



Détection et classification de signatures temporelles CAN pour l'aide à la maintenance de sous-systèmes d'un véhicule de transport collectif

Nicolas Cheifetz

► To cite this version:

Nicolas Cheifetz. Détection et classification de signatures temporelles CAN pour l'aide à la maintenance de sous-systèmes d'un véhicule de transport collectif. Autre. Université Paris-Est, 2013. Français. NNT : 2013PEST1077 . tel-01066894

HAL Id: tel-01066894

<https://theses.hal.science/tel-01066894>

Submitted on 22 Sep 2014

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

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



IFSTTAR

INSTITUT FRANÇAIS
DES SCIENCES
ET TECHNOLOGIES
DES TRANSPORTS,
DE L'AMÉNAGEMENT
ET DES RÉSEAUX



Thèse de doctorat

Soutenue le 9 septembre 2013

Détection et classification de signatures temporelles CAN pour l'aide à la maintenance de sous-systèmes d'un véhicule de transport collectif

par **Nicolas CHEIFETZ**

en vue de l'obtention du titre de docteur de
l'**Université Paris-Est** dans le cadre de
l'**école doctorale n° 532 – MSTIC**
Convention CIFRE n° 2009 / 592

Structure de recherche d'accueil : GRETTIA



UNIVERSITÉ PARIS-EST
ÉCOLE DOCTORALE MSTIC
MATHÉMATIQUES ET SCIENCES ET TECHNOLOGIES
DE L'INFORMATION ET DE LA COMMUNICATION

THÈSE

pour obtenir le titre de

Docteur en Sciences

de l'Université Paris-Est

Spécialité : SIGNAL, IMAGE, AUTOMATIQUE

Présentée et soutenue par

Nicolas CHEIFETZ

Détection et classification de signatures temporelles CAN pour l'aide à la maintenance de sous-systèmes d'un véhicule de transport collectif

Thèse préparée à l'IFSTTAR, laboratoire GRETTIA
et Veolia Environnement Recherche & Innovation

soutenue publiquement le 9 septembre 2013 devant le jury composé de :

M. Patrice AKNIN <i>Directeur de Recherche à l'IFSTTAR</i>	Directeur de thèse
M. Jean-Laurent FRANCHINEAU <i>Directeur de Programme à l'IEED VeDeCoM</i>	Invité
Mme. Nadine MARTIN <i>Directrice de Recherche CNRS à Gipsa-Lab</i>	Examinatrice
M. Mohamed NADIF <i>Professeur à l'Université Paris Descartes</i>	Rapporteur
M. Igor NIKIFOROV <i>Professeur à l'Université de Technologie de Troyes (UTT)</i>	Rapporteur
M. Allou SAMÉ <i>Chargé de Recherche à l'IFSTTAR</i>	Encadrant de thèse

Détection et Classification de Signatures Temporelles CAN pour l'Aide à la Maintenance de Sous-Systèmes d'un Véhicule de Transport Collectif

Copyright © Nicolas CHEIFETZ, Université Paris-Est, IFSTTAR-GRETTIA, Veolia Environnement Recherche & Innovation

L'Université Paris-Est, l'IFSTTAR-GRETTIA et Veolia Environnement Recherche & Innovation ont le droit, perpétuel et sans limites géographiques, d'archiver et publier cette thèse grâce à des copies reproduites sur papier ou sous format numérique, ou par tout autre moyen connu ou à être inventé, et exclusivement à des fins non lucratives de recherche et d'enseignement. L'auteur et les coauteurs le cas échéant conservent la propriété du droit d'auteur et des droits moraux qui protègent ce document.

Manuscrit mis en page avec $\text{\LaTeX}2_{\epsilon}$ à partir d'une version personnalisée du template *ThesisDIFCTUNL*.

A Yasmina et mes parents.

Remerciements

Tout d’abord, je remercie vivement M. Patrice AKNIN et M. Allou SAMÉ, respectivement directeur et encadrant de thèse, pour m’avoir initié aux problèmes de diagnostic et de maintenance liés aux systèmes de transport au sein de l’INRETS devenu IFSTTAR. La confiance et les conseils rigoureux qu’ils m’ont donnés, ont su me guider jusqu’à l’aboutissement de ce travail. Un très grand merci à Allou pour son soutien scientifique, son infinie patience et la disponibilité dont il a fait preuve tout au long de ma thèse.

J’ai le plaisir de remercier tous les membres du jury pour m’avoir fait l’honneur de participer à l’évaluation de mes travaux de thèse. Je remercie notamment M. Igor NIKIFOROV et M. Mohamed NADIF, qui ont accepté de rapporter ces travaux, pour leur lecture attentive du manuscrit et leurs observations toujours constructives. Je remercie également Mme Nadine MARTIN, examinatrice et présidente de ce jury, pour l’attention et l’intérêt qu’elle a porté à mes travaux. Aussi, je remercie M. Jean-Laurent FRANCHI-NEAU, invité du jury, et M. Emmanuel DE VERDALLE pour avoir lancé ces travaux au sein de l’entité recherche de Veolia (VERI) et m’avoir donné la chance de découvrir la recherche en entreprise. Malgré la cessation des activités de transport à VERI, ils ont continué à suivre mon travail et soutenir mes efforts de jeune chercheur. Cette recherche n’a pas toujours été un long fleuve tranquille et je tiens à remercier M. Damien CHENU de VERI pour avoir su me diriger en eau trouble ainsi que pour sa relecture minutieuse et vaillante du manuscrit.

Plus généralement, je remercie l’ensemble de mes collègues et des partenaires impliqués dans les travaux du projet européen EBSF qui a permis de financer une partie de mes travaux. Je pense notamment à l’équipe de maintenance au réseau d’autobus de la STRAV (en Ile-de-France) et en particulier, M. Matthieu OMBRABELLA et M. Manuel RODRIGUES pour m’avoir soutenu dans mes investigations, ainsi que M. Michaël BEDROSSIAN pour m’avoir pris sous son aile au début de la thèse. Mes remerciements vont également à M. Mohamed BRACI et Mme Lydie NOUVELIÈRE pour les discussions que nous avons pu avoir sur la modélisation longitudinale d’un véhicule. Ayant eu l’occasion de présenter plusieurs fois une partie de mes travaux, je remercie tous les chercheurs avec lesquels j’ai pu échanger en conférences et autres séminaires.

Je souhaite maintenant remercier chacun de mes collègues bien qu'ils n'aient pas été directement impliqués dans ce travail de recherche. Je m'excuse par avance de ne pas tous les citer. Je remercie tous mes collègues d'enseignement et ceux qui ont accompagné mes années au sein du laboratoire GRETTIA et en particulier Latifa, Étienne, Olivier, Laurent, et également Annie et Mustapha. Je tiens à remercier amicalement, pour les bons moments partagés, mes collègues thésards réunis sous la bannière du « RaDium » au premier étage du nouveau bâtiment Bienvenue : les anciens réticents Wissam, Hani, Johanna et Rony, les jeunes convertis Josquin, Dihya et Andry, ainsi que les nouveaux anciens Guillaume et Manuel L. P. Les « Hey » matinaux me manqueront, c'est certain ! Je remercie également les anciens du labo Inès, Faicel, Roland, Zohra, Raïssa et Cyril (Spiderman est resté à mes côtés !) tout comme les stagiaires, en particulier Nadir. Je remercie notamment Sébastien D. pour l'instrumentation pneumatique. En tant que thésard CIFRE, j'ai pu profiter d'un second environnement de recherche, en entreprise, et ainsi rencontrer de nombreux chercheurs avec des problématiques scientifiques qui m'étaient étrangères. Merci à M. Philippe DOYEN de m'avoir accueilli dans son équipe au début de ma thèse. Je n'oublierai jamais mes collègues de VERI, notamment le premier étage de Rueil et nos matchs rituels de mots fléchés après déjeuner. Je remercie sincèrement mes co-bureaux du bureau 116, Karl pour ses réparations innombrables de mes casques audio et Philippe B. pour sa grande curiosité scientifique. Merci au thésard plein d'avenir Vincent S. Enfin, je remercie tous mes collègues du pôle CAD pour les nombreux échanges savants et la bonne ambiance à chaque réunion !

Je terminerai ce préambule en remerciant ma famille et mes amis qui m'ont supporté pendant toutes ces années et ont su me faire lâcher prise lorsqu'il le fallait. Mes remerciements vont particulièrement à mes parents pour leurs encouragements infaillibles durant toutes mes années. Je remercie enfin Yasmina, sans laquelle ces travaux n'auraient pu converger, pour son soutien constant et à qui cet ouvrage est également dédié.

Nicolas Cheifetz
Paris, le 17 septembre 2013

Résumé

De nos jours, les systèmes de transport ont la nécessité d'être exploités à un haut niveau de sécurité et de fiabilité. On constate notamment un besoin croissant d'outils de surveillance et d'aide à la maintenance préventive de certains sous-systèmes critiques. Menée dans le cadre du projet EBSF (European Bus System of the Future, FP7), cette thèse aborde la problématique de surveillance de deux sous-systèmes d'autobus impactant fortement la disponibilité des véhicules et leurs coûts de maintenance : le système de freinage et les portes d'accès. L'approche à base de reconnaissance des formes, qui s'appuie sur l'analyse de données collectées en exploitation, est privilégiée dans ces travaux de thèse. La méthodologie de détection appliquée aux freins s'appuie sur une modélisation dynamique du véhicule afin d'extraire des descripteurs physiques liés au freinage. Des cartes de contrôle multivariées opérant sur ces descripteurs sont alors employées afin de détecter toute anomalie. La stratégie de détection appliquée aux portes traite directement les courbes, qui sont collectées par des capteurs embarqués pendant des cycles d'ouverture et fermeture. On propose un test séquentiel du rapport de vraisemblance généralisé, qui repose sur un modèle paramétrique de régression à segmentation latente, pour résoudre ce problème de détection. Une approche d'aide au diagnostic à base de détecteurs locaux est également mise en œuvre. Les résultats obtenus à partir de données réelles collectées sur une flotte de bus ont permis de mettre en évidence l'efficacité des méthodes proposées aussi bien pour l'étude des freins que celle des portes.

Mots-clés: Diagnostic industriel, détection de changement, cartes de contrôle, tests séquentiels d'hypothèses, segmentation de courbes, modèles de mélange fini, algorithme EM, modélisation longitudinale d'un véhicule, système de freinage, portes d'autobus.

Table des matières

Table des matières	xi
Notations et Abréviations	xv
1 Introduction	1
1.1 Introduction générale	2
1.1.1 Contexte et problématique	2
1.1.2 Positionnement et objectifs de la thèse	2
1.1.3 Organisation du manuscrit	4
1.2 Contexte et enjeux industriels	6
1.2.1 Collaboration industrielle	6
1.2.2 Projet européen EBSF	6
1.3 Cadre expérimental	7
1.3.1 Étude et objectifs	7
1.3.2 Architecture télématique et instrumentation	8
1.3.3 Spécification des données	11
1.3.4 Agrégation des connaissances réelles d'exploitation	12
1.4 Étude bibliographique	14
1.4.1 Système de freinage	14
1.4.2 Système des portes	18
1.5 Conclusion	20
2 Méthodes de diagnostic par reconnaissance des formes intégrées aux processus de maintenance	21
2.1 Introduction	22
2.1.1 Stratégies de maintenance	22
2.1.2 Principales approches de diagnostic de systèmes industriels	24
2.2 Diagnostic à base de reconnaissance des formes	26
2.2.1 Démarche générique	26
2.2.2 Apprentissage statistique	28

2.2.3	Les modèles de mélange de lois et leur estimation	35
2.3	Diagnostic à partir de données fonctionnelles	40
2.3.1	Classification de données fonctionnelles	41
2.3.2	Modèle de mélange de régressions pour la modélisation de courbes à changement de régimes	42
2.3.3	Cas d'étude réel	49
2.4	Conclusion	52
3	Détection statistique de défaillances	53
3.1	Formalisation du problème de la détection de points de changement . . .	54
3.1.1	Quelques généralités sur les tests d'hypothèses	55
3.1.2	Détection séquentielle de changements sous forme de test d'hypo- thèses	56
3.2	Algorithmes classiques de détection à base de test d'hypothèses	58
3.2.1	Quelques tests pour le problème de détection séquentielle	58
3.2.2	Résultats classiques d'optimalité d'un détecteur	64
3.2.3	Contrainte de fausses alarmes et estimation du seuil de détection .	67
3.2.4	Évaluation d'un détecteur	70
3.3	Cartes de contrôle et autres tests séquentiels	72
3.3.1	Principales cartes de contrôle multivariées	72
3.3.2	Quelques extensions à base de test séquentiel	78
3.4	Détection à base de reconnaissance des formes	80
3.4.1	Détection d'anomalies et classification mono-classe	80
3.4.2	Détection de changements dans des modèles markoviens	82
3.4.3	Détection de changements par la sélection de modèles	84
3.4.4	Détection de changements dans une séquence de données fonction- nelles	85
3.5	Conclusion	87
4	Détection d'anomalies à l'aide d'un modèle physique sur un système de frein- nage	89
4.1	Introduction	90
4.1.1	Différents systèmes de freinage	90
4.1.2	Description du système de freinage étudié	91
4.1.3	Démarche des travaux scientifiques	93
4.2	Modélisation dynamique du véhicule et des freins	94
4.2.1	Modélisation longitudinale du bus	94
4.2.2	Estimation de la pente	102
4.2.3	Représentation du fonctionnement normal du système de freinage	104
4.3	Stratégie séquentielle de détection d'anomalies	106
4.3.1	Description globale	106

4.3.2	Détection de défauts	107
4.4	Expérimentations sur données réelles	108
4.4.1	Acquisition des données de fonctionnement	108
4.4.2	Description des données réelles collectées	110
4.4.3	Dégradations simulées du système de freinage	113
4.4.4	Résultats expérimentaux et estimation du seuil de détection	116
4.5	Conclusion	119
5	Surveillance des portes d'accès à partir d'une séquence de courbes	121
5.1	Introduction	122
5.1.1	Différents systèmes de porte	122
5.1.2	Description du système de porte étudié	123
5.1.3	Démarche des travaux scientifiques	125
5.2	Un modèle de régression pour des courbes multivariées	126
5.2.1	Modèle de régression multivarié à processus caché	126
5.2.2	Cas d'étude simulé dans un cadre non supervisé	132
5.2.3	Apprentissage semi-supervisé des paramètres	136
5.2.4	Cas d'étude simulé dans un cadre semi-supervisé	138
5.3	Stratégie de détection séquentielle dans un flux de courbes	140
5.3.1	Rappel de quelques règles de détection en ligne	140
5.3.2	Approche globale : détection d'anomalies	142
5.3.3	Approche locale : localisation des défauts sur la signature	144
5.3.4	Estimation des seuils de détection et de localisation	146
5.4	Évaluation du détecteur sur données simulées réalistes	149
5.4.1	Travaux expérimentaux	149
5.4.2	Réglage de la structure du modèle de régression	150
5.4.3	Résultats expérimentaux	153
5.5	Expérimentations sur données réelles	156
5.5.1	Acquisition des données de fonctionnement	156
5.5.2	Description des données réelles collectées	156
5.5.3	A priori physique et sélection du modèle	159
5.5.4	Interprétation des défauts réels	161
5.5.5	Résultats en terme de détection, localisation et interprétation . . .	163
5.6	Conclusion	168
6	Conclusions et perspectives	171
	Annexes	177
A	Rappel de quelques outils mathématiques	177
A.1	Quelques définitions	178
A.2	Quelques opérateurs et inégalités	180

Bibliographie	181
Liste des figures	197
Liste des tableaux	201
Liste des algorithmes	203
Liste des publications	205

Notations et Abréviations

Estimation paramétrique et Détection séquentielle

d	Dimension du vecteur d'observation ;
$\mathbf{x} \in \mathbb{R}^d$	Vecteur d'observation ;
$\mathbf{x}^T$	Transposée du vecteur d'observation ;
$\mathbf{x} = (x_1, \dots, x_n)$	Vecteur de n observations ;
$\mathbf{X} = (X_1, \dots, X_n)$	Vecteur aléatoire dont $\mathbf{x}$ est une réalisation ;
$\boldsymbol{\theta} \in \Theta$	Vecteur paramètre à valeur dans le domaine de définition Θ ;
$\mathcal{P}_\theta$	Distribution/loi de probabilité paramétrée par θ ;
$P(\mathbf{x}; \theta)$	Fonction de masse (probabilité) de $\mathbf{X}$ associée à la distribution $\mathcal{P}_\theta$;
$p(\mathbf{x}; \theta)$	Densité de probabilité de la variable $\mathbf{X}$ associée à la distribution $\mathcal{P}_\theta$;
$\mathbb{E}(\mathbf{X}; \theta)$	Espérance mathématique de $\mathbf{X}$ associée à la distribution $\mathcal{P}_\theta$;
$\text{Var}(\mathbf{X}; \theta)$	Matrice de variance-covariance de $\mathbf{X}$ associée à la distribution $\mathcal{P}_\theta$;
$L(\theta; \mathbf{x})$	Vraisemblance de θ connaissant $\mathbf{x}$;
$\mathcal{L}(\theta; \mathbf{x})$	Log-vraisemblance de θ connaissant $\mathbf{x}$;
$\mathcal{L}_c(\theta; \mathbf{x})$	Log-vraisemblance complétée de θ connaissant $\mathbf{x}$;
$\mathcal{N}(\cdot; \mu, \Sigma)$	Loi normale/gaussienne de moyenne μ et matrice de covariance Σ ;
$\mathcal{M}(\cdot; \pi)$	Loi multinomiale de paramètre π ;
$\mathcal{H}$	Hypothèse dans un test statistique ;
t_c	Temps de changement, point de rupture ;
t_a	Temps d'alarme (signalé par un détecteur) ;
FAR ou α	<i>False Alarm Rate</i> , ou taux de fausses alarmes ;
$\chi_d^2(\alpha)$	Quantile d'ordre $(1 - \alpha)$ de la loi de χ^2 avec d degrés de liberté ;
$F_{d,e}(\alpha)$	Quantile d'ordre $(1 - \alpha)$ de la loi de Fisher avec (d, e) degrés de liberté ;
$B_{d,e}(\alpha)$	Quantile d'ordre $(1 - \alpha)$ de la loi de Beta avec (d, e) degrés de liberté ;
ARL ou ARL2FA	<i>Average Run Length</i> , ou temps moyen entre deux fausses alarmes ;
ADD	<i>Average Delay of Detection</i> , ou retard moyen à la détection ;
LCL	<i>Lower Control Limit</i> , ou limite de contrôle inférieure ;
UCL	<i>Upper Control Limit</i> , ou limite de contrôle supérieure .

p.s.	Presque sûrement ;
$f \sim g \text{ qd } x \mapsto \infty$	La fonction f est asymptotiquement équivalente à la fonction g : $\lim_{x \mapsto \infty} \frac{g(x)}{f(x)} = 1 ;$
$f(x) = \mathcal{O}(g(x))$	La fonction f est asymptotiquement dominée par la fonction g : $\exists N, C : \forall x > N, f(x) \leq C \cdot g(x) ;$
$f(x) = o(g(x))$	La fonction f est asymptotiquement négligeable devant la fonction g : $\lim_{x \mapsto \infty} \frac{f(x)}{g(x)} = 0 .$

Acronymes projet

EBSF	Projet European Bus System of the Future (EBSF Consortium, 2012) ;
UC de Brunoy	<i>Use Case</i> , ou cas d'étude du projet EBSF dédié au télédiagnostic ;
CAN	<i>Controller Area Network</i> , bus de communication dans des véhicules ;
SAE J1939	Protocole de communication et spécifications de données implémentables sur le CAN (SAE Technical Standards, 2006) ;
SPN	<i>Suspect Parameter Number</i> , donnée spécifiée dans le J1939 ;
PGN	<i>Parameter Group Number</i> , groupe d'un ou plusieurs SPN ;
bus-FMS	<i>bus Fleet Management System</i> , sous-ensemble des données du J1939. Les six principaux constructeurs européens d'autobus ont défini l'interface bus-FMS en 2009 dans sa seconde version ;
ATON TM	Calculateur embarqué dans les autobus, illustré par la figure 1.3, propriété de Veolia Environnement. Les données associées au système de freinage sont collectées sur ce calculateur ;
VIDAC III TM	Calculateur embarqué dans les autobus, illustré par la figure 1.3, propriété de DIGIGROUP. Les données associées aux portes sont collectées sur ce calculateur.

Autres notations

La notation des variables physiques concernant les travaux sur le système de freinage est définie dans le tableau 4.4. Nous avons fait le choix de les mentionner au début du chapitre 4 car cette notation n'apparaît que dans ce chapitre là.

Introduction

Sommaire

1.1	Introduction générale	2
1.1.1	Contexte et problématique	2
1.1.2	Positionnement et objectifs de la thèse	2
1.1.3	Organisation du manuscrit	4
1.2	Contexte et enjeux industriels	6
1.2.1	Collaboration industrielle	6
1.2.2	Projet européen EBSF	6
1.3	Cadre expérimental	7
1.3.1	Étude et objectifs	7
1.3.2	Architecture télématique et instrumentation	8
1.3.3	Spécification des données	11
1.3.4	Agrégation des connaissances réelles d'exploitation	12
1.4	Étude bibliographique	14
1.4.1	Système de freinage	14
1.4.2	Système des portes	18
1.5	Conclusion	20

1.1 Introduction générale

1.1.1 Contexte et problématique

De nos jours, les opérateurs de transport sont confrontés à de fortes contraintes liées à la qualité de leur offre de services. En outre, les systèmes de transport ont la nécessité d'être exploités avec un haut niveau de sécurité et de fiabilité. Pour répondre à l'attente des collectivités et à celle des passagers, chaque opérateur doit mettre en place une stratégie capable de maintenir en état fonctionnel les systèmes de transport mis en exploitation. Il est également nécessaire de pouvoir garantir leur adaptation face à une évolution des conditions opérationnelles. Une stratégie de maintenance trop laxiste des systèmes entraîne des coûts de maintenance prohibitifs et implique des conséquences néfastes à l'encontre de certains sous-systèmes critiques. A l'inverse, des actions de maintenance systématiques peuvent se révéler inutiles et s'avèrent chronophages pour la disponibilité des véhicules. Une politique de maintenance efficace s'appuie sur une surveillance préventive des sous-systèmes et repose généralement sur l'analyse d'*indicateurs* (AFNOR, 2007) représentant l'état de fonctionnement de ces sous-systèmes (voir chapitre 2).

Les sous-systèmes industriels susceptibles d'être les plus surveillés sont ceux qui présentent de fortes exigences de sécurité et une faible *fiabilité intrinsèque* (AFNOR, 2010). Dans un autobus, le système de freinage et les portes correspondent à ce type de description. Une défaillance sur l'un de ces organes critiques entraîne l'immobilisation immédiate du véhicule ce qui peut impacter la disponibilité des véhicules sur l'ensemble du réseau. De surcroît, l'image de l'opérateur de transport peut être affectée lorsque les usagers sont invités à quitter le véhicule immobilisé. Aujourd'hui, la maintenance des freins et des portes de bus se base généralement sur les recommandations d'entretien des constructeurs et sur l'expérience des agents de maintenance. Dans ce cas, l'émergence de défauts n'est pas toujours détectable et ni la signature des défauts ni leur gravité ne peuvent être caractérisées systématiquement.

1.1.2 Positionnement et objectifs de la thèse

Prenant pour terrain d'essais un réseau d'autobus en Ile-de-France, l'objectif visé dans ce travail de thèse est la mise au point de méthodes innovantes pour la surveillance de sous-systèmes critiques d'autobus. On s'intéresse au suivi du mode de fonctionnement de sous-systèmes dont l'évolution inhérente est non déterministe. Les sous-systèmes étudiés sont supposés évoluer d'un état à un autre par des changements plus ou moins rapides. Les méthodes proposées doivent permettre de détecter un changement brusque du modèle représentant le sous-système, ce qui signifierait le passage d'un état de fonctionnement à un autre. On s'intéresse notamment à la détection de *changements structurels* (Davis et Franke, 2008) dans des modèles génératifs paramétriques (voir chapitre 2).

A titre d'illustration, la figure 1.1 représente l'évolution de deux processus simulés au cours du temps avec un changement structurel. La série temporelle de la figure 1.1(a) décrit une séquence d'observations univariées avec un changement de moyenne au temps 5 001 et chacune de ces observations est la réalisation d'une variable aléatoire de $\mathbb{R}$. Dans beaucoup de systèmes réels, plusieurs indicateurs sont mesurés simultanément et le problème devient multivarié. La série temporelle de la figure 1.1(b) décrit une séquence de courbes bidimensionnelles avec un changement au temps 51. Ainsi, chaque donnée correspond à un vecteur de dimension 2×300 dans notre exemple. Il est à noter que la première séquence est composée de 10 000 points de mesure et la seconde est composée de 60 000 ($= 2 \times 300 \times 100$) points.

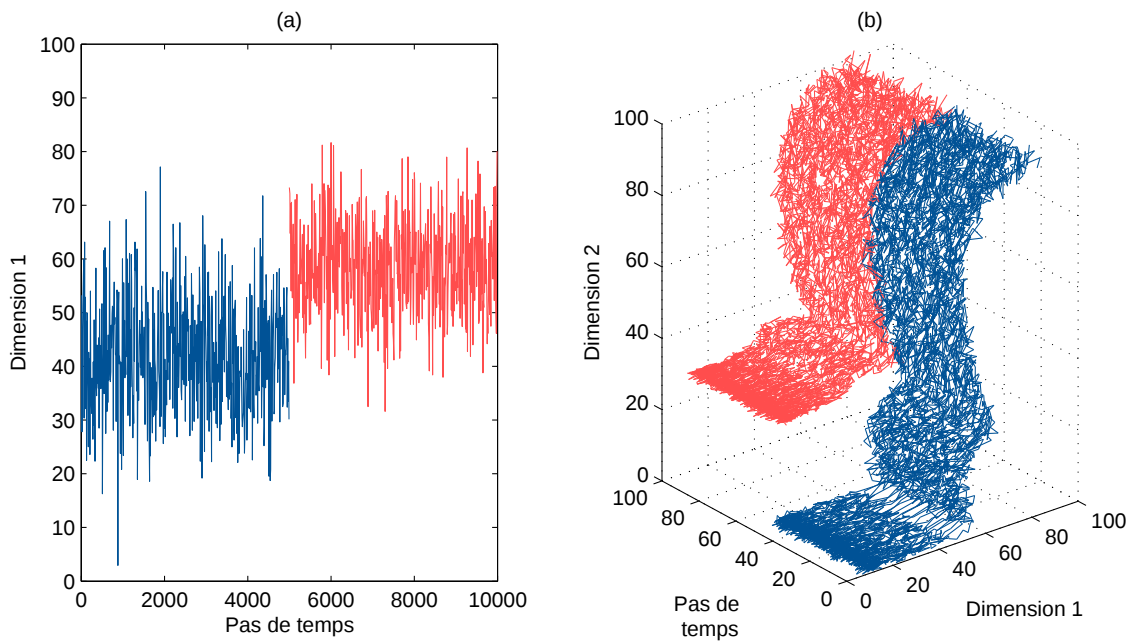


FIGURE 1.1 – Évolution de deux processus au cours du temps avec un changement structurel : dans une séquence de données unidimensionnelles (a) et dans une séquence de données fonctionnelles bivariées (b).

La spécificité de nos travaux réside dans l'approche macroscopique et générique adoptée pour le développement d'outils d'aide à la maintenance. Les outils proposés dans ce mémoire sont dédiés à la surveillance des systèmes de freinage et de porte d'autobus urbains en exploitation. Pour le système de freinage, on propose une méthodologie de diagnostic à base de modèle qui s'appuie d'une part sur un modèle de la dynamique des bus qui permet de générer des variables physiquement interprétables en sortie et, d'autre part sur une stratégie de détection opérant sur ces dernières variables. La méthode utilisée pour la surveillance des portes n'exploite pas de modèle physique mais opère directement sur des courbes mesurées lors de cycles consécutifs d'ouverture/fermeture. La nature fonctionnelle des données nous a conduits à proposer une méthode de

détection couplée à un modèle de régression adapté à ces courbes selon une approche de diagnostic par reconnaissance des formes.

Afin de construire une stratégie de diagnostic, les premières étapes de collecte et de mise en forme des données sont décrites pour nos travaux. Même si ces travaux ne font pas l'objet de développements théoriques, nous les considérons comme une contribution. En effet, avant de pouvoir valider les modèles proposés dans cette thèse et d'analyser les données réelles, il a fallu construire des bases de données (mesurées sur chaque système) et de connaissances expertes (échanges avec l'équipe de maintenance) avec « le plus d'exhaustivité possible ».

Les contributions opérationnelles de ce travail de thèse visent à répondre à quelques problèmes pratiques dont certains sont listés ci-après :

- Acquérir une expertise sur des données nouvellement collectées, ces données n'ayant été que très peu analysées par l'exploitant du réseau jusqu'à présent ;
- Proposer des méthodes de surveillance sensibles à des changements structurels mais robustes face à la variabilité des processus observés en condition nominale ;
- Surveiller automatiquement un système pendant son fonctionnement, sans interventions de contrôle contraignantes ni destructives ;
- Développer des outils qui soient facilement implémentables sur des calculateurs (embarqués ou au dépôt), exécutables en temps quasi réel et facilement interprétables par les équipes de maintenance.

1.1.3 Organisation du manuscrit

Ce mémoire commence, en **chapitre 1**, par l'introduction du contexte industriel et notamment l'introduction d'un cas d'étude sur lequel porte ce travail de thèse. Ce premier chapitre se termine par une étude bibliographique des méthodes actuelles pour la modélisation d'un système de freinage d'une part, et celle d'un système de porte d'autre part. Le reste du manuscrit est décomposable en deux chapitres état de l'art suivis de deux chapitres centrés sur nos contributions pour la surveillance des systèmes de freinage et de porte.

Le **chapitre 2** présente les principales stratégies de diagnostic d'un système industriel, et notamment l'approche de diagnostic à base de reconnaissance des formes dans le cadre d'une maintenance conditionnelle. L'objectif de cette approche est de surveiller régulièrement l'état de fonctionnement d'un système sur la base de tests non destructifs. Le diagnostic s'appuie sur l'apprentissage d'un modèle statistique à partir d'une base historisée de données mesurées sur le système à surveiller. Les principaux paradigmes associés à l'estimation de ce type de modèle sont passés en revue, et on décrit notamment les modèles de mélange de lois, adaptés pour la modélisation de données observées peu

supervisées dans des groupes homogènes. Dans ce cadre, le problème de la classification de données fonctionnelles (courbes) est posé. Un mélange de régression adapté à ce type de données est spécifié puis illustré.

Dans nos travaux, on s'intéresse particulièrement à l'étape de détection, phase capitale et préalable au processus de diagnostic introduit au chapitre précédent. Le **chapitre 3** dresse un état de l'art des méthodes en détection séquentielle dont l'objectif est de signaler le plus tôt possible un changement parmi un flux de données observées. On détaille notamment les méthodes séquentielles à base de tests d'hypothèses et leurs propriétés d'optimalité. En particulier, les méthodes classiques dites de cartes de contrôle sont décrites pour des données multidimensionnelles qui seront employées pour le travail sur les freins. Ce chapitre se termine par la description d'une liste de problèmes et méthodes liées à la détection de changement et on s'intéresse notamment à cette problématique dans le cadre d'une séquence de courbes.

Le **chapitre 4** introduit notre première contribution dans ce travail de thèse. Ce chapitre présente une méthodologie pour la surveillance séquentielle d'un système de freinage pneumatique basée sur une modélisation dynamique du véhicule. On extrait des indicateurs à partir de cette modélisation longitudinale pendant des phases de freinage. Plusieurs méthodes de détection sont testées sur la base d'une modélisation de ces indicateurs physiques. Dans un contexte où les situations dégradées n'étaient pas disponibles, l'extraction d'indicateurs physiquement interprétables nous a permis de développer les détecteurs dans un espace de représentation pertinent. Les résultats obtenus sur des données réelles d'exploitation avec dégradations simulées sont présentés, ainsi qu'une stratégie d'estimation des seuils de détection.

Le **chapitre 5** est consacré au développement d'une approche de surveillance des portes, ce qui constitue notre seconde contribution dans ce travail de thèse. Plus particulièrement, on propose une stratégie de surveillance séquentielle dans un flux de courbes multidimensionnelles. Ce travail est motivé par la forme des données mesurées sur les systèmes de portes pneumatiques d'autobus. Chaque courbe correspond à une séquence de pressions et de positions, pendant une ouverture ou une fermeture de porte. Pour cela, on propose un nouveau modèle paramétrique de régression adapté à des données fonctionnelles multivariées. On décline alors une stratégie séquentielle à base d'hypothèses qui s'appuie sur le modèle de régression évoqué afin de détecter tout comportement anormal dans le flux de courbes observées. On propose également la formulation de statistiques de localisation afin d'aider à caractériser un défaut détecté. Cette méthodologie est testée sur des courbes simulées, et sur des données réelles collectées en exploitation.

Enfin, le manuscrit s'achève, dans le **chapitre 6**, par nos remarques de conclusion et explore quelques directions futures de recherche.

1.2 Contexte et enjeux industriels

1.2.1 Collaboration industrielle

Ce travail s'inscrit dans le cadre d'une collaboration entre Veolia Environnement Recherche & Innovation (VERI) et le laboratoire GRETTIA de l'Institut Français des Sciences et Technologies des Transports, de l'Aménagement et des Réseaux (IFSTTAR). Cette collaboration porte sur la surveillance de deux sous-systèmes critiques d'autobus, et la thèse a obtenu un financement CIFRE de la part de l'ANRT. Les travaux ont été déployés au sein d'un réseau de transport collectif d'Ile de France exploité par Transdev¹ (anciennement Veolia Transdev) – 1^{er} opérateur privé mondial de transport public.

1.2.2 Projet européen EBSF

La thèse s'inscrit dans le cadre du **projet européen EBSF**². Sous l'égide de l'Union Européenne et coordonné par l'UITP³, EBSF fait partie du 7^e Programme Cadre de Recherche et Développement (PCRD). Le projet a débuté en 2008 pour une durée de quatre ans. Un consortium de 47 partenaires, répartis sur dix pays, participe au projet dont les 5 principaux constructeurs, équipementiers, opérateurs et centres de recherche européens du secteur.



Aujourd'hui, 80% des passagers de transport public à travers le monde voyagent en bus⁴. Toutefois, le bus est souvent perçu comme un système moins « attrayant » que le tramway urbain moderne ou le chemin de fer métropolitain sur des critères de design, confort, accessibilité ou encore information voyageur. Il devient donc indispensable d'accélérer la modernisation ou la « renaissance » de l'autobus. En réponse à cette problématique, le projet EBSF vise à valoriser et promouvoir l'image de l'autobus européen.

Ce projet d'envergure apporte une nouvelle vision du « système bus » en proposant des innovations sur le véhicule, l'infrastructure et l'exploitation. L'intérêt de participer aux recommandations des prochaines générations de bus est évident pour un opérateur de transport public comme Transdev. En 2011, le bus urbain était déjà le premier mode de transport géré et exploité par l'opérateur. La tendance ne devrait pas changer puisque selon l'OMS, 60% de la population mondiale vivra dans les villes en 2030. En outre, un bus reste moins coûteux que le tramway et, atout supplémentaire, les délais de livraison de l'infrastructure sont plus courts.

1. Transdev : <http://www.transdev.net/>.

2. European Bus System of the Future : <http://www.ebsf.eu/>.

3. Union internationale des transports publics : <http://www.uitp.org/>.

4. UITP – Bus : <http://www.uitp.org/public-transport/bus/>.

Remarque 1.1 (Bus à Haut Niveau de Service – BHNS)

*En France, le tramway est réintroduit dès les années 80 avec le soutien de l'État. Mais il apparaît rapidement qu'une solution hybride entre tramway et autobus soit préférable tant en terme économique que des besoins opérationnels. Le concept de **BHNS**, né en 2005, répond à cette attente. A l'image du réseau TEOR à Rouen (exploité par Transdev), un réseau de BHNS est un mode de transport collectif en site propre urbain notamment caractérisé par de fortes exigences sur la vitesse, le confort, la régularité, la fréquence de passage, l'accessibilité et l'information aux voyageurs (Certu, 2009). Dans nos travaux, les bus ne sont pas des BHNS; toutefois, on cherche à ajouter de nouvelles fonctionnalités aux bus afin d'améliorer les actions de maintenance.*

Une partie importante du projet EBSF est dédiée à la mise en pratique de certaines innovations dont le design de nouveaux véhicules⁵ ou encore un nouveau concept de station de bus⁶ « Osmose » installée depuis mai 2012 à Paris. Un cas d'étude du projet vise à mettre au point des outils de diagnostic et d'aide à la maintenance des bus par des échanges télématiques ; on parle de **télédiagnostic**. Le réseau urbain de la STRAV⁷ (Transdev) a été choisi pour mener ce travail sur les systèmes freins et portes. Dans le reste du manuscrit, on désignera ce cas d'étude par le **UseCase de Brunoy**.

1.3 Cadre expérimental

1.3.1 Étude et objectifs

Le réseau de la STRAV, situé dans le Val de Marne, a mis à disposition son infrastructure et son expertise en tant qu'exploitant pour mener les expérimentations du UseCase de Brunoy. L'étude s'appuie sur une flotte de dix autobus Citelis (moteur diesel respectant la norme d'émission Euro 4) du constructeur IRISBUS-IVECO. La flotte, composée de six bus articulés (18m) et quatre bus standards (12m), a été équipée d'une instrumentation dédiée. Afin de ne pas perdre de vue le besoin de faisabilité, les dix véhicules sont restés opérationnels et exploités normalement sur toute la durée de l'étude sur les lignes J1 et J2 (trajet de 15km, 38 arrêts, forte topographie, plus grande affluence du réseau).

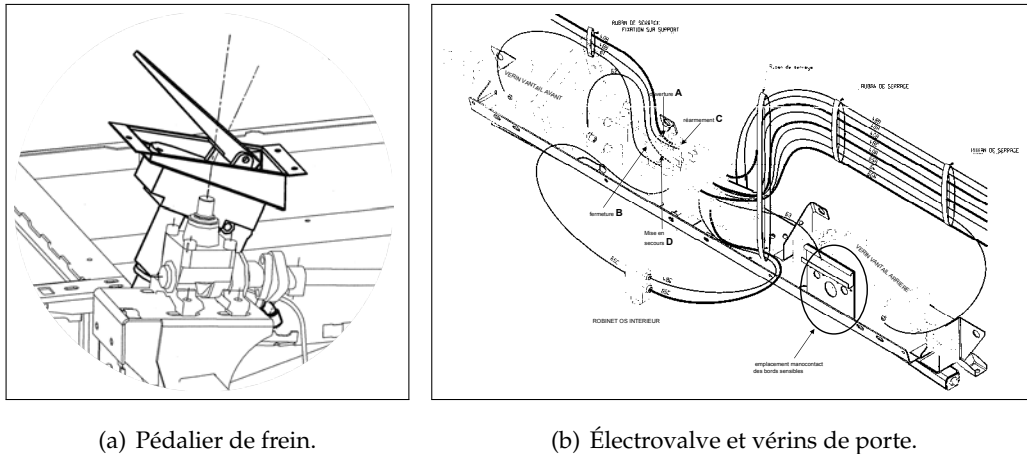
Suite à la conception des systèmes mécaniques, chaque constructeur émet des préconisations pour l'entretien de ses systèmes. Ces recommandations théoriques sont relatives à un fonctionnement plus ou moins standard. Un opérateur de transport opère des véhicules dans des conditions particulières liées aux besoins d'exploitation. Dans cette étude, on propose des solutions adaptées à de telles conditions et toujours complémentaires aux préconisations des constructeurs. A long terme, l'intérêt pour l'exploitant est double : améliorer la disponibilité des véhicules et limiter leurs pannes en service.

5. Prototype Volvo à Göteborg : <http://www.ebsf.eu/index.php/project-development/use-cases/72-gothenburg>.

6. http://www.ratp.fr/fr/ratp/r_65980/osmose-quelle-station-de-bus-pour-demain/.

7. Société de Transports Automobiles et de Voyages : <http://www.idf.veolia-transport.fr/reseau-bus-strav/>.

Dans le cadre de cette étude, l'objectif est de mettre en place des outils pour la surveillance de deux sous-systèmes pneumatiques : le système de freinage et les portes. Ces organes impactent les temps de disponibilité des bus et les coûts de maintenance. La figure 1.2 illustre ces deux sous-systèmes par un pédalier de frein (a) et le schéma de la partie supérieure d'une porte (b). L'analyse fine de données collectées via des capteurs cherche à suivre l'état de santé des systèmes et à détecter toute dérive de fonctionnement.



(a) Pédalier de frein.

(b) Électrovalve et vérins de porte.

FIGURE 1.2 – Deux actionneurs pneumatiques : frein (a) et porte (b).

A ce jour, les seules actions de maintenance menées au dépôt de bus sont *systématiques* (après une certaine période temporelle et/ou un certain kilométrage des bus) ou *correctives* (après la défaillance d'un système). Il n'existe pas encore de maintenance préventive *conditionnelle* / *prévisionnelle* au dépôt. Un diagramme des différents modes de maintenance est présenté dans la partie 2.1. Les outils proposés dans cette thèse cherchent à aider les agents de maintenance à prévenir des défauts et à détecter tout comportement anormal. Il s'agit bien de compléter la politique actuelle de maintenance et non de se substituer aux agents de maintenance.

1.3.2 Architecture télématique et instrumentation

Aujourd'hui, les véhicules de transport public sont généralement équipés en série d'un matériel télématique de première génération. Les systèmes d'information sont souvent hétérogènes. Ces équipements propriétaires sont difficiles à configurer et à maintenir. Un des objectifs du projet EBSF est de proposer une architecture standardisée sur IP intitulée *Smart Bus*⁸, pour les prochaines générations de bus. Ainsi, les opérateurs et autorités organisatrices pourront acheter des véhicules pourvus d'un équipement de base et installer eux-mêmes des services spécifiques « plug and play », à l'image des applications sur smartphone. Le télédiagnostic représentera l'une de ces fonctionnalités et facilitera la mise en place d'une politique de maintenance conditionnelle dont l'objectif est de réduire le nombre de pannes en exploitation.

8. EBSF IT Standard : www.ebsf.eu/index.php/ebsf-it-standard/.

Le bus de communication le plus répandu dans les véhicules automobiles est le bus CAN (Controller Area Network) (Vlasic et al., 2001, chapitre 2), standardisé par Bosch dans les années 80. Le CAN est conçu pour permettre aux différents boîtiers électroniques embarqués, appelés *Electronic Control Unit / Input Output Unit* (ECU/IOU), de communiquer les uns avec les autres par des messages de 8 octets de données.

La figure 1.3 décrit l'emplacement des principaux ECU des autobus disponibles pour l'étude. On cite notamment les deux calculateurs ATON et VIDAC (1^{er} voussoir gauche) qui collectent les données de l'étude. Tous les bus sont multiplexés par le bus CAN (appelé *CAN constructeur*, ou encore *multiplex*) sur lequel transite un grand nombre de données relatives au fonctionnement du véhicule.

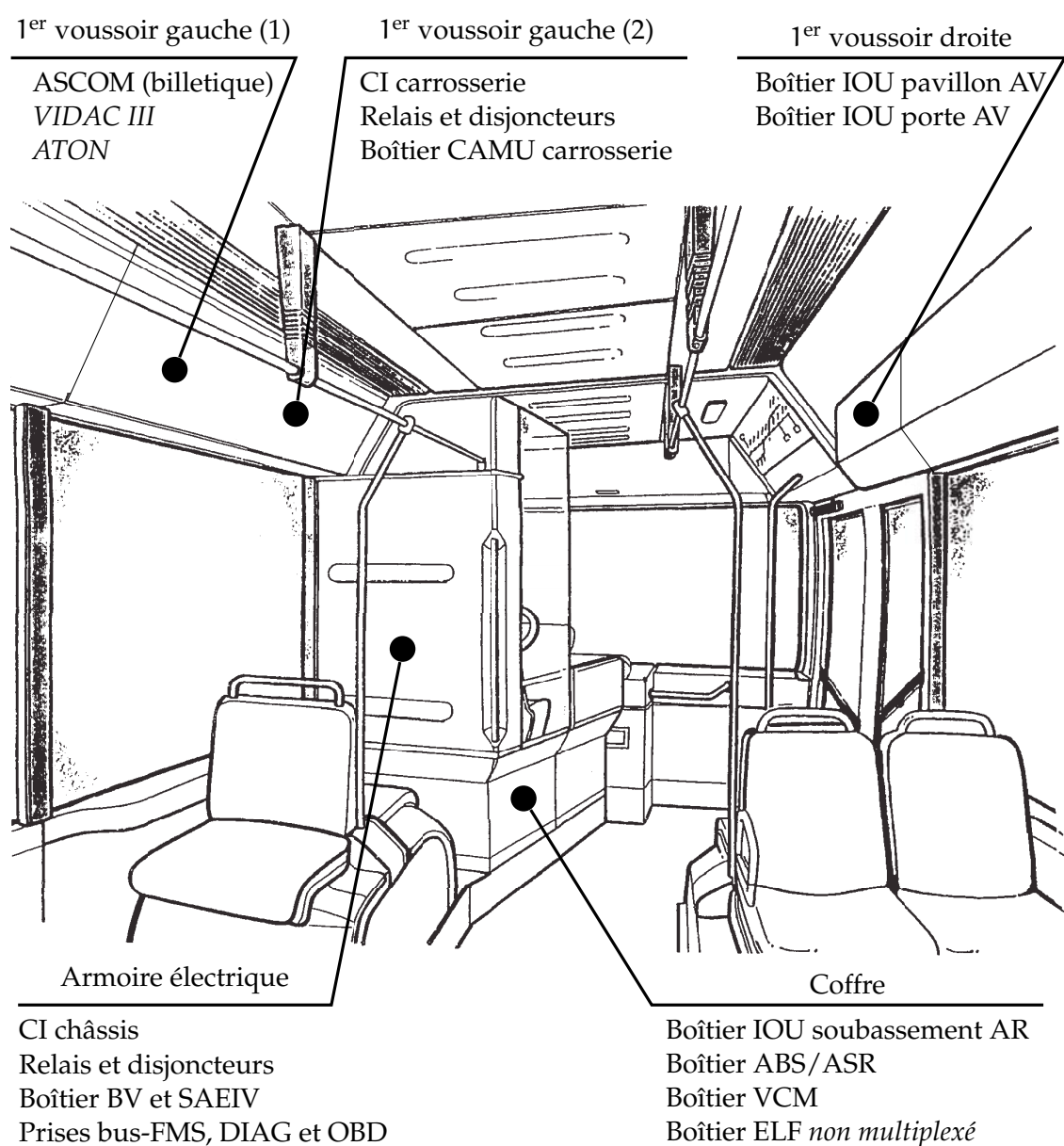


FIGURE 1.3 – Localisation des principaux boîtiers électroniques embarqués.

Dans le cadre du projet EBSF, une architecture *Smart Bus*⁹ a été préconisée. Les différents systèmes d'information (billettique, information trafic,...) dialoguent entre eux via un langage commun sur un unique réseau IP. Un accès au CAN constructeur est prévu au travers d'une prise, dite bus-FMS (voir section suivante). A terme, le système de diagnostic préconisé aura accès aux informations qui transitent sur le réseau IP et à celles accessibles par le bus-FMS.

Dans le cadre du UseCase de Brunoy, une partie de l'architecture IT préconisée par EBSF a pu être implémentée sur dix véhicules en partenariat avec le constructeur IRIS-BUS et l'équipementier DIGIGROUP. L'équipement est composé d'une partie embarquée (calculateurs, connectique et capteurs existants/additionnels) et une partie débarquée (serveur au dépôt de bus) dite *Back office*. Les données de fonctionnement des véhicules ont été collectées à partir de capteurs additionnels clés et du CAN constructeur. La quantité importante de données nous a contraints à les stocker temporairement dans des calculateurs embarqués lorsque les autobus sont en exploitation. La figure 1.4 décrit l'architecture mise en place pour le UseCase de Brunoy. Pour les besoins de nos travaux, deux calculateurs interconnectés (en jaune sur la figure) ont été installés : ATON (VERI) et VIDAC (DIGIGROUP), pour le suivi respectivement des sous-systèmes de freinage et de porte.

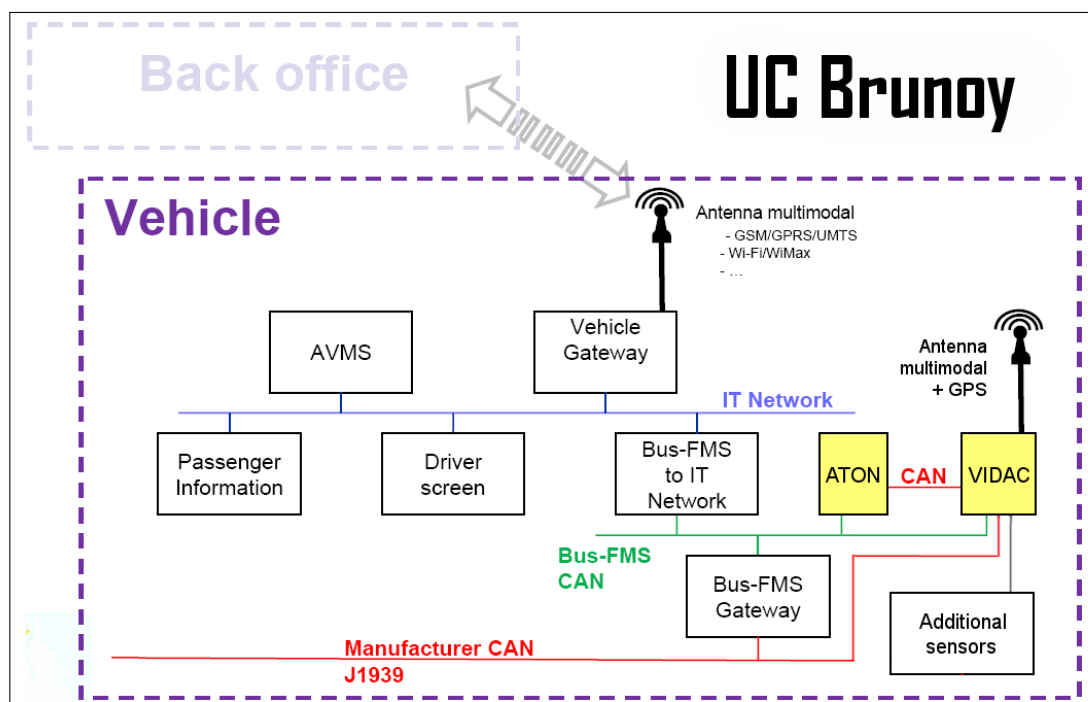


FIGURE 1.4 – Architecture télématique implémentée pour le UseCase de Brunoy.

9. *Smartbus*, the EBSF IT Architecture : <http://ebsf.eu/index.php/ebsf-it-standard>.

1.3.3 Spécification des données

Le **J1939**, introduit par la Society of Automotive Engineers (SAE) ou [SAE Technical Standards \(2006\)](#), est un protocole de haut niveau utilisé sur le bus de communication CAN. Les trames de données qui circulent sur le CAN sont notamment caractérisées par deux numéros. Le premier numéro est appelé **PGN** (Parameter Group Number) et permet d'identifier un groupe de paramètres réunissant des informations de la même « catégorie » (par exemple : les fluides du moteur, les températures moteur...). A l'intérieur de ces groupes, chaque paramètre peut être identifié par son numéro **SPN** (Suspect Parameter Number) et caractérisé par une description, une taille (octets), une résolution (unité/bit) et une gamme de valeur. Il est à noter que toutes les données spécifiées dans le J1939 ne sont pas présentes sur les trames CAN d'un véhicule.

Un exploitant d'autobus peut accéder aux données de fonctionnement du véhicule via des installations supplémentaires : capteurs additionnels et/ou interface bus-FMS. L'interface **FMS**¹⁰ (Fleet Management System) a pour objet de donner à un tiers, un accès à certaines données du véhicule. Les principaux constructeurs européens Daimler Buses - EvoBus GmbH, NEOMAN Bus GmbH, Scania CV, Volvo Bus Corporation (y compris Renault), IRIBUS IVECO et VDL Bus International B.V. ont mis au point l'interface bus FMS-Standard v2 en (2009) afin d'évaluer des données de fonctionnement et tester des applications indépendantes de celles fournies par les constructeurs. Une fois implémenté, ce protocole rend accessible plus facilement une centaine de variables. Il est à noter que le groupe de données disponibles via le bus-FMS correspond à un sous-ensemble de données spécifiées dans le J1939. Cette interface, illustrée par la figure 1.3, a bien été implémentée par IRISBUS sur l'ensemble de la flotte comprenant les dix autobus.

Avant de commencer la collecte et le traitement des données, un long travail d'identification des variables pertinentes a été nécessaire afin de caractériser le fonctionnement des systèmes. Un large nombre de données aurait pu être proposé avec pour chacune d'elles, plus ou moins d'information sur le fonctionnement et le contexte des systèmes à surveiller. Cette étude préalable a été contrainte par l'absence de certaines données de fonctionnement et par la difficulté d'accès à certaines autres en lien avec des problèmes de propriété industrielle. Dans le cadre du projet EBSF, il a été possible de répondre partiellement à ces limitations par la pose de capteurs additionnels avec l'accord de l'exploitant du réseau et l'acquisition de données CAN supplémentaires via le bus FMS, avec l'autorisation du constructeur. L'acquisition de ces variables de fonctionnement a fortement conditionné nos choix de modélisation.

10. Bus-FMS-Standard : <http://www.bus-fms-standard.com/>.

Comme le précise la figure 1.4, un calculateur spécifique est utilisé pour l'enregistrement des mesures de fonctionnement de chaque système. En outre, les modèles de surveillance sont alimentés par des flux de données différents :

Système de freinage : le modèle de détection s'appuiera sur un modèle physique intégrant la dynamique longitudinale du véhicule. Un nombre relativement important de données de fonctionnement est nécessaire pour alimenter le modèle. Le détail des variables retenues pour ce système est présenté dans la figure 4.10.

Système des portes : le détecteur se base uniquement sur le suivi de deux variables de fonctionnement. L'approche *à base de données* a été motivée par l'absence de modèle physique disponible représentant le système. Le détail des variables retenues pour ce système est présenté dans la partie 5.5.2.

1.3.4 Agrégation des connaissances réelles d'exploitation

L'objectif de ces travaux est de proposer de nouveaux outils aux exploitants de lignes de transport afin d'améliorer leurs actions de maintenance. La finalité sous-jacente est de concevoir des outils adaptées aux conditions réelles d'exploitation et les outils proposés devront être aisément interprétables par les agents de maintenance. Aujourd'hui, les exploitants de bus accumulent souvent une quantité importante de données sur la gestion de leur maintenance (imputations, réparations, références, noms des agents). Cela permet de suivre les actions réalisées et de faciliter la gestion des pièces. Dans le meilleur des cas, des données de type log-alarmes sont stockées (alertes et alarmes remontées à partir d'outils des constructeurs/équipementiers) mais on trouve assez rarement des informations sur le fonctionnement interne des systèmes.

Sur le réseau de la STRAV dont le dépôt de Limeil est illustré par figure 1.5, un outil de GMAO (Gestion de la Maintenance Assistée par Ordinateur), dénommé **WINATEL**, nous a permis d'observer la répartition des différentes opérations sur l'ensemble des véhicules du réseau. Ce logiciel permet de gérer la configuration des équipements et assure les tâches de logistique opérationnelle (approvisionnement des pièces et magasinage). Transdev préconise l'utilisation de cet outil sur chacun de ses réseaux ; toute intervention est ainsi caractérisée par plusieurs informations (date, pièces utilisées, coût des pièces et de la main d'œuvre,...) sous la forme d'un rapport historisé sur WINATEL. Chaque intervention doit également être associée à l'un des 12 groupes d'imputation (moteur, carrosserie, servitudes,...) et un des 10 groupes de réparation (visites, accidents, vandalisme,...). Néanmoins, chaque réseau est libre de choisir la liste des imputations et réparations ainsi que leurs affectations dans les différents groupes. Selon l'agent, les actions de maintenance qu'il réalise sont décrites plus ou moins brièvement et remontent rarement jusqu'à la cause de l'anomalie.



FIGURE 1.5 – Autobus articulé Citelis Euro IV devant le dépôt de la STRAV à Limeil-Brévannes (Val de Marne, 94).

Suite à des extractions régulières de la base WINATEL, il a été possible de renseigner des événements concernant les freins et les portes par la création de synthèses d'intervention sur ces systèmes. Les listes détaillées des modes de dysfonctionnement n'étaient pas disponibles mais des discussions informelles avec les agents du dépôt ont permis d'éliciter quelques défauts récurrents sur la base de leur expérience. En nous appuyant sur le savoir-faire des équipes de maintenance, certaines de ces défaillances ont pu être provoquées volontairement sans toutefois dégrader les bus de manière irréversible afin de constater leur incidence sur les variables mesurées. Ce type d'information est essentiel pour compléter le diagnostic des systèmes après détection d'une anomalie. L'identification des défauts réels que nous avons pu recenser est détaillée dans la sous-section 5.5.4 pour les portes. En revanche, l'identification des défauts sur les freins n'est pas présente dans ce mémoire car il n'a pas été exploité pour le travail concernant le suivi du système de freinage.

Depuis leur première mise en circulation, environ 60% des interventions sur les bus du réseau sont des actions de *maintenance corrective* et près de 40% représentent des mesures de *maintenance systématique* (visites réglementaires). Nécessitant un investissement initial plus important, une approche conditionnelle permettrait toutefois de réduire les coûts d'exploitation et de maintenance à moyen terme. Il est à noter que cette stratégie requiert un grand nombre de données historisées et un suivi spécifique des systèmes à maintenir. Les différentes stratégies de maintenance sont présentées dans la partie 2.1.1. Au dépôt, notre démarche est actuellement sans précédent et constitue la seule méthode d'analyse des données de fonctionnement pour les systèmes à surveiller.

1.4 Étude bibliographique

La dernière section de ce chapitre présente une vue d'ensemble des méthodes actuelles pour la modélisation et la surveillance des systèmes de freinage et de porte.

1.4.1 Système de freinage

Le système de freinage constitue un organe de sécurité évident pour tout véhicule (conduite du véhicule, distance d'arrêt...). Dès 2006, le projet Large Truck Crash Causation Study (LTCCS, données collectées de 2001 à 2003 aux USA) avait mis en évidence l'impact du système de freinage sur les accidents : 30% des accidents mortels étaient liés à des défauts ou dérèglages des freins pour les véhicules lourds ([Federal Motor Carrier Safety Administration \(FMCSA\), 2006](#)). Le système de freinage a ainsi été classé comme premier facteur de cause parmi les accidents rapportés. Il y a quelques mois, le Commercial Vehicle Safety Alliance (CVSA) Roadcheck 2012¹¹ a mené une campagne d'inspections de 3 jours sur autoroutes en Amérique du Nord. Presque un véhicule commercial sur quatre fut déclaré « hors service » pour des raisons mécaniques ; au total, 74 072 camions et autobus furent inspectés. Comme aux Roadcheck précédents, les défauts liés au système de freinage représentaient près de la moitié des violations constatées. Suite à des résultats similaires au Roadcheck 2008, [Freund et Skorupski \(2009\)](#) ont préconisé l'utilisation de capteurs pour la collecte de mesures liées à la force de freinage à chaque roue. La plupart des travaux publiés sur le système de freinage pneumatique établissent une corrélation entre la force de freinage, la décélération du véhicule et le couple de freinage d'une part, et la pression dans les cylindres de frein d'autre part ([Johnson et al., 1978](#); [Shaffer et Alexander, 1995b](#); [Hedrick et Uchanski, 2001](#); [Yu et al., 2010](#)). La remontée dynamique de ces données et leur analyse en continu permettrait de mieux répartir les actions de freinage à chaque roue d'une part et ouvrirait la voie à une maintenance conditionnelle, voir prévisionnelle d'autre part.

[Freund et Skorupski \(2009\)](#) et [Van Order et al. \(2009\)](#) ont décrit un cas d'étude dont l'objectif est de contrôler le système de freinage à partir de nouveaux équipements (capteurs additionnels, collecte de 59 variables à 50Hz par un ordinateur embarqué...). L'étude a été menée sur un semi-remorque puis une flotte de 12 bus urbains (de novembre 2006 à novembre 2007) à Washington DC, USA. Il a été observé une forte corrélation entre la force de freinage mesurée et la décélération du véhicule (et donc sa vitesse). Le poids du véhicule a également impacté la performance du système de surveillance (10 scénarios de dégradation). Le nombre d'inspections au dépôt, la maintenance et la fiabilité des systèmes ont été évalués. Le nouvel équipement (capteurs et algorithmes) a réduit les temps d'inspection et un seul défaut de capteur a été observé sur les 12 mois d'étude ([Van Order et al., 2009](#), pages 87-88). Cette approche est semblable à celle que nous avons proposée dans le cadre de cette thèse et présentée dans le chapitre 4.

11. 25^e CVSA's Roadcheck (juin 2012) : http://www.cvsa.org/news/2012_press.php.

Certains constructeurs émettent directement des instructions d'entretien et de test (basées sur la conception) afin de faciliter la maintenance des systèmes. [WABCO \(2007\)](#) propose par exemple l'utilisation d'arbres de défaillance pour l'identification de défauts sur les freins dans des véhicules agricoles ou forestiers, tout en précisant que certaines instructions peuvent varier en fonction de la configuration du système. Dans notre étude, le robinet de frein principal est de marque WABCO mais les étriers, cylindres et autres valves ont été conçus par d'autres constructeurs et sont accompagnés de recommandations pour la plupart. En pratique, il est difficile d'agréger ces recommandations locales d'entretien après conception afin d'établir une stratégie de surveillance globale.

La stratégie à base de modèles internes est une approche populaire mise en œuvre depuis plusieurs décennies (voir la sous-section 2.1.2). Ces modèles présentent l'avantage d'offrir une interprétation physique des systèmes. Un tel outil permet souvent de simuler le comportement du système et ainsi de réduire les coûts par rapport à des essais en condition réelle. [Weeks et Moskwa \(1995\)](#) présente un modèle de contrôle dynamique d'un moteur basé sur l'estimation de son couple moyen et implémenté sous MATLAB/SIMULINK sous la forme de blocs inter-connectés. Ce type de modèle est également utilisé afin de dimensionner le niveau de complexité nécessaire du modèle pour obtenir un compromis entre simplicité et précision. En effet, un modèle interne efficace doit être assez précis pour capturer l'essentiel de la dynamique du système tout en étant suffisamment simple pour faciliter son implémentation en temps réel. Les travaux de [Fisher \(1970\)](#) introduisent une modélisation dynamique pour des freins à tambour et à disque. Le modèle mathématique sous-jacent est formulé mais souffre d'une forte complexité. [Khan et al. \(1994\)](#) ont réduit ce modèle à dix états sur la base d'une modélisation du *vacuum booster* (compresseur), du cylindre principal et des cylindres aux roues (voir figure 1.6). [Gerdes \(1996\)](#) propose un modèle du vacuum booster pour un système hydraulique. A chaque fois, les auteurs mesurent la pression dans les cylindres aux roues et cherchent à prédire les variations de cette pression avec leurs modèles. Les données d'entrée comprennent la force appliquée par le conducteur à la pédale de freinage F_{br} , la pression du compresseur P_c et des constantes physiques liées au système de freinage. Ces modèles peuvent s'avérer assez complexe et requièrent toujours une forte expertise pour l'estimation de certains paramètres physiques. Pour surveiller le système de freinage, une stratégie de contrôle devrait idéalement suivre son fonctionnement de la pédale de frein jusqu'aux roues du véhicule.

Développer un outil de diagnostic du système de freinage est difficile à cause du nombre important de ses composants. Afin d'obtenir une bonne répartition de la force de freinage, la plupart des constructeurs ajoutent même des « soupapes de freins » au circuit de freins ([Limpert, 2011](#)). La figure 1.6 représente un système de freinage hydraulique très simplifié. On remarque notamment que le circuit principal est séparé en deux sous-circuits indépendants (croisés) pour des raisons de sécurité. Dans notre application, le

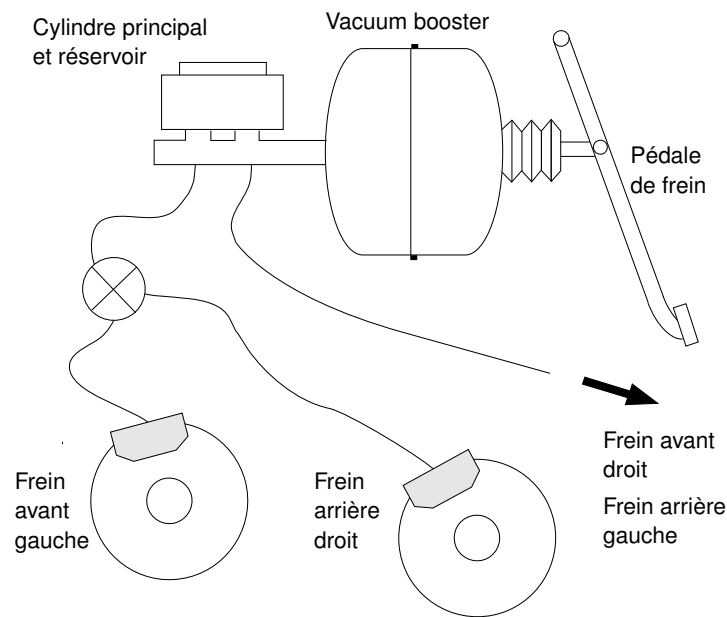


FIGURE 1.6 – Diagramme simplifié d'un système de freinage hydraulique.

Source : (Hedrick et al., 1997).

circuit est également séparé en deux mais celui-ci est partitionné différemment : sous-circuit avant et arrière.

Récemment, un nouveau modèle physique de freinage pneumatique a été développé à des fins de diagnostic (Subramanian et al., 2004; Natarajan et al., 2007; Karthikeyan et al., 2010; Sonawane et al., 2011) mais s'avère assez complexe. Ce modèle est utilisé afin de corréliser la pression des cylindres avec le déplacement de la tige de commande et la force appliquée au niveau de la pédale de frein; la modélisation s'appuie sur les travaux de thèse de Subramanian (2006). Le diagnostic porte sur un système de freinage à tambour avec des cames en S. La figure 1.7(a) illustre cette technologie de frein. Dans ce type de freins, la came en S agit sur les mâchoires et les garnitures de frein et les applique contre le tambour. Cette technologie est la plus commune dans les véhicules commerciaux aux USA alors que les freins à disque ont lentement remplacé les freins à tambour en Europe (Subramanian, 2006).

Dans notre étude, les autobus sont tous équipés d'un système de freinage à disque et au lieu du système à came, c'est une « vis de commande » qui est utilisée et qui agit comme une vis de serrage, de façon à ce que les garnitures répartissent également la force des deux côtés du disque. La figure 1.7(b) illustre cette technologie de frein. Enfin, on trouve d'autres modélisations physiques assez détaillées pour le contrôle de freins pneumatiques (Bu et Tan, 2007) ou électropneumatiques (Karthikeyan et al., 2011) de véhicules commerciaux.

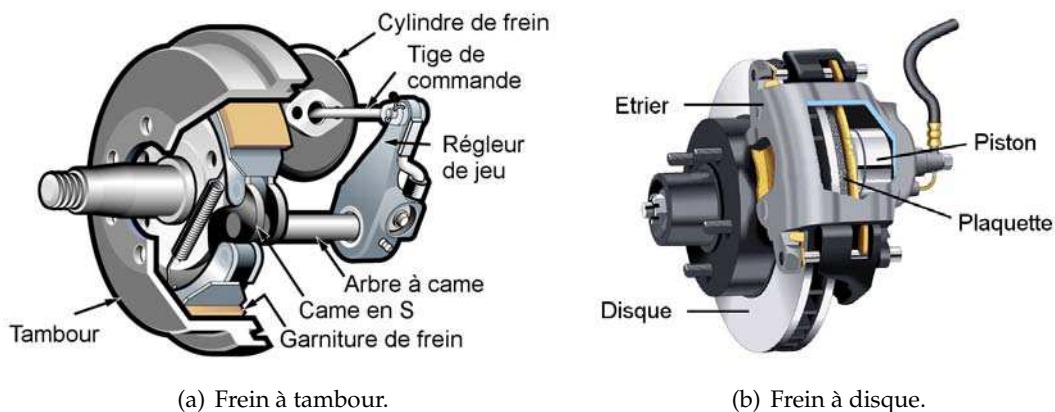


FIGURE 1.7 – Mécanismes de frein à tambour (a) et à disque (b).

Vlacic et al. (2001, chapitre 12) présente une modélisation globale d'un système de freinage hydraulique. Les auteurs suggèrent que cette modélisation est suffisamment simple pour l'implémentation d'un système de contrôle tout en parvenant à capturer l'essentiel de la dynamique du phénomène. Les formulations des modèles longitudinaux des véhicules et de contrôle des freins sont décrits avec plus de détail dans (Maciuca et al., 1994; Hedrick et al., 1997; Hedrick et Uchanski, 2001). Plusieurs modèles sont proposés avec des réductions justifiées par des principes physiques tout en conservant des relations électriques non linéaires. Sous certaines conditions et hypothèses, il y a une relation quasi-linéaire entre la pression P_b dans les cylindres des sous-circuits et le couple de freinage total K_b . Des simulations et résultats expérimentaux sont présentés et valident l'approche proposée. Ces travaux s'appuient sur la modélisation du vacuum booster et des freins hydrauliques décrits dans les thèses de Gerdes (1996) et Maciuca (1997) dans le cadre du PATH (Partners for Advanced Transit and Highways) au sein du laboratoire VDL¹². Huang et Ren (1999) introduisent un algorithme pour contrôler le délai entre accélération et freinage. Plusieurs méthodes de contrôle non linéaire sont présentées dans l'ouvrage de Song et Hedrick (2011).

Cette modélisation longitudinale, déclinée pour un système de freinage pneumatique, a récemment été utilisée pour le contrôle d'autobus standards (12m) ou articulés (18m). Song et Hedrick (2004) valident cette approche pour le contrôle de la vitesse et de la distance sur des bus à San Diego (Californie, USA). Dans le cadre du projet ANGO (PREDIT 4), Nouvelière et Braci (2007) ont développé un système d'assistance à la conduite longitudinale déployé sur une ligne du réseau TEOR à Rouen (France). Cette dernière étude constitue les prémisses des travaux menés pour le projet EBSF. La stratégie de détection de fautes, présentée dans le chapitre 4, repose sur une modélisation de la dynamique longitudinale d'un véhicule décrite dans ces travaux.

12. Thèses et publications du Vehicle Dynamic Lab (VDL) sur <http://vehicle.me.berkeley.edu/Publications/publications.htm>.

1.4.2 Système des portes

Les portes d'un véhicule de transport public sont particulièrement sollicitées (nombreux cycles) et exposées à des interactions avec les passagers (tenant ou s'appuyant sur la porte). Il est à noter que les défauts rencontrés sont souvent d'ordre mécanique dus à l'usure de ses composants. Un autobus est habituellement équipé de deux à quatre portes qui peuvent être pneumatiques ou électriques. Dans notre étude, les portes sont de type pneumatique et leur nombre dépend du type de véhicule :

- *Autobus standard (12m)* : deux portes ;
- *Autobus articulé (18m)* : trois portes.

Dans cette partie, on présente quelques travaux concernant la modélisation et la surveillance de systèmes de portes pneumatiques ou électriques, installés dans des autobus ou des trains.

[McCullers et Smith \(2001\)](#) détaillent un cas d'étude allant de la conception jusqu'au test d'un prototype pour la maintenance conditionnelle de portes électriques dans des autobus d'aéroport. Les auteurs proposent deux approches de surveillance en-ligne avant l'occurrence de défauts non réparables. Le principal signe d'une dégradation mis en avant est l'augmentation de la force de friction, qui pourrait mener à la casse du moteur. Il est supposé que la dégradation est mesurable et augmente en fonction du temps. En pratique, un type de dégradation est simulé par une barre de métal boulonnée à l'extérieur du véhicule et une pression est exercée pour créer des forces de frottement dans le sens contraire de celui de la porte. Les données d'entrée sont le courant du moteur, la tension aux bornes du moteur et les durées de chaque cycle de porte. Un premier modèle analytique est testé mais s'est rapidement révélé très compliqué à mettre en œuvre dans des calculateurs embarqués. Une seconde approche, dite empirique, s'avère plus pratique et plus efficace. Plusieurs réseaux de neurones sont mis en compétition (perceptron multi-couches, cascade corrélation, réseau RBF) à partir d'une base historisée comprenant quelques situations de dégradation. Le réseau cascade obtient la plus faible erreur de classification. Cette étude est également présentée dans [\(Smith et al., 2002\)](#) et il est bien précisé que ce type d'approche ne nécessite pas de lourds investissements financiers ni d'installation très contraignante dans les véhicules. Plus récemment, [Smith et al. \(2010\)](#) revisitent la même étude et précise les conditions de cinq situations de dégradation.

Par ailleurs, [Miguelanez et al. \(2008\)](#) proposent une approche de diagnostic à base d'ontologies appliquée à un système de porte pneumatique de train. Ce type de méthodes nécessite d'avoir à disposition une base de connaissance (formalisée en ontologies) la plus exhaustive possible. Ainsi, cette base est organisée sous la forme de concepts liés les uns aux autres, traduisant des dépendances entre composants. Un ensemble de règles basées sur ces relations permet de classer différents événements. Dans cette thèse, nous n'avons pas choisi cette approche dont la réussite dépend fortement de la connaissance experte contenue dans la base.

D'autres travaux ont été menés sur des portes de train équipées d'une technologie électrique ou pneumatique. Ces dernières années, l'université de Birmingham a proposé de nombreuses méthodes pour la surveillance des portes, qui sont présentées brièvement par la suite. On cite notamment une approche de diagnostic d'un équipement mécanique proposée pendant la thèse de [Lehrasab \(1999\)](#). Sur la base de cette étude, [Lehrasab et al. \(2002\)](#) présentent une méthode de diagnostic (détection et isolation) pour les portes pneumatiques d'un train. Des caractéristiques sont extraites de données collectées à partir d'une porte de train instrumentée pour l'étude :

- *Caractéristiques dynamiques* : pressions dans les cylindres, vitesses d'ouverture et de fermeture, délais entre l'activation des électrovannes et mouvement physique des portes. Les vitesses sont corrélées à la pression dans les cylindres ;
- *Caractéristiques spectrales* : coefficients d'une transformée de Fourier discrète sur le déplacement angulaire de la porte et sa densité spectrale de puissance.

Il est à noter que les caractéristiques spectrales sont utilisées pour identifier les défauts de capteur et les caractéristiques abstraites, pour détecter des fuites dans les cylindres. La classification des défauts détectés est effectuée par des réseaux de neurones (RBF avec une distance euclidienne et carte SOM). D'autre part, ([Roberts et al., 2002](#)) propose une démarche similaire pour le diagnostic d'une porte de train cette fois-ci électrique. Cette méthode est alimentée par le même type de caractéristiques (pressions, vitesses d'ouverture et de fermeture). L'identification des défauts est effectuée par un modèle de classification à base de logique floue, introduite dans les travaux de thèse de [Dassanayake \(2002\)](#). L'auteur suggère de surveiller un système de portes pneumatiques par des capteurs additionnels mesurant pressions et positions pendant le fonctionnement des portes. Cette démarche a été adoptée dans nos travaux et décrite dans la sous-section 5.1.2.

Plus récemment, [Dassanayake et al. \(2009\)](#) décrit une méthode d'extraction de paramètres physiques pour le diagnostic de portes électriques de train dans le cadre d'une maintenance conditionnelle. Le système de porte est modélisé par un modèle dynamique, démontré dans ([Dassanayake, 2002](#)), et le moteur est décrit par une équation différentielle. On extrait de ces systèmes d'équation des paramètres mécaniques (position de la porte, moments d'inertie et quelques coefficients liés à la force de frottement) d'une part, et d'autre part, des paramètres électriques du moteur (courant induit notamment). Cette approche suppose que les défauts du système de porte sont liés à la force de frottement du moteur (force opposée au mouvement des battants).

Enfin, on peut également citer les travaux de thèse de [Silmon \(2009, chapitre 5\)](#) qui détaille un cas d'étude sur des portes pneumatiques de train. Les données (pressions, positions) ont été collectées en laboratoire ainsi que la réalisation de situations de dégradation du système. Il est à noter que les données collectées ([Silmon, 2009, figures 5.3 et 5.4](#)) sont très similaires à celles de notre étude. En revanche, l'approche choisie par l'auteur ne considère pas la nature fonctionnelle des données.

1.5 Conclusion

Ce chapitre introduit le contexte industriel dans lequel s'inscrit la thèse et plus particulièrement un cas d'étude du projet [EBSF](#) concernant la surveillance de sous-systèmes critiques pendant leurs fonctionnements. Pour notre étude, les systèmes de freinage et de porte ont été sélectionnés pour leurs impacts sur les coûts de maintenance et la disponibilité des véhicules. Nos travaux visent à créer des outils d'aide à la maintenance de ces systèmes afin d'améliorer les actions prises par l'exploitation. Dans ce contexte, une liste des données à surveiller a été préconisée, où chaque donnée est spécifiée dans le protocole J1939. Afin de collecter ces données, une architecture télématique a été déployée sur une flotte d'autobus d'un réseau urbain en Ile de France. Il est à noter que, lorsque la phase d'acquisition des données a débuté, une période non négligeable de validation a été nécessaire avant de pouvoir collecter des données exploitables sur les différents organes des véhicules à surveiller. Ce chapitre s'achève sur une étude bibliographique pour la modélisation et la surveillance des systèmes de freinage et de portes.

Ces travaux de thèse proposent de nouveaux outils pour l'aide à la maintenance d'un système de freinage et de portes pneumatiques. La suite du manuscrit présente les approches retenues et les modélisations élaborées pour le traitement des données puis le suivi de fonctionnement de ces deux organes d'autobus.

Chapitre 2

Méthodes de diagnostic par reconnaissance des formes intégrées aux processus de maintenance

Sommaire

2.1	Introduction	22
2.1.1	Stratégies de maintenance	22
2.1.2	Principales approches de diagnostic de systèmes industriels	24
2.2	Diagnostic à base de reconnaissance des formes	26
2.2.1	Démarche générique	26
2.2.2	Apprentissage statistique	28
2.2.3	Les modèles de mélange de lois et leur estimation	35
2.3	Diagnostic à partir de données fonctionnelles	40
2.3.1	Classification de données fonctionnelles	41
2.3.2	Modèle de mélange de régressions pour la modélisation de courbes à changement de régimes	42
2.3.3	Cas d'étude réel	49
2.4	Conclusion	52

2.1 Introduction

Les opérateurs de transport, particulièrement sensibles au contexte économique, cherchent à augmenter leur part de marché en proposant des prestations à haute qualité de service tout en réduisant leurs coûts d'exploitation par une amélioration des processus de maintenance. La mise en place d'une politique de maintenance efficace s'appuie habituellement sur des techniques de contrôle non destructif permettant de surveiller l'état des systèmes suivi d'un diagnostic précis des défauts.

Après une description succincte des différentes stratégies de maintenance, ce chapitre présente les principales approches de diagnostic de systèmes industriels. L'approche par reconnaissance des formes (RdF) et les paradigmes sous-jacents sont particulièrement détaillés. La situation classique du diagnostic par RdF concerne l'analyse d'observations multidimensionnelles. Dans ce chapitre, on s'intéresse également au diagnostic à partir de données fonctionnelles (réalisations discrétisées de fonctions inconnues). Cette approche est notamment utilisée pour les travaux menés sur les portes, et présentée dans le chapitre 5.

2.1.1 Stratégies de maintenance

La norme [AFNOR \(2010\)](#) a pour objet de définir un ensemble des termes génériques liés au concept de maintenance (hormis celle des logiciels). La maintenance regroupe l'ensemble de toutes les actions durant le cycle de vie d'un système (conception, exploitation et fin de vie) afin de le maintenir/rétablir dans un état de fonctionnement opérationnel. Une politique de maintenance efficace doit satisfaire les conditions suivantes :

- Assurer la disponibilité du bien pour la fonction requise, au coût optimal ;
- Tenir compte des exigences obligatoires (sécurité) relatives au bien ;
- Tenir compte des répercussions sur l'environnement ;
- Améliorer la durabilité du bien et/ou la qualité du produit/service.

En condition d'exploitation, le processus de maintenance consiste à contrôler régulièrement des indicateurs liés à la sûreté de fonctionnement du système ([AFNOR, 2010](#)). Lorsqu'un fonctionnement anormal est détecté, une liste d'actions est préconisée.

Les politiques de maintenance, illustrées par la figure 2.1 sont généralement regroupées en deux grandes familles ([AFNOR, 2010](#), annexe A) :

Maintenance Corrective : Cette politique de maintenance est la plus simple dans la mesure où elle est exécutée après une défaillance. Aucune étude analytique préalable ou stratégie de surveillance n'est nécessaire. On distingue la **maintenance palliative** qui vise à réparer provisoirement le système afin d'éviter des conséquences inacceptables, et la **maintenance curative** destinée à remettre un système dans un état fonctionnel permanent.

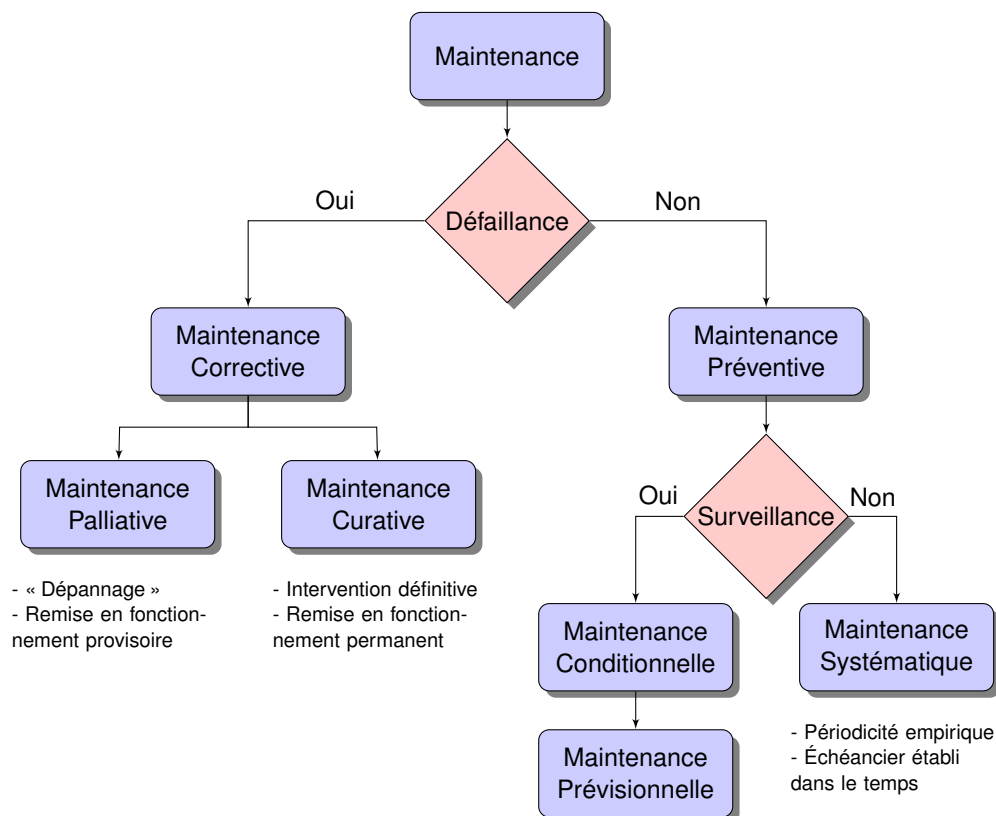


FIGURE 2.1 – Aperçu général des différentes politiques de maintenance.

Maintenance Préventive : Ce type de maintenance est exécuté en permanence et des actions peuvent être déclenchées à intervalles prédéterminés (**maintenance systématique**) ou en fonction d'indicateurs destinés à réduire la probabilité de défaillance/dégradation (**maintenance conditionnelle**). La maintenance systématique appliquée à une flotte d'autobus correspond à la planification de visites d'entretien après une durée ou un kilométrage prédéfini. [Haghani et Shafahi \(2002\)](#) proposent de modéliser ce problème sous la forme d'un programme linéaire en nombres entiers (PLNE) testé sur une flotte de 181 bus. D'autre part, la stratégie conditionnelle requiert une surveillance du système généralement basée sur une tâche de diagnostic. Récemment, [Ben Salem \(2008\)](#) et [Donat \(2009\)](#) ont proposé un formalisme basé sur des modèles graphiques probabilistes pour la maintenance conditionnelle de composants de l'infrastructure ferroviaire.

Enfin, on parle de **maintenance prévisionnelle** ([Cocheteux, 2010](#)) lorsque l'on intègre un modèle de prévision à la tâche de diagnostic. Ce type de processus est alors appelé *pronostic* ([Ribot, 2009](#)). Une telle politique nécessite au préalable d'avoir correctement formalisé l'évolution de l'état du système (modèle de dégradation).

La maintenance de la flotte de bus étudiée dans le cadre de cette thèse se partage notamment en 60% d'interventions correctives et 40% d'interventions systématiques.

2.1.2 Principales approches de diagnostic de systèmes industriels

Dans le cadre de la maintenance conditionnelle, on cherche à mettre en œuvre une stratégie de surveillance de certains systèmes basée sur un processus de diagnostic automatique. Le diagnostic est défini comme l'ensemble des actions menées pour la détection de la panne, sa localisation et d'identification de ses causes (AFNOR, 2010). La fonction de détection (voir chapitre 3) a pour objectif de signaler la présence d'une anomalie du système le plus tôt possible. Elle s'appuie sur la connaissance du fonctionnement nominal du système et tout changement de son état sous-jacent sera considéré comme une défaillance. La localisation de panne regroupe l'ensemble des actions menées en vue d'identifier à quel niveau d'arborescence du système en panne se situe le fait générateur de la panne (AFNOR, 2010). Après détection, les fonctions de localisation et d'identification nécessitent la connaissance *a priori* des états de panne du système et de pouvoir les caractériser (Charbonnier, 2006). Le tableau 2.1 résume les connaissances nécessaires afin d'établir un diagnostic.

	Fonction	Connaissance
1.	Détection	Données avant défaut
2.	Localisation / identification	Données après défaut Signatures des anomalies

TABLE 2.1 – Surveillance d'un système : les connaissances nécessaires.

On distingue trois approches pour établir le diagnostic d'un système industriel : systèmes experts, modèles analytiques, et reconnaissance des formes. Ces approches de diagnostic sont décrites dans plusieurs documents (Basseville et Cordier, 1996; Dubuisson, 2001a,b; Kempowsky, 2004; Aknin, 2008), et celles-ci sont résumées par la suite :

- La première famille de méthodes, appelées **systèmes experts**, datent des années 80. Cette approche cherche à reproduire le raisonnement exercé par un expert humain (Zwingelstein, 2002), souvent à partir de données symboliques. Dans cette famille, on retrouve notamment les méthodes d'arbre de défaillances et les méthodes de type AMDEC (Analyses des Modes de Défaillances, de leur Effets et de leurs Criticités). En particulier pour le domaine de l'automobile, on peut citer les travaux de thèse de Faure (2001) sur des arbres de diagnostic construits à partir de données de conception fournies par le constructeur. Dans cette approche, il est nécessaire d'avoir identifié au préalable tous les défauts/dysfonctionnements que peut subir le système sous la forme d'une base de connaissance. Ces méthodes sont dépendantes du système étudié, limitées à la détection de quelques défauts et désarmées face à de nouveaux modes de fonctionnement car le recensement des défauts doit être le plus exhaustif possible (Dubuisson, 2001b). Cette approche n'a donc pas été retenue pour nos travaux.

- Les **modèles analytiques**, également appelés méthodes internes, appartiennent usuellement au domaine de l'automatique (Isermann, 1984, 2011). Les données sont plutôt quantitatives et le fonctionnement interne (physique) du système est modélisé mathématiquement. La connaissance du modèle mathématique du système permet de générer des quantités (écarts entre sorties mesurées et celles estimées par le modèle), appelées *résidus*, qui sont quasi nulles en fonctionnement normal et non nulles en cas de défaillance. Cette approche est largement utilisée dans l'industrie. Elle offre l'avantage de donner un sens physique aux paramètres du modèle, ce qui facilite l'interprétation du diagnostic. L'analyse des résidus conduit à générer une alarme si un écart significatif est détecté (Basseville et Cordier, 1996) et le cas échéant, à localiser puis identifier le défaut. La figure 2.2 décrit toutes ces étapes :

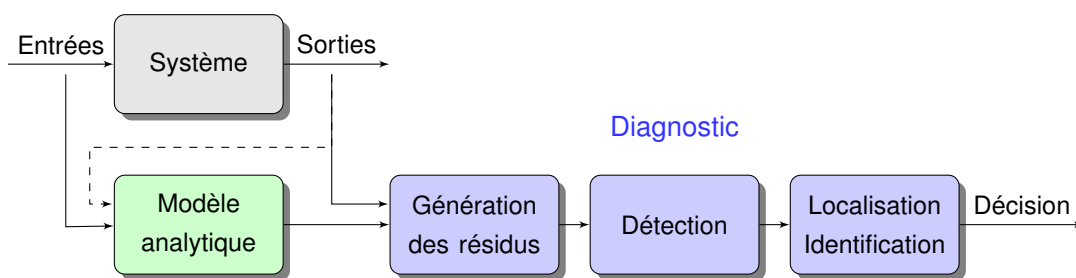


FIGURE 2.2 – Principales étapes pour établir un diagnostic automatique à base de modèle analytique.

L'étude sur le système de freinage, présentée dans le chapitre 4, propose une stratégie intégrant la modélisation longitudinale d'un bus. La mise en équation de ce modèle dynamique nous a permis d'estimer des indicateurs physiques non mesurés par notre étude. Cette formulation s'est fortement inspirée des travaux de Nouvelière et Braci (2007) et de Hedrick et Uchanski (2001, chapitre 2). La section 4.2 présente cette modélisation dédiée à la dynamique longitudinale d'un autobus.

- L'approche par méthodes externes à base de **reconnaissance des formes (RdF)** établit un diagnostic presque exclusivement à partir de données mesurées sur le système industriel. L'analyse de ces observations consiste à retrouver d'anciens modes de fonctionnement du système ou en découvrir de nouveaux. Ces observations sont parfois appelées données historisées (Kempowsky, 2004) car la collecte d'une quantité importante de données, dans une base d'apprentissage, nous aide à décrire le comportement du système. Plus cette base est riche d'informations, plus il est possible d'en extraire de la connaissance. Ces méthodes présentent donc l'avantage de ne pas avoir à construire un modèle physique complexe a priori. Si un tel modèle est disponible, il est toutefois possible d'utiliser les sorties du modèle physique comme des variables d'entrée pour une nouvelle méthode externe, comme dans notre étude sur le système de freinage (chapitre 4).

La prochaine section détaille l'approche de diagnostic par reconnaissance des formes.

2.2 Diagnostic à base de reconnaissance des formes

L'approche à base de Reconnaissance des Formes (RdF) (Duda et al., 2001; Theodoridis et Koutroumbas, 2008) a été privilégiée pour nos travaux. Cette démarche est particulièrement adaptée lorsque le comportement des systèmes à diagnostiquer est complexe et difficile à modéliser entièrement et précisément. Ce type de diagnostic s'appuie sur des techniques issues essentiellement de domaines tels que l'apprentissage statistique (Vapnik, 1999; Bishop, 2006; Hastie et al., 2009), l'analyse de données (Govaert, 2003; Saporita, 2006), la théorie de l'information (Cover et Thomas, 2006) et l'optimisation (Boyd et Vandenberghe, 2004; Minoux, 2007). Cette section a pour but de présenter les principaux paradigmes liés à la reconnaissance des formes. Il ne s'agit pas d'une description exhaustive de méthodes utilisées. On mettra notamment l'accent sur les approches paramétriques pour des problèmes statistiques d'estimation de densité, de régression et de classification.

2.2.1 Démarche générique

La construction d'un diagnostic à base de reconnaissance des formes est guidé par la validation de plusieurs étapes successives. Les choix des modélisations utilisées à chaque étape impactent la qualité des étapes suivantes et ces choix sont motivés par la nature des données mises à disposition. La démarche suit un raisonnement inductif : l'objectif est d'analyser un échantillon de données afin de caractériser l'état du système et de statuer sur de nouveaux exemples. La figure 2.3 présente les principales étapes qui composent un diagnostic à base de reconnaissance des formes.

Les premières étapes visent à caractériser au mieux l'état du système en fonctionnement. La spécification et l'**acquisition des données** sont généralement menées par des experts du système à surveiller. Dans notre cas d'étude, nous avons participé à l'élaboration des spécifications et à l'installation des instruments de collecte. Les données mesurées ont le plus souvent été soumises à diverses transformations (discrétisation, compression, conversion,...) auxquelles il convient de prêter attention car elles peuvent entraîner des pertes d'information. On applique ensuite différents **prétraitements** basés sur la connaissance du système et l'expérience du modélisateur (débruitage, normalisation...) afin d'extraire les caractéristiques les plus représentatives du système (Dubuisson, 2001b). Sur la base de ces indicateurs, on se place dans un **espace de représentation** contenant un maximum d'informations, ce qui nous permet de distinguer plus facilement les modes de fonctionnement du système. Toutefois, un espace trop vaste (de dimension trop élevée) implique une exploration difficile voir impossible ; on parle de fléau de la dimension (Bellman, 1961). L'espace final de description (*feature space*) est donc le plus souvent un sous-espace de l'espace de départ (input space). Dans ce but, on distingue deux approches concurrentes (Theodoridis et Koutroumbas, 2008, chapitres 5 à 7) de réduction de dimension :

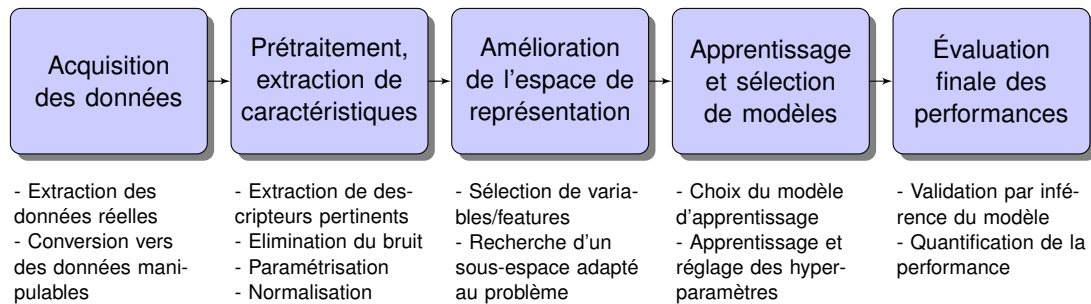


FIGURE 2.3 – Principales étapes pour établir un diagnostic automatique à base de reconnaissance des formes (Theodoridis et Koutroumbas, 2008, chapitre 1).

- *Sélection de variables (feature selection)* : on sélectionne un sous-ensemble des variables de départ. En général, il s'agit d'évaluer l'impact de l'ajout/la suppression d'une variable et de supprimer les variables les moins pertinentes (Guyon et Elisseeff, 2003). Les techniques exactes avec garantie d'optimalité (branch and bound, programmation dynamique,...) ne peuvent être envisagées que lorsque le nombre de variables n'est pas trop élevé. Les techniques approchées (recuit simulé, algorithmes génétiques,...) sont à base d'heuristiques et il n'est pas garanti que ces méthodes convergent vers un optimum global mais peuvent être envisagées pour la sélection d'un grand nombre de variables.
- *Projection (feature extraction)* : on projette les variables dans un sous-espace qui conserve un maximum d'information ; ainsi, les variables projetées perdent un caractère interprétable dans le nouvel espace. La technique la plus populaire est l'Analyse en Composantes Principales (ACP) introduite par Pearson (1901). La projection sur des axes orthogonaux de cette méthode maximise la variance expliquée (empirique) des données (Jackson, 2004) : la direction d'un axe est fournie par un vecteur propre de la matrice de covariance et la variance expliquée par cet axe est donnée par la valeur propre associée à ce même vecteur propre. Les travaux de Borga et al. (1992) fournissent une description concise et compacte de plusieurs projections linéaires dont l'ACP, la PLS (Partial Least Squares) et la CCA (Canonical Correlation Analysis).

Ces étapes préliminaires sont fondamentales pour construire un modèle capable d'apprendre une représentation précise et concise du système.

Sur la base d'un espace de représentation, l'objectif est de statuer sur le mode de fonctionnement/comportement du système. Le diagnostic s'appuie sur une machine dite d'apprentissage (*machine learning*, en anglais) qui génère des décisions sur la base des données observées dans l'espace ad-hoc. Dans le reste de la section, quelques paradigmes liés à ces outils sont présentés dans la sous-section 2.2.2 puis l'analyse de données fonctionnelles est décrite dans la section 2.3.

2.2.2 Apprentissage statistique

Une **méthode d'apprentissage** est un algorithme qui estime une fonction ϕ entre les entrées $\mathbf{x}$ et sorties $\mathbf{y}$ d'un système à partir de données disponibles tel que décrit par l'équation 2.1 :

$$\phi : \begin{matrix} \mathcal{X} & \rightarrow & \mathcal{Y} \\ \mathbf{x} & \mapsto & \mathbf{y} \end{matrix} . \quad (2.1)$$

Une fois la fonction estimée (étape d'apprentissage), la méthode peut être utilisée (étape d'inférence) pour *prédire* les sorties du système ou également *analyser* et *interpréter* des valeurs d'entrée mises à disposition. Ces méthodes s'appuient traditionnellement sur des outils statistiques et informatiques : les observations sont considérées comme des réalisations de variables aléatoires, et ces observations peuvent être représentées par différentes structures de données. Le domaine de l'apprentissage statistique est notamment décrit dans les ouvrages de Vapnik (1999), Bishop (2006), Cherkassky et Mulier (2007) et Hastie et al. (2009).

Principaux problèmes posés en apprentissage statistique

Dans cette partie, on présente succinctement trois catégories de problèmes statistiques parmi les plus répandues : la classification, la régression et l'estimation de densité.

Classification : le problème de classification consiste à apprendre une fonction $\hat{\phi}$ afin de prédire la classe y à valeur discrète d'une entrée observée x . Dans ce cas, le domaine $\mathcal{Y}$ est un ensemble fini de catégories et il est souvent choisi $\mathcal{Y} = \{1, \dots, K\}$, où K est le nombre de classes. Par exemple pour un client de messagerie, il s'agit de classer un courrier électronique en spam ou le déposer dans la boîte de réception. Cette fonction est souvent appelée un discriminateur (statistique) ou un classifieur (reconnaissance des formes). L'ouvrage de Raudys (2001) dénombre près de 500 classifieurs. Le choix d'un type de modèle de classification n'est pas une tâche facile et doit donc être motivé à partir de fondations statistiques solides et d'une description fondée du système.

Régression : la régression consiste à apprendre une fonction $\hat{\phi}$ afin de prédire une sortie à valeur continue correspondant à une entrée observée. Dans ce cas, l'ensemble $\mathcal{Y}$ est souvent choisi comme l'ensemble des réels $\mathcal{Y} = \mathbb{R}$. Par exemple, une agence immobilière veut pouvoir prédire le prix d'un appartement en fonction de sa localisation, sa superficie, etc. Les problèmes de régression et de classification sont liés dans le sens où ils nécessitent de considérer des paires d'entrée-sortie. Il est parfois possible de reformuler un modèle utilisé pour un problème de classification afin de résoudre un problème de régression. On cite notamment la *régression à vecteurs supports* (SVR), extension des *machines à vecteurs supports* (SVM), dont une large introduction a été proposée par Schölkopf et Smola (2002).

Estimation de densité : l'estimation ou l'approximation de densité pour l'analyse de données est un problème statistique étudié depuis plusieurs décennies (Silverman,

1986). Supposons que l'ensemble des observations soient les réalisations d'une variable aléatoire $\mathbf{X}$. Le problème est d'utiliser ces réalisations afin de construire la densité de probabilité de $\mathbf{X}$ et de repérer dans cette densité des modes pertinents (en terme de diagnostic). Ce problème assez général peut être difficile à résoudre à cause de la nature assez complexe des données ou de leur dimension. Ces spécificités conduisent à des vraisemblances complexes (non convexes) dont la maximisation mène à des solutions sous-optimales (minima locaux, voir section A.1). Il est à noter que l'estimation de densité peut être utilisée pour des problèmes de régression et classification, notamment dans l'approche générative décrite ci-après.

Selon le problème statistique étudié, le modèle est habituellement appris à partir d'un ensemble de données observées *indépendantes et identiquement distribuées (i.i.d.)*, appelé *base d'apprentissage*. Ces observations sont supposées avoir été générées par une loi de probabilité jointe $P(\mathbf{x}, \mathbf{y})$, où $\mathbf{x}$ et $\mathbf{y}$ représentent respectivement les entrées et les sorties du système observé. Dans l'approche paramétrique, cette probabilité est notée $P(\mathbf{x}, \mathbf{y}; \theta)$, où θ appartient à un ensemble de paramètres noté Θ . On assimile Θ à une famille paramétrique de modèles. Après avoir défini une famille de modèles, il convient de choisir un *critère d'ajustement* (par exemple, une fonction de perte) qui mesure l'adéquation du modèle à l'échantillon de données disponibles. Le problème de l'apprentissage se résume donc comme l'estimation du meilleur paramètre θ à partir d'une base d'apprentissage et en fonction d'un critère d'ajustement. Afin d'éviter le sur-apprentissage (sous-section 2.2.2), on cherche à sélectionner un modèle spécifique aux données tout en conservant une bonne capacité de généralisation (Hastie et al., 2009).

Dans ces travaux de thèse, on s'intéresse à des modèles paramétriques dont les paramètres sont estimés par la méthode du maximum de vraisemblance. En d'autres termes, l'apprentissage des modèles est effectué par ce *principe du maximum de vraisemblance (MV)*. Pour un échantillon de paires d'observations i.i.d. $(\mathbf{x}, \mathbf{y}) = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_n, \mathbf{y}_n)\}$, la densité de probabilité, appelée *vraisemblance des données* et notée $L(\theta; \mathbf{x}, \mathbf{y})$, désigne le critère d'ajustement pour de nombreuses machines d'apprentissage :

$$p(\mathbf{x}, \mathbf{y}; \theta) = \prod_{i=1}^n p(\mathbf{x}_i, \mathbf{y}_i; \theta). \quad (2.2)$$

La mise en œuvre du MV définit finalement un problème d'optimisation qui cherche à maximiser une fonction de vraisemblance (continue et différentiable, voir annexe A.1), ou fonction de densité jointe définie sur l'ensemble des données disponibles :

$$\hat{\theta}_{MV} = \arg \max_{\theta \in \Theta} p(\mathbf{x}, \mathbf{y}; \theta). \quad (2.3)$$

Ce type de problème peut être résolu par de nombreux algorithmes d'optimisation (Boyd et Vandenberghe, 2004). On peut citer par exemple, les méthodes du premier ordre (descentes de type gradient), ou d'ordre deux (méthodes de type Newton-Raphson).

D'autre part, on peut également s'appuyer sur le formalisme bayésien qui fait l'hypothèse que le paramètre θ est issu d'une variable aléatoire de loi a priori $\mathcal{P}_\theta$. Dans ce cas, on estime θ par l'estimateur du *maximum a posteriori* (MAP) :

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta \in \Theta} p(\theta | \mathbf{x}, \mathbf{y}) . \quad (2.4)$$

En vertu de la règle de Bayes et en passant au logarithme, on remarque que l'estimateur du MAP est obtenu après maximisation d'une log-vraisemblance pénalisée :

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta \in \Theta} \{ \log p(\mathbf{x}, \mathbf{y} | \theta) + \log p(\theta) \} . \quad (2.5)$$

On poursuit par une rapide introduction au paradigme de l'apprentissage statistique en présentant les principales approches pour l'apprentissage d'un modèle. Afin de faciliter la lecture de cette sous-section et sans perte de généralité, on se place dans le cas d'un problème de classification (voir sous-section précédente). Après l'apprentissage d'une fonction ϕ entre données d'entrée $\mathbf{x}$ et de sortie $\mathbf{y}$, l'étape d'inférence consiste à prédire la classe (ou catégorie) d'une nouvelle donnée d'entrée, parmi un ensemble $\mathcal{Y}$ fini. Il est à noter que la donnée de sortie $\mathbf{y}$ est ici appelée étiquette ou label.

Supervision de l'apprentissage

On distingue trois principaux modes de supervision pour l'apprentissage d'un modèle de classification à partir d'une base de taille n :

Supervisé : toutes les étiquettes de la base d'apprentissage sont renseignées et associées à une classe particulière : $\text{BA} = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_n, \mathbf{y}_n)\}$. Suivant l'application, l'étiquetage complet des observations peut être laborieux et coûteux.

Non supervisé : la base d'apprentissage ne contient pas de labels (aucune donnée d'entrée n'est étiquetée) : $\text{BA} = \{\mathbf{x}_1, \dots, \mathbf{x}_n\}$. Ce type d'apprentissage est également appelé regroupement automatique, auto-apprentissage ou clustering.

Semi-supervisé : chaque étiquette peut être soit vide, soit labellisée : $\text{BA} = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_l, \mathbf{y}_l), \mathbf{x}_{l+1}, \dots, \mathbf{x}_n\}$, où l est le nombre d'observations labellisées. Ce paradigme est à mi-chemin entre un mode supervisé et non supervisé. Un large panorama de méthodes adaptées à ce mode de supervision est dressé par [Chapelle et al. \(2010\)](#). Ce formalisme est adopté pour l'estimation des paramètres dans la partie 5.2.3.

Enfin, il existe bien d'autres modes de supervision plus généraux. On cite notamment l'*apprentissage partiellement supervisé* ([Ambroise et al., 2001](#)) où chaque étiquette peut être labellisée par un ensemble de classes. Récemment, une *supervision douce* a été proposée pour le diagnostic d'un composant de l'infrastructure ferroviaire ([Côme, 2009](#)). Dans ces travaux, l'incertitude de la classe d'une étiquette est représentée par un degré de plausibilité ou une fonction de croyance sur $\mathcal{Y}$.

Approches génératives et approches discriminatives

Cette partie décrit succinctement les deux approches classiques d'apprentissage dans le cadre d'un problème de classification. On pose $\mathcal{Y} = \{c_1, \dots, c_K\}$, l'ensemble des classes. D'autre part, les probabilités sur les variables sont assimilées à des densités évaluées ponctuellement ; on adopte désormais la notation $p(x)$ pour $P(x)$.

Dans l'approche générative (Bishop, 2006, section 4.2), l'objectif est d'estimer la loi jointe sur les observations, définie par la densité $p(x, y)$, avec $x \in \mathcal{X}$ et $y \in \mathcal{Y}$. Ce type de modélisation est fortement lié au problème statistique d'estimation de densité (voir les principaux problèmes décrits précédemment). A l'étape d'inférence, les classes sont prédites à partir de la densité (conditionnelle) a posteriori $p(y|x)$, obtenue par la *règle de Bayes* :

$$p(y|x) = \frac{p(x, y)}{p(x)} = \frac{p(x|y) \cdot p(y)}{\sum_{y^*} p(y^*)p(x|y^*)}, \quad (2.6)$$

où $p(x)$ est la densité marginale des observations et $p(y)$ est la densité a priori de la classe y . En pratique, il n'est pas nécessaire d'estimer la densité $p(x)$ mais uniquement $p(x|y)$. Cette approche permet de prendre en compte quelques a priori pendant l'apprentissage, d'estimer un modèle riche sur le processus de génération des données et si besoin, de générer de nouvelles observations. On peut citer par exemple, les analyses discriminantes linéaire ou quadratique (LDA ou QDA) qui sont décrites rapidement par la suite.

Supposons que $\mathcal{X} = \mathbb{R}^d$, avec un nombre entier $d \geq 1$. Les méthodes LDA et QDA considèrent que la densité d'appartenance à chaque classe $p(x|y = c_k)$ est modélisée par une densité gaussienne (normale) multivariée, définie par :

$$\mathcal{N}(x; \mu_k, \Sigma_k) = \frac{1}{(2\pi)^{d/2} \sqrt{\det \Sigma_k}} \exp \left(-\frac{1}{2} (x - \mu_k)^T \Sigma_k^{-1} (x - \mu_k) \right) \quad \forall x \in \mathbb{R}^d, \quad (2.7)$$

où les paramètres $\mu_k \in \mathbb{R}^d$ et $\Sigma_k \in \mathbb{R}^{d \times d}$ désignent les deux premiers moments de la densité de la classe k , respectivement la moyenne et la matrice de variance-covariance. La différence entre les deux analyses LDA et QDA repose essentiellement sur des hypothèses différentes concernant la matrice de covariance. En LDA, toutes les matrices de covariances sont les mêmes $\Sigma_1 = \dots = \Sigma_K$ ce qui induit une frontière de décision linéaire sur $\mathcal{X}$. D'autre part en QDA, les matrices de covariances sont estimées pour chaque classe ce qui induit une frontière de décision quadratique sur $\mathcal{X}$.

En LDA et QDA, la classe de x est estimée par la règle du maximum a posteriori (MAP) donnée dans l'équation 2.8, dont la densité conditionnelle est définie en 2.6.

$$\hat{y} = \arg \max_{c_k \in \mathcal{Y}} \frac{\frac{n_k}{n} \cdot \mathcal{N}(x; \mu_k, \Sigma_k)}{\sum_{h=1}^K \frac{n_h}{n} \cdot \mathcal{N}(x; \mu_h, \Sigma_h)}, \quad (2.8)$$

avec $n_k = \#\{i|y_i = c_k\}$, le nombre d'éléments de la classe k .

On peut également résoudre un problème de classification par une *approche discriminative* (Bishop, 2006, section 4.3) qui vise à modéliser et à estimer directement la loi conditionnelle, définie par la densité $p(y|x)$. Ainsi, l'étape d'inférence s'appuie directement sur cette loi conditionnelle sans faire appel à la règle de Bayes (équation 2.6) pour labelliser toute nouvelle observation. Ce type de modèles présente l'avantage d'admettre moins de degrés de libertés, en comparaison avec l'approche générative, et répond à une problématique plus simple que l'estimation d'une densité globale sur les données observées. En effet, les paramètres du modèle servent uniquement à définir une frontière de décision entre les classes. Nécessitant peu d'hypothèses sur la forme des données, l'apprentissage discriminatif a tendance à être plus robuste à l'imprécision du modèle estimé (Bengio et Frasconi, 1996).

On peut citer quelques méthodes discriminatives historiques présentées dans l'ouvrage de Bishop (2006) : le perceptron de F. Rosenblatt (section 4.1), l'analyse discriminante de Fisher (section 4.1), la régression logistique (section 4.4), les réseaux de neurones (chapitre 5) et les machines à vecteur support (section 7.1). On cite en particulier, les méthodes utilisant l'astuce du noyau (*kernel trick*). Ce type de machines d'apprentissage s'appuie sur le calcul de produits scalaires entre données d'entrée. Notamment présentées par Schölkopf et Smola (2002) ou Shawe-Taylor et Cristianini (2004), ces méthodes utilisent un noyau (défini sur un espace fonctionnel spécifique, appelé RKHS) afin de projeter les données d'entrée vers un espace de représentation (potentiellement de très grande dimension) et d'utiliser un produit scalaire dans ce nouvel espace. L'avantage de ces méthodes est donc de projeter (facilement) les données puis de travailler dans le nouvel espace muni d'un produit scalaire. De telles méthodes sont dites non paramétriques car les paramètres de la machine ne sont pas estimés explicitement. Cependant, ces méthodes nécessitent de choisir un noyau adapté aux données et au problème à résoudre. Les *machines à vecteur supports* (SVM) sont des méthodes populaires utilisant l'astuce des noyaux. Historiquement, Boser et al. (1992); Vapnik (1999) ont proposé ce type de méthodes pour un problème de classification binaire (deux classes) dans un cadre supervisé.

Récemment, Widodo et Yang (2007) dressent un état de l'art des machines à vecteur support pour la maintenance conditionnelle via un processus de diagnostic appliqué au suivi de fonctionnement de plusieurs objets (table 2). Les limitations de ce type de méthodes se résument par leur manque de flexibilité dans le sens où il est souvent difficile d'inclure des connaissances a priori sur les données (sans prise en compte de la vraisemblance des données). Enfin de par leur nature, les méthodes discriminatives ne peuvent pas être utilisées pour générer de nouvelles observations.

Le lecteur pourra se référer aux travaux de Jebara (2004) ou Bouchard (2005) pour plus de détails sur les paradigmes d'apprentissage génératif et discriminatif. La prochaine partie présente comment réaliser un compromis entre simplicité et précision pour dimensionner un modèle via des critères de sélection de modèles d'une part, ou un apprentissage pénalisé d'autre part.

Sélection de modèle

Cette sous-section est dédiée au problème de dimensionnement d'un modèle statistique pour un problème donné. En d'autres termes, il s'agit d'identifier l'instance d'une famille de modèles qui, après estimation à partir d'une base d'apprentissage, présenterait un minimum d'erreurs (risque de prédiction) sur de nouvelles observations. On parle d'*erreur de généralisation* (Hastie et al., 2009, section 2.9). Quelques critères de sélection de modèles après apprentissage sont présentés puis quelques méthodes pour estimer un modèle directement à partir d'un critère pénalisé. Chacune de ces approches nécessite de fixer des paramètres qui contrôlent la complexité du modèle. Ces paramètres qui dimensionnent chaque modèle sont appelés *hyperparamètres*.

Tout d'abord, l'idée du compromis entre simplicité et précision d'un modèle statistique est stipulée. Un modèle trop précis sur la base d'apprentissage généralise mal sur de nouvelles observations. Ce compromis est également appelé *compromis biais/variance* (Cherkassky et Mulier, 2007, section 3.4), (Hastie et al., 2009, section 2.9). Ainsi, un modèle trop précis a un biais plus faible mais une variance plus importante. Dans ce cas précis, on parle de *sur-apprentissage* d'un modèle. Afin de réaliser un tel compromis biais/variance, il faudrait idéalement connaître les erreurs en apprentissage et en généralisation d'une famille de modèles, en fonction de sa complexité. Il est possible d'approcher ces erreurs à partir de techniques dites de *ré-échantillonnage* (Efron et Tibshirani, 1994). On cite notamment la technique de *validation croisée* qui nécessite de fixer un entier $K > 1$:

1. Séparer (aléatoirement) les données disponibles en K parties de même taille.
2. Pour un k fixé, apprendre le modèle sur les $(K - 1)$ parties des données et calculer l'erreur d'apprentissage. Calculer l'erreur de généralisation sur la k^e partie des données.
3. Faire ce qui précède pour $k = 1, \dots, K$. Combiner les K estimations d'erreurs.

Ce type de méthodes permet d'approcher les erreurs d'apprentissage et de généralisation afin de sélectionner efficacement un modèle statistique. Toutefois, ces méthodes souffrent de deux principales limitations : fixer la valeur de K et réaliser un grand nombre de simulations pour un coût souvent prohibitif.

On peut également faire appel à des critères de *sélection de modèles* (Bishop, 2006, partie 1.3) issues de la théorie de l'information (Cover et Thomas, 2006), après l'étape d'apprentissage, afin de comparer plusieurs configurations de modèles d'une même famille (variation des hyperparamètres). A cause du surapprentissage, la performance sur la base d'apprentissage n'est pas un bon critère pour garantir une bonne prédiction sur de nouvelles observations. On choisit donc le modèle qui minimise une erreur d'apprentissage pénalisée. Historiquement, le premier critère de sélection de modèle est le critère *Akaike Information Criterion* (AIC) de Akaike (1974) et le critère AIC3 est proposé par Bozdogan (1987) quelques années après. Dans le cas d'un modèle m de paramètre θ_m à variables latentes appris sur une base d'apprentissage de taille n , les équations 2.9 et 2.10 présentent

des formulations des critères AIC et AIC3. La sélection de modèle consiste à minimiser ces critères.

$$\text{AIC}(m) = -2\mathcal{L}(\hat{\theta}_m) + 2v_m \quad (2.9)$$

$$\text{AIC}_3(m) = -2\mathcal{L}_c(\hat{\theta}_m) + 3v_m. \quad (2.10)$$

où v_m est la taille du paramètre $\hat{\theta}_m$ estimé par le principe du maximum de vraisemblance, et $\mathcal{L}(\hat{\theta}_m)$ représentent la log-vraisemblance après estimation du modèle m .

Il existe d'autres critères et les critères *Bayesian Information Criterion (BIC)* (Schwarz, 1978) et *Integrated Classification Likelihood (ICL)* (Biernacki et al., 2000) sont également présentés. Ces deux critères sont définis par :

$$\text{BIC}(m) = -2\mathcal{L}(\hat{\theta}_m) + v_m \log(n) \quad (2.11)$$

$$\text{ICL}(m) = -2\mathcal{L}_c(\hat{\theta}_m) + v_m \log(n), \quad (2.12)$$

où $\mathcal{L}_c(\hat{\theta}_m)$ représente la log-vraisemblance complétée (voir la sous-section 2.2.3) après estimation du modèle m . La maximisation de ces deux critères, multipliés par un facteur $-\frac{1}{2}$, mène à la sélection d'un modèle de mélange dans la section 2.3.3 et pour l'étude des portes dans la sous-section 5.4.2.

Par ailleurs dans un contexte supervisé, un critère discriminant comme le *Bayesian Entropy Criterion (BEC)* peut être envisagé (Bouchard, 2005). Dans un contexte de segmentation en régression, Kim et al. (2009) propose une étude comparative entre une procédure basée sur des rapports de vraisemblance, des critères informatifs (AIC, BIC) et une méthode de validation croisée.

D'autre part, on peut directement guider le problème d'optimisation pendant l'étape d'apprentissage avec une fonction objectif pénalisée. On parle alors de manière analogue de *pénalisation*, *régularisation*, ou *shrinkage*. L'idée est que l'algorithme d'apprentissage converge alors vers une solution plus simple en pénalisant les modèles trop complexes. L'estimateur du MAP (équation 2.5) par exemple, est obtenu après la maximisation d'une log-vraisemblance pénalisée par la loi a priori des paramètres. Pour un problème de régression, la technique de la *régression ridge* (Hoerl et Kennard, 1970) introduit un critère pénalisé par la norme ℓ_2 sur les paramètres. La méthode du LASSO (Least Absolute Shrinkage and Selection Operator), introduite par Tibshirani (1996) utilise le critère des moindres carrés pénalisé par la norme ℓ_1 sur les paramètres ce qui annule certains paramètres tout en préservant la convexité du problème (quadratique). L'algorithme LARS, proposé par Efron et al. (2004), est une des méthodes les plus populaires pour résoudre le problème du LASSO.

Notons que les méthodes avec pénalisation nécessitent de fixer un hyperparamètre qui contrôle la complexité du modèle estimé. Il est à noter que les méthodes à noyaux, comme les SVM, intègrent automatiquement un terme de régularisation.

2.2.3 Les modèles de mélange de lois et leur estimation

Définition

Dans un contexte non supervisé, les modèles à variables latentes constituent une modélisation de choix. En particulier, les *modèles de mélange*, dans le cadre des approches génératives, restent une référence en la matière (McLachlan et Peel, 2000; Frühwirth-Schnatter, 2006; Mengersen et al., 2011). Ce modèle suppose que les données sont constituées de groupes homogènes dont les classes d'appartenance sont modélisées par des variables cachées (ou variables latentes). L'estimation des paramètres de ce modèle s'effectue par maximisation de la vraisemblance sur les données d'entrée disponibles (équation 2.3) supposées indépendantes. En vertu de la loi des probabilités totales, la densité d'un modèle de mélange fini est la combinaison linéaire des densités marginales de ses K composantes :

$$p(x; \theta) = \sum_{k=1}^K \underbrace{p(Y = c_k; \theta)}_{\pi_k} \cdot \underbrace{p(x|Y = c_k; \theta)}_{\varphi(x; \theta_k)}, \quad (2.13)$$

où Y est la variable (latente) de sortie associée à l'entrée $x \in \mathcal{X}$, π_k est la densité a priori de la k^e composante (classe) et $\varphi(x; \theta_k)$ est la densité de probabilité marginale de la k^e composante. En vertu de la règle de Bayes (équation 2.6) et une fois que le modèle de mélange a été estimé, on peut prédire la classe $\hat{y}$ d'une observation x par la règle du maximum a posteriori (MAP) :

$$\hat{y} = \arg \max_{1 \leq k \leq K} P(Y = c_k | x; \theta) \quad (2.14)$$

$$= \arg \max_{1 \leq k \leq K} \frac{\pi_k \cdot \varphi(x; \theta_k)}{\sum_{l=1}^K \pi_l \cdot \varphi(x; \theta_l)}. \quad (2.15)$$

Soit $BA = \{x_1, \dots, x_n\}$, la base d'apprentissage composée de n observations non labellisées et supposées indépendantes. D'après les équations 2.2 et 2.13, la vraisemblance des données s'écrit :

$$L(\theta; \mathbf{x}) = \prod_{i=1}^n \sum_{k=1}^K \pi_k \cdot \varphi(x_i; \theta_k). \quad (2.16)$$

Le paramètre θ est estimé par maximisation de cette vraisemblance ou de la fonction de log-vraisemblance :

$$\mathcal{L}(\theta; \mathbf{x}) = \sum_{i=1}^n \log \sum_{k=1}^K \pi_k \cdot \varphi(x_i; \theta_k). \quad (2.17)$$

Les fonctions données par les équations 2.16 et 2.17 sont habituellement maximisées par un algorithme de type EM. Cette stratégie consiste à maximiser une fonction de vraisemblance, dite complétée, plus simple à calculer. Le principe de ce type de méthodes est décrit dans la partie suivante.

Estimation par l'algorithme EM

La maximisation de la fonction de vraisemblance peut se révéler complexe à cause de la forme non convexe (voir définition A.3) de cette fonction. Il est également impossible de calculer les racines de la différentielle de la vraisemblance. On utilise alors l'algorithme *Expectation Maximisation (EM)* (McLachlan et Krishnan, 2008) qui est une méthode itérative adaptée au modèle de mélange. Historiquement, cet algorithme a été introduit par Dempster et al. (1977) et représente une solution classique pour l'estimation des paramètres d'un modèle à variables latentes.

Soit $\mathbf{Y}$, l'ensemble des variables latentes associées à l'ensemble des données d'entrée $\mathbf{x}$. L'algorithme EM s'appuie sur la réécriture suivante de la log-vraisemblance en fonction des variables latentes (équation de Chapman-Kolmogorov) :

$$\mathcal{L}(\theta; \mathbf{x}) = \log p(\mathbf{x}, \mathbf{Y}; \theta) - \log p(\mathbf{Y}|\mathbf{x}; \theta). \quad (2.18)$$

L'algorithme EM estime le paramètre θ en maximisant itérativement l'espérance de la log-vraisemblance sur la variable latente $\mathbf{Y}$ de chaque observation. Ainsi, à chaque itération (q), on cherche à maximiser la quantité suivante :

$$\mathbb{E}_{\mathbf{Y}} \left[\mathcal{L}(\theta; \mathbf{X}) | \mathbf{x}; \theta^{(q)} \right] = \underbrace{\mathbb{E}_{\mathbf{Y}} \left[\log p(\mathbf{X}, \mathbf{Y}; \theta) | \mathbf{x}; \theta^{(q)} \right]}_{Q(\theta, \theta^{(q)})} - \underbrace{\mathbb{E}_{\mathbf{Y}} \left[\log p(\mathbf{Y}|\mathbf{X}; \theta) | \mathbf{x}; \theta^{(q)} \right]}_{H(\theta, \theta^{(q)})}. \quad (2.19)$$

Démontrons que la fonction H diminue au cours des itérations :

$$H(\theta^{(q+1)}, \theta^{(q)}) - H(\theta^{(q)}, \theta^{(q)}) = \sum_{\mathbf{y} \in \mathcal{Y}^n} p(\mathbf{y}|\mathbf{x}; \theta^{(q)}) \log \frac{p(\mathbf{y}|\mathbf{x}; \theta^{(q+1)})}{p(\mathbf{y}|\mathbf{x}; \theta^{(q)})} \quad (2.20)$$

$$\begin{aligned} &\leq \log \sum_{\mathbf{y} \in \mathcal{Y}^n} \frac{p(\mathbf{y}|\mathbf{x}; \theta^{(q+1)})}{p(\mathbf{y}|\mathbf{x}; \theta^{(q)})} \frac{p(\mathbf{y}|\mathbf{x}; \theta^{(q)})}{p(\mathbf{y}|\mathbf{x}; \theta^{(q)})} \quad (2.21) \\ &= \log \sum_{\mathbf{y} \in \mathcal{Y}^n} p(\mathbf{y}|\mathbf{x}; \theta^{(q+1)}) \\ &= 0. \end{aligned}$$

Le passage de l'équation 2.20 à 2.21 s'appuie sur l'inégalité de Jensen (voir définition A.6) et sur la concavité de la fonction logarithme (McLachlan et Krishnan, 2008). On en déduit bien que la fonction H diminue en fonction des itérations de l'algorithme : $H(\theta^{(q+1)}, \theta^{(q)}) \leq H(\theta^{(q)}, \theta^{(q)})$. $\square$

Il est donc suffisant de maximiser la fonction Q , définie dans l'équation 2.19 pour estimer le paramètre du modèle θ . La densité $\log p(\mathbf{x}, \mathbf{y}; \theta)$ est appelée log-vraisemblance complétée et notée $\mathcal{L}_c(\theta; \mathbf{x}, \mathbf{y})$.

Dans le cas d'un modèle de mélange, la log-vraisemblance complétée s'écrit :

$$\begin{aligned}\mathcal{L}_c(\boldsymbol{\theta}; \mathbf{x}, \mathbf{y}) &= \log \prod_{i=1}^n \prod_{k=1}^K [\pi_k \cdot \varphi(x_i; \boldsymbol{\theta}_k)]^{y_{ik}} \\ &= \sum_{i=1}^n \sum_{k=1}^K y_{ik} \cdot \log [\pi_k \cdot \varphi(x_i; \boldsymbol{\theta}_k)] \\ &\quad \text{avec } y_{ik} = \begin{cases} 1 & \text{si } y_i = k \\ 0 & \text{sinon} \end{cases}.\end{aligned}\quad (2.22)$$

On écrit alors la fonction Q comme l'espérance de la log-vraisemblance complétée conditionnellement aux données observées et au paramètre courant :

$$\begin{aligned}Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(q)}) &= \mathbb{E}_Y \left[\mathcal{L}_c(\boldsymbol{\theta}; \mathbf{X}, \mathbf{Y}) \mid \mathbf{x}; \boldsymbol{\theta}^{(q)} \right] \\ &= \sum_{i=1}^n \sum_{k=1}^K \underbrace{\mathbb{E}_Y \left(Y_{ik} \mid \mathbf{x}; \boldsymbol{\theta}^{(q)} \right)}_{\tau_{ik}^{(q)}} \cdot \left(\log \pi_k + \log \varphi(x_i; \boldsymbol{\theta}_k) \right),\end{aligned}\quad (2.23)$$

où $\tau_{ik}^{(q)}$ représente la loi *a posteriori* qui mesure le degré d'appartenance de l'observation x_i à la k^e composante :

$$\begin{aligned}\tau_{ik}^{(q)} &= p(Y_i = c_k \mid \mathbf{x}; \boldsymbol{\theta}^{(q)}) \\ &= p(x_i, Y_i = c_k; \boldsymbol{\theta}^{(q)}) / p(x_i; \boldsymbol{\theta}^{(q)}) \\ &= \frac{\pi_k \cdot \varphi(x_i; \boldsymbol{\theta}_k)}{\sum_{h=1}^K \pi_h \cdot \varphi(x_i; \boldsymbol{\theta}_h)}.\end{aligned}\quad (2.24)$$

L'algorithme EM adapté aux modèles de mélange alterne itérativement les deux étapes suivantes :

1. **Étape E (estimation)** : Calcul des probabilités *a posteriori* $\tau_{ik}^{(q)}$;
2. **Étape M (maximisation)** : Mise à jour du vecteur de paramètres $\boldsymbol{\theta}^{(q+1)}$ en maximisant la fonction Q définie par l'équation 2.23 ;

jusqu'à convergence du critère sur la vraisemblance défini par l'équation suivante :

$$\left| \frac{L(\boldsymbol{\theta}^{(q+1)}; \mathbf{x}) - L(\boldsymbol{\theta}^{(q)}; \mathbf{x})}{L(\boldsymbol{\theta}^{(q)}; \mathbf{x})} \right| < \epsilon, \quad \forall \epsilon > 0. \quad (2.25)$$

Dans cette formulation dédiée à un modèle de mélange, cette procédure respecte les conditions de l'algorithme EM (Dempster et al., 1977) et la convergence de la vraisemblance vers un maximum local est donc garantie. La qualité de l'estimation par cet algorithme dépend de son initialisation. Plusieurs stratégies d'initialisation ont été proposées afin d'éviter l'estimation d'un maximum local. Dans ces travaux de thèses, nous avons choisi d'exécuter plusieurs fois l'algorithme de type EM avec initialisation aléatoires et on garde le paramètre associé à la plus grande vraisemblance (Biernacki, 2004).

Exemple classique pour un mélange de lois gaussiennes

Dans cette partie, on s'intéresse à un modèle de mélange fini où les densités marginales sont modélisées par des densités gaussiennes multidimensionnelles, définies par l'équation 2.7. A partir d'un échantillon $\mathbf{x}$ de taille n , la densité du modèle de mélange (vraisemblance) est définie par l'équation 2.26.

$$p(\mathbf{x}; \boldsymbol{\theta}) = \prod_{i=1}^n \sum_{k=1}^K \pi_k \cdot \mathcal{N}(\mathbf{x}_i; \boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k) \quad \forall \mathbf{x}_i \in \mathbb{R}^d \text{ et } d \geq 1, \text{ en entier.} \quad (2.26)$$

L'apprentissage d'un tel mélange est décrit par l'algorithme 2.1 de type EM. A chaque itération (q), cet algorithme recalcule les probabilités a posteriori (ligne 3) et met à jour le vecteur paramètre $\boldsymbol{\theta}$ en maximisant la fonction Q définie par l'équation suivante :

$$Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(q)}) = \sum_{i=1}^n \sum_{k=1}^K \tau_{ik}^{(q)} \cdot \log [\pi_k \cdot \mathcal{N}(\mathbf{x}_i; \boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k)]. \quad (2.27)$$

Il est à noter que [Celeux et Govaert \(1995\)](#) proposent des modèles de mélanges gaussiens parcimonieux basés sur différentes décompositions de la matrice de covariance.

Algorithme 2.1 : Pseudo-code de l'algorithme EM appliqué à un mélange gaussien.

Entrées : Échantillon de données indépendantes $\mathbf{x} = \{\mathbf{x}_1, \dots, \mathbf{x}_n\}$, et nombre K de composantes du mélange.

// Initialisation aléatoire de $\boldsymbol{\theta}^{(0)}$ si non fourni

1 $q \leftarrow 0$

Répéter

// Étape E : calcul de la loi a posteriori (équation 2.24)

2 **Pour** $(i, k) \in \{1, \dots, n\} \times \{1, \dots, K\}$ **faire**

3

$$\tau_{ik}^{(q)} \leftarrow \frac{\pi_k^{(q)} \cdot \mathcal{N}(\mathbf{x}_i; \boldsymbol{\mu}_k^{(q)}, \boldsymbol{\Sigma}_k^{(q)})}{\sum_{h=1}^K \pi_h^{(q)} \cdot \mathcal{N}(\mathbf{x}_i; \boldsymbol{\mu}_h^{(q)}, \boldsymbol{\Sigma}_h^{(q)})} \quad (2.28)$$

// Étape M : maximisation de la fonction Q définie par l'équation 2.27

4 **Pour** $k \in \{1, \dots, K\}$ **faire**

5

$$\pi_k^{(q+1)} \leftarrow \sum_{i=1}^n \tau_{ik}^{(q)} / n \quad (2.29)$$

$$\boldsymbol{\mu}_k^{(q+1)} \leftarrow \sum_{i=1}^n \tau_{ik}^{(q)} \cdot \mathbf{x}_i / \sum_{j=1}^n \tau_{jk}^{(q)} \quad (2.30)$$

$$\boldsymbol{\Sigma}_k^{(q+1)} \leftarrow \sum_{i=1}^n \tau_{ik}^{(q)} (\mathbf{x}_i - \boldsymbol{\mu}_k^{(q+1)}) (\mathbf{x}_i - \boldsymbol{\mu}_k^{(q+1)})^T / \sum_{j=1}^n \tau_{jk}^{(q)} \quad (2.31)$$

6 $q \leftarrow q + 1$

7 **Jusqu'à convergence** (équation 2.25);

Sortie : Paramètre estimé $\hat{\boldsymbol{\theta}} = (\hat{\pi}_1, \dots, \hat{\pi}_K, \hat{\boldsymbol{\mu}}_1, \dots, \hat{\boldsymbol{\mu}}_K, \hat{\boldsymbol{\Sigma}}_1, \dots, \hat{\boldsymbol{\Sigma}}_K)$.

A titre d'illustration, on étudie un cas bidimensionnel où les données sont générées par quatre lois gaussiennes. Les vecteurs de moyennes sont donnés par l'équation 2.32 et les matrices (pleines) de covariances sont données par l'équation 2.33.

$$\mu = \left\{ \begin{pmatrix} 30 \\ 30 \end{pmatrix}, \begin{pmatrix} 70 \\ 30 \end{pmatrix}, \begin{pmatrix} 70 \\ 70 \end{pmatrix}, \begin{pmatrix} 30 \\ 70 \end{pmatrix} \right\}, \quad (2.32)$$

$$\Sigma = \left\{ \begin{pmatrix} 55 & -40 \\ -40 & 55 \end{pmatrix}, \begin{pmatrix} 55 & 40 \\ 40 & 55 \end{pmatrix}, \begin{pmatrix} 55 & -40 \\ -40 & 55 \end{pmatrix}, \begin{pmatrix} 55 & 40 \\ 40 & 55 \end{pmatrix} \right\}. \quad (2.33)$$

On simule un échantillon de 4000 points indépendants, illustré par la figure 2.4(a). Après estimation du modèle de mélange par l'algorithme 2.1 avec $K = 4$, la densité mélange est représentée par la figure 2.4(b).

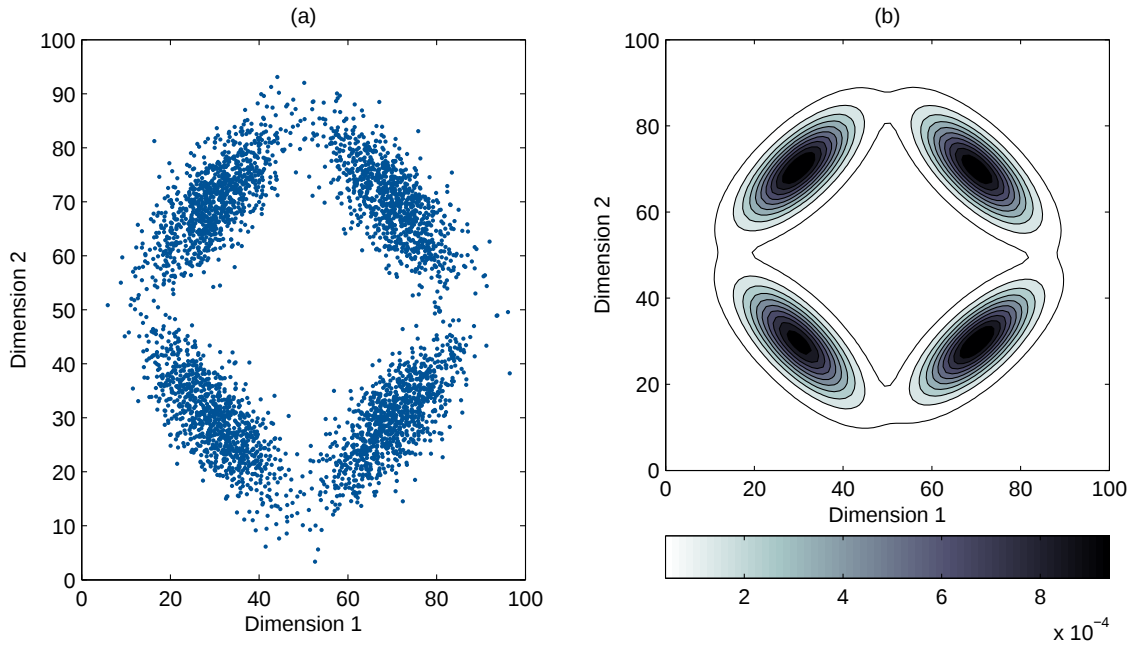


FIGURE 2.4 – Échantillon de 4000 points en deux dimensions générés à partir de quatre lois gaussiennes bivariées (a) : 1000 points par gaussienne. Densité de probabilité d'un mélange à quatre composantes, avec isocontours (b).

Quelques extensions de l'algorithme EM

De nombreux états de l'art ont été proposés dans la littérature. Le lecteur pourra se référer à l'ouvrage de [McLachlan et Krishnan \(2008\)](#) et plus récemment, à la synthèse informelle de [Roche \(2012\)](#). En particulier, on cite quelques extensions en-lignes de l'algorithme EM disponibles dans ([Titterton, 1984](#); [Sato, 2000](#); [Samé et al., 2007](#); [Cappé et Moulines, 2009](#); [Cappé, 2011](#)). L'algorithme CEM ([Celeux et Govaert, 1992](#)), également très performant, ajoute une étape de classification entre les étapes E et M. Enfin, une formulation en-ligne de l'algorithme CEM a été proposée par [Samé et al. \(2005\)](#).

2.3 Diagnostic à partir de données fonctionnelles

Depuis la fin du XX^e siècle, les moyens informatiques permettent de collecter et traiter des quantités massives de données. Dans de nombreux domaines d'application, chaque échantillon est souvent fortement structuré (séquence, arbre, graphe, image). Dans ce contexte, les données sont alors des valeurs discrétisées de fonctions (temporelles ou spatiales) non observables appelées par la suite « données fonctionnelles ». *L'analyse de données fonctionnelles* (Ramsay et Silverman, 2005) présente un panel de méthodes adaptées au traitement de ce type de données, également appelées *courbes*, *trajectoires*, ou parfois *signatures*. Cette section introduit la problématique de traitement de données fonctionnelles et présente notamment un modèle de mélange de régressions polynomiales pour la segmentation, la représentation et la classification de courbes univariées. On illustre cette méthode par une étude sur des données réelles.

La forme des mesures collectées sur les portes, illustrées par la figure 2.5, motive notre choix de les modéliser comme des données fonctionnelles. Pour chaque ouverture/fermeture, on collecte une série temporelle de pressions et de positions sur une fenêtre d'environ 5 secondes. Avec une fréquence de 100 Hz, on observe donc une courbe bivariée d'une longueur d'environ 500 points à chaque fonctionnement d'un battant de porte. Nous présentons une stratégie de détection séquentielle adaptée à ce type de données dans le chapitre 5 qui s'appuie sur l'extension multivariée du modèle de mélange présenté dans la partie 2.3.2.

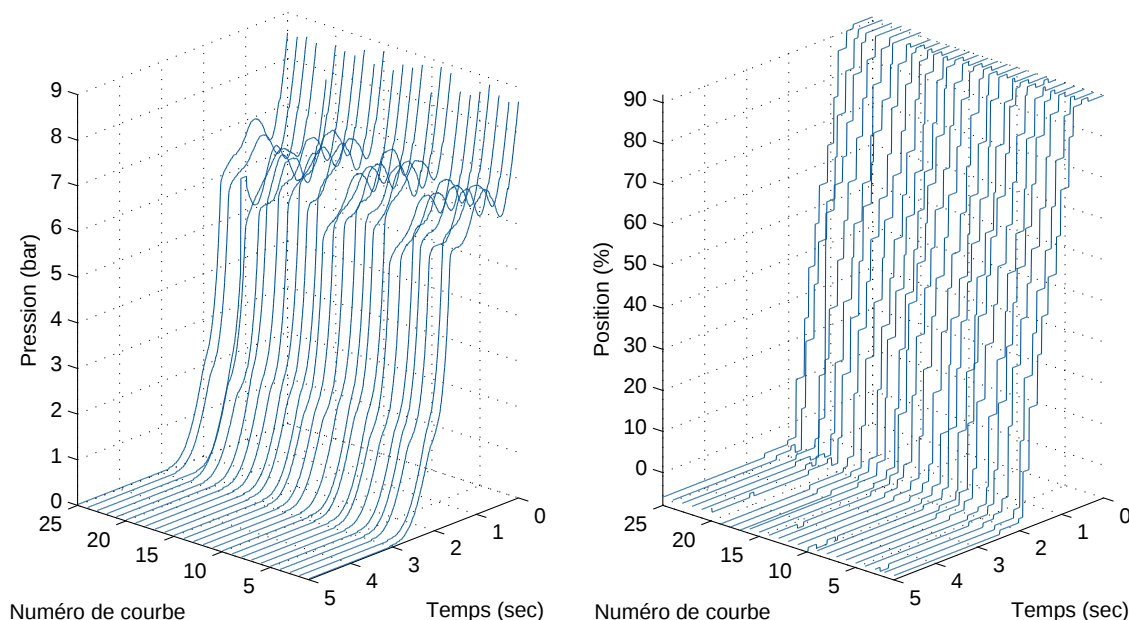


FIGURE 2.5 – Échantillon de 25 courbes pendant des ouvertures de porte.
Pression FAV et position collectées le 30/01/2012 sur le battant avant
(voir la sous-section 5.1.2).

2.3.1 Classification de données fonctionnelles

Le diagnostic par reconnaissance des formes à partir de courbes est habituellement formalisé comme un problème de classification dans un mode supervisé ou non supervisé. L'objectif est donc d'apprendre un modèle afin de prédire la classe d'une courbe observée. L'analyse de données fonctionnelles (Ramsay et Silverman, 2005; Vieu et Ferraty, 2006; Ramsay et al., 2009)¹, suppose que chaque courbe est la réalisation d'une variable aléatoire fonctionnelle univariée $x \in \mathcal{X}$ (section 2.3), ou multivariée $x \in \mathcal{X}^d$ avec $d > 1$ (section 5.2). Chacune de ces données est donc une fonction et $\mathcal{X}$ est un espace de grande dimension (voir infinie) ce qui constitue l'une des principales difficultés liées à ce type de données (Jacques, 2012, chapitre 4). Les problèmes statistiques d'estimation de densité, régression et, en particulier, de classification requièrent des méthodes dédiées car il est établi que les méthodes classiques alimentées par des données multidimensionnelles échouent sur des données fonctionnelles (Ferraty et Romain, 2011).

La première approche consiste à extraire des caractéristiques puis à classer ces dernières dans l'espace euclidien associé. Les méthodes de projection (ACP fonctionnelle,...) permettent de réduire l'espace de recherche et ainsi d'analyser les courbes projetées dans des sous-espaces de plus faibles dimensions. Les méthodes de filtrage (splines, Fourier, ondelettes) projettent chaque courbe sur une base de dimension finie puis analysent les coefficients des bases. Récemment, Meynet (2012) propose une procédure paramétrée par l'estimateur d'une vraisemblance pénalisée (Lasso) et basée sur une projection sur une base d'ondelettes. Dans un problème de classification non supervisée, ces approches peuvent s'avérer inefficaces car elles supposent que les informations discriminantes sont toujours présentes dans ces nouveaux espaces de recherche.

Dans la seconde approche, la discrimination des courbes est obtenue par maximisation de la probabilité a posteriori obtenue à partir de la règle de Bayes où les distributions des courbes sont issues de modèles de régression (splines, RHLP, modèles de Markov). Pour les courbes qui ne sont pas observées sur la même grille temporelle, James et Sugar (2003) ont notamment proposé une méthode de clustering basée sur des splines dont les coefficients suivent un modèle avec effet aléatoire, ce qui peut s'avérer utile lorsque les courbes sont clairsemées. On cite également les travaux de Gaffney (2004) sur un mélange de régressions polynomiales dont l'apprentissage est réalisé par un algorithme de type EM. D'autre part, Jacques (2012) introduit un modèle de mélange pour des courbes multivariées, appelé MFunClust, dont le paramètre est estimé par un algorithme de type EM. Il est à noter que cette méthode utilise une ACP fonctionnelle conditionnellement aux probabilités a posteriori afin d'effectuer la tâche de discrimination dans un sous-espace dédié. Récemment, Chamroukhi et al. (2010) propose des méthodes dédiées à la discrimination de courbes univariées à changement de régime.

1. Le site <http://www.functionaldata.org> regroupe quelques logiciels, exemples et références avec différents niveaux de lecture, liées à l'analyse de données fonctionnelles.

La pierre angulaire des deux approches citées précédemment est la modélisation des courbes. On décrit par la suite une méthode de régression pour la modélisation de courbes multivariées où chaque courbe est décomposable en plusieurs morceaux.

2.3.2 Modèle de mélange de régressions pour la modélisation de courbes à changement de régimes

On reprend ici les travaux de [Chamroukhi \(2010\)](#) sur le *modèle de régression à processus caché (RHLP)*. Le modèle RHLP est un mélange de régressions polynomiales où le passage d'un modèle de régression à un autre est régi par un processus latent. Ce modèle génératif peut être utilisé pour la segmentation, l'interprétation, ou la classification de courbes univariées. L'extension multidimensionnelle de cette modélisation pour des courbes de longueurs différentes sera présentée dans la section 5.2. Dans cette sous-section, nous détaillons la spécification du modèle RHLP univarié et nous démontrons la flexibilité de la transformation logistique associée.

Définition du modèle

Chaque courbe $\mathbf{x} = (x_1, \dots, x_m) \in \mathbb{R}^m$ est associée à un vecteur temps $\mathbf{t} = (t_1, \dots, t_m)$ avec $t_1 < \dots < t_m$; le vecteur temps n'est pas forcément à pas constant. A chaque instant t_j , on admet que l'observation x_j appartient de façon privilégiée à un segment parmi K et chaque segment est décrit par un modèle de régression polynomial d'ordre r . Pour tout t_j , le modèle est défini par :

$$x_j = \beta_{z_j}^T \cdot \mathbf{T}_j + \epsilon_j, \quad \forall j = 1, \dots, m, \quad (2.34)$$

où $z_j \in \{1, \dots, K\}$ représente le segment auquel appartient l'observation x_j , $\epsilon_j \sim \mathcal{N}(0, \sigma_{z_j}^2)$ est un bruit blanc univarié de variance $\sigma_{z_j}^2$, la matrice β_{z_j} est le vecteur des $(r+1)$ coefficients polynomiaux et $\mathbf{T}_j = (1, t_j, \dots, (t_j)^r)^T$ est le vecteur des covariables. La variable latente Z_j modélise l'appartenance de x_j aux différents segments. Cette variable Z_j est supposée être distribuée selon une loi multinomiale $\mathcal{M}(1, \pi_1(t_j; \mathbf{w}), \dots, \pi_K(t_j; \mathbf{w}))$:

$$P(Z_j = k | t_j; \mathbf{w}) = \pi_k(t_j; \mathbf{w}), \quad \forall k \in \{1, \dots, K\}, \quad (2.35)$$

avec π_k , une fonction logistique définie par les équations suivantes :

$$\pi_k(t_j; \mathbf{w}) = \frac{\exp(\mathbf{w}_k^T \mathbf{v}_j)}{1 + \sum_{h=1}^{K-1} \exp(\mathbf{w}_h^T \mathbf{v}_j)}, \quad \forall k \in \{1, \dots, K-1\} \quad (2.36)$$

$$\pi_K(t_j; \mathbf{w}) = 1 / (1 + \sum_{h=1}^{K-1} \exp(\mathbf{w}_h^T \mathbf{v}_j)), \quad (2.37)$$

où $\mathbf{v}_j = (1, t_j, \dots, t_j^u)^T$ est le vecteur de covariable de degré u associé à la transformation logistique et $\mathbf{w}$ est la matrice paramètre de taille $(u+1) \times K$, avec $\mathbf{w}_K = \mathbf{0}$. Le degré u contrôle le nombre de variations de chaque fonction logistique ([Chamroukhi, 2010](#)). Les coefficients π_k pondèrent de façon probabiliste les degrés d'appartenance aux différents segments. Le modèle graphique présenté dans la figure 2.7 résume le modèle RHLP.

Dans notre étude, on se place dans le cas d'une fonction logistique de degré 1 ($u = 1$), de manière à assurer une apparition unique de chaque segment :

$$\pi_k(t_j; \mathbf{w}) = \frac{\exp(w_{k0} + w_{k1}t_j)}{\sum_{h=1}^K \exp(w_{h0} + w_{h1}t_j)}, \quad \forall k \in \{1, \dots, K\}. \quad (2.38)$$

La proposition 2.1 expose le système d'équations pour l'estimation de la matrice paramètre $\mathbf{w}$ de la fonction logistique à partir de la localisation et de la vitesse des transitions entre les segments, pour le cas multi-classes ($K \geq 2$).

Proposition 2.1 (Estimation de la matrice paramètre $\mathbf{w}$)

On suppose que chaque courbe x est composée de K segments contigus. Soient $\gamma = (\gamma_1, \dots, \gamma_{K-1}) \in \mathbb{R}^{K-1}$, le vecteur des points de changement entre les K segments. Alors, il existe une matrice paramètre $\mathbf{w}$ de taille $2 \times K$ telle que

$$w_{k0} = \sum_{h=k}^{K-1} \gamma_h (w_{h+1,1} - w_{h1}) \quad , \quad \forall k \in \{1, \dots, K-1\}. \quad (2.39)$$

On obtient une matrice paramètre unique en ajoutant $K-1$ contraintes supplémentaires. En particulier, les contraintes :

$$w_{k1} = \sum_{h=k}^{K-1} \lambda_h \quad , \quad \forall k \in \{1, \dots, K-1\}, \quad (2.40)$$

où chaque $\lambda_k < 0$ permet de caractériser la vitesse de transition entre le k^e segment et le $(k+1)^e$ segment. Dans ce cas particulier, λ_k définit la pente de π_k en γ_k , et on a $\lambda_k = w_{k1} - w_{k+1,1}$.

Preuve. Par définition, $(\gamma_k)_{1 \leq k \leq K-1}$ est le point de changement entre le k^e segment et le $(k+1)^e$ segment. Les fonctions π_k et π_{k+1} se croisent donc en γ_k , et d'après l'équation 2.38, on a :

$$1 = \frac{\pi_k(\gamma_k; \mathbf{w})}{\pi_{k+1}(\gamma_k; \mathbf{w})} = \frac{\exp(w_{k0} + w_{k1}\gamma_k)}{\exp(w_{k+1,0} + w_{k+1,1}\gamma_k)} \quad (2.41)$$

$$\Leftrightarrow w_{k0} = \gamma_k(w_{k+1,1} - w_{k1}) + w_{k+1,0} \quad , \quad \forall k \in \{1, \dots, K-1\}. \quad (2.42)$$

De même, les fonctions π_{K-1} et π_K se croisent en γ_{K-1} . On obtient alors :

$$w_{K-1,0} = \gamma_{K-1}(0 - w_{K-1,1}), \quad (2.43)$$

car $w_{K0} = w_{K1} = 0$. □

On illustre la flexibilité de la transformation logistique par un exemple simulé avec des transitions abruptes ou souples entre 5 segments. Chaque courbe de la figure 2.6 représente la probabilité $(\pi_k)_{1 \leq k \leq 5}$ d'appartenance au k^e segment. Les vecteurs γ et λ contrôlent la localisation et la vitesse des transitions entre les 5 segments :

$$\gamma = (5, 10, 15, 20) \quad \text{et} \quad \lambda = (-2, -1, -0.5, -0.25). \quad (2.44)$$

A partir de la proposition 2.1, on obtient le vecteur paramètre de la fonction logistique :

$$\mathbf{w} = \begin{pmatrix} 32.50 & 22.50 & 12.50 & 5.00 & 0 \\ -3.75 & -1.75 & -0.75 & -0.25 & 0 \end{pmatrix}. \quad (2.45)$$

qui permet d'obtenir les fonctions logistiques représentées par la figure 2.6.

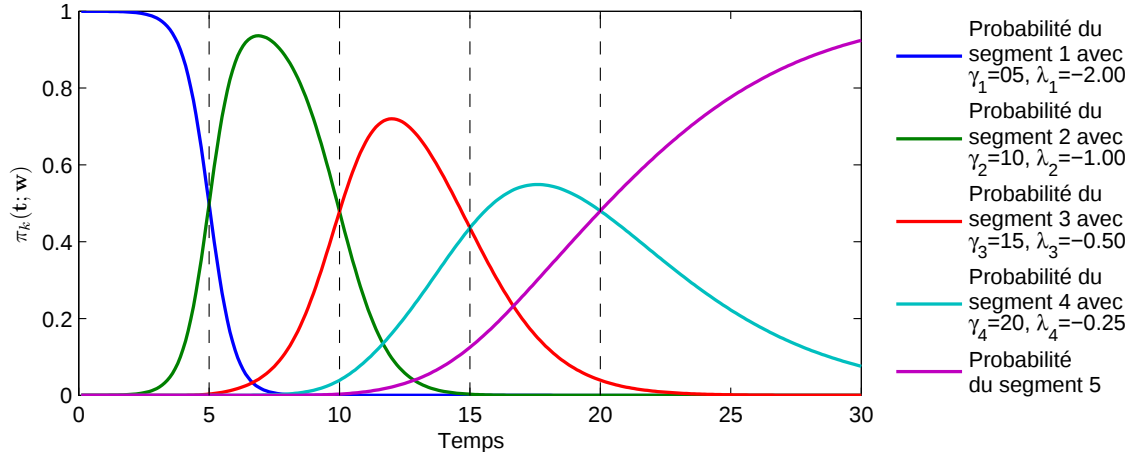


FIGURE 2.6 – Exemple de fonction logistique qui régit les transitions entre 5 segments. Les paramètres γ et λ contrôlent la localisation et la vitesse des transitions.

Le modèle peut être étendu simplement à un ensemble de t courbes $\{\mathbf{x}_1, \dots, \mathbf{x}_t\}$, où chaque courbe $\mathbf{x}_i$ est une séquence $(x_{i1}, \dots, x_{im})$ d'observations modélisées par l'équation 2.34. On fait donc l'hypothèse que :

$$x_{ij} = \beta_{zij}^T \cdot \mathbf{T}_j + \epsilon_{ij}, \quad \forall (i, j) = \{1, \dots, t\} \times \{1, \dots, m\}. \quad (2.46)$$

Estimation des paramètres

On estime le vecteur paramètre de ce modèle $\theta = (\mathbf{w}, \beta_1, \dots, \beta_K, \sigma_1^2, \dots, \sigma_K^2)$ par maximisation de la vraisemblance via un algorithme EM (Dempster et al., 1977) dédié. Conditionnellement au k^e modèle de régression, on fait l'hypothèse que l'observation x_{ij} suit une loi normale de moyenne $\beta_{zij}^T \mathbf{T}_j$ et de variance σ_{zij}^2 :

$$p(x_{ij}|Z_{ij} = k, t_j; \theta_k) = \mathcal{N}(x_{ij}; \beta_{zij}^T \mathbf{T}_j, \sigma_{zij}^2), \quad \forall (i, j) = \{1, \dots, t\} \times \{1, \dots, m\}. \quad (2.47)$$

A chaque pas de temps t_j , l'observation x_{ij} est donc distribuée par un modèle de mélange (sous-section 2.2.3) à K composantes :

$$p(x_{ij}|t_j; \theta) = \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2). \quad (2.48)$$

En supposant l'indépendance des courbes $\mathbf{x}_i$ conditionnellement aux $\mathbf{t}_i$, on cherche à maximiser la log-vraisemblance :

$$\mathcal{L}(\theta; \mathbf{x}_1, \dots, \mathbf{x}_t, \mathbf{t}) = \sum_{i=1}^t \sum_{j=1}^m \log \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2). \quad (2.49)$$

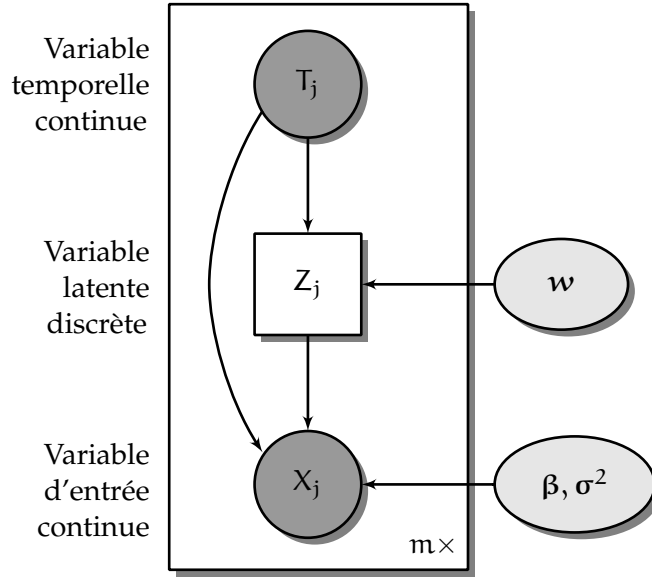


FIGURE 2.7 – Représentation graphique du modèle probabiliste RHLP.

Comme décrit dans la sous-section 2.2.3, l'algorithme EM estime le paramètre θ en maximisant itérativement l'espérance de la log-vraisemblance sur la variable latente Z . A chaque itération (q), on cherche à maximiser la quantité suivante :

$$\mathbb{E}_Z \left[\mathcal{L}(\theta; \mathbf{X}, \mathbf{T}) | \mathbf{x}, \mathbf{t}; \theta^{(q)} \right] = \underbrace{\mathbb{E}_Z \left[\log p(\mathbf{X}, \mathbf{Z} | \mathbf{T}; \theta) | \mathbf{x}, \mathbf{t}; \theta^{(q)} \right]}_{Q(\theta, \theta^{(q)})} - \underbrace{\mathbb{E}_Z \left[\log p(\mathbf{Z} | \mathbf{X}, \mathbf{T}; \theta) | \mathbf{x}, \mathbf{t}; \theta^{(q)} \right]}_{H(\theta, \theta^{(q)})}. \quad (2.50)$$

La fonction H décroît au cours des itérations d'après l'inégalité de Jensen (McLachlan et Krishnan, 2008). Il suffit donc de maximiser la fonction Q pour faire croître la vraisemblance. Cette fonction Q est l'espérance de la log-vraisemblance complétée :

$$\begin{aligned} \mathcal{L}_c(\theta; \mathbf{x}, \mathbf{z}, \mathbf{t}) &= \sum_{i=1}^t \sum_{j=1}^m \log p(x_{ij}, z_{ij} | t_j; \theta) \\ &= \sum_{i=1}^t \sum_{j=1}^m \sum_{k=1}^K z_{ijk} \cdot \log \left[\pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2) \right] \end{aligned} \quad (2.51)$$

avec $z_{ijk} = \begin{cases} 1 & \text{si } z_{ij} = k \\ 0 & \text{sinon} \end{cases}.$

On écrit alors la fonction Q comme l'espérance de la log-vraisemblance complétée conditionnellement aux données observées et au paramètre courant (voir équation 2.23) :

$$Q(\theta, \theta^{(q)}) = \sum_{i,j,k} \underbrace{\mathbb{E}_Z \left(Z_{ijk} | \mathbf{x}, \mathbf{t}; \theta^{(q)} \right)}_{\tau_{ijk}^{(q)}} \cdot \left(\log \pi_k(t_j; \mathbf{w}) + \log \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2) \right), \quad (2.52)$$

où $\tau_{ijk}^{(q)}$ représente la loi a posteriori qui mesure le degré d'appartenance de l'observation x_{ij} au k^e segment.

L'étape E consiste à calculer les probabilités a posteriori à partir des paramètres disponibles à l'itération courante :

$$\begin{aligned}\tau_{ijk}^{(q)} &= P(Z_{ij} = k | x_{ij}, t_j; \theta^{(q)}) \\ &= \frac{\pi_k(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2)^{(q)}}{\sum_{h=1}^K \pi_h(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_h^T \mathbf{T}_j, \sigma_h^2)^{(q)}}.\end{aligned}\quad (2.53)$$

D'autre part, l'étape M consiste à mettre à jour chaque paramètre en maximisant la fonction Q . D'après l'équation 2.52, il convient de maximiser deux fonctions Q_1 et Q_2 , respectivement pour la mise à jour des paramètres $\mathbf{w}$ et $(\beta_k, \sigma_k^2)_{1 \leq k \leq K}$:

$$\begin{aligned}\theta^{(q+1)} &= \arg \max_{\theta} \{Q(\theta, \theta^{(q)})\} \\ &= \arg \max_{\theta} \left\{ \underbrace{\sum_{i,j,k} \tau_{ijk}^{(q)} \cdot \log \pi_k(t_j; \mathbf{w})}_{Q_1(\theta, \theta^{(q)})} + \underbrace{\sum_{i,j,k} \tau_{ijk}^{(q)} \cdot \log \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2)}_{Q_2(\theta, \theta^{(q)})} \right\}.\end{aligned}\quad (2.54)$$

On détaille brièvement par la suite, la mise à jour des différents paramètres.

Mise à jour de la matrice $\mathbf{w}$

La matrice paramètre $\mathbf{w}$ de la transformation logistique est mise à jour via l'algorithme multi-classes *Iterative Reweighted Least Squares* (IRLS) (Green, 1984; Chamroukhi, 2010) afin de maximiser la fonction $Q_1(\theta, \theta^{(q)})$. A chaque itération (c), cette stratégie de type Newton-Raphson met à jour le paramètre $\mathbf{w}$:

$$\mathbf{w}^{(c+1)} = \mathbf{w}^{(c)} - \left[\nabla_{\mathbf{w}}^2 Q_1(\theta^{(c)}, \theta^{(q)}) \right]^{-1} \cdot \nabla_{\mathbf{w}} Q_1(\theta^{(c)}, \theta^{(q)}), \quad (2.55)$$

où $\nabla_{\mathbf{w}} Q_1$ et $\nabla_{\mathbf{w}}^2 Q_1$ représentent respectivement le vecteur gradient et la matrice hessienne de la fonction Q_1 . On peut montrer que la fonction Q_1 est concave et différentiable deux fois. L'algorithme IRLS converge itérativement vers un optimum global jusqu'à une stabilisation de la norme (euclidienne) sur son gradient : $\|\nabla_{\mathbf{w}} Q_1\|$. A partir de l'équation 2.54, on peut calculer les dérivées partielles premières de la fonction Q_1 :

$$\frac{\partial}{\partial w_k} Q_1(\theta, \theta^{(q)}) = \sum_{i,j} \sum_{h=1}^K \frac{\tau_{ijk}^{(q)}}{\pi_h(t_j; \mathbf{w})} \cdot \frac{\partial}{\partial w_k} \pi_h(t_j; \mathbf{w}) \quad (2.56)$$

$$= \sum_{i,j} \sum_{h=1}^K \frac{\tau_{ijk}^{(q)}}{\pi_h(t_j; \mathbf{w})} \cdot \pi_h(t_j; \mathbf{w}) [\delta_{hk} - \pi_k(t_j; \mathbf{w})] \mathbf{v}_j \quad (2.57)$$

$$= \sum_{i,j} [\tau_{ijk}^{(q)} - \pi_k(t_j; \mathbf{w})] \mathbf{v}_j \quad \forall k \in \{1, \dots, K\}, \quad (2.58)$$

où δ_{hk} est le symbole de Kronecker et $\mathbf{v}_j = (1, t_j)^T$ est le vecteur de covariables de la logistique de degré 1.

De même, on peut calculer les dérivées partielles secondes : $\forall (k, h) \in \{1, \dots, K\}^2$,

$$\frac{\partial^2}{\partial \mathbf{w}_k \partial \mathbf{w}_h^T} Q_1(\boldsymbol{\theta}, \boldsymbol{\theta}^{(q)}) = \frac{\partial}{\partial \mathbf{w}_h^T} \sum_{i,j} \left[\tau_{ijk}^{(q)} - \pi_k(t_j; \mathbf{w}) \right] \mathbf{v}_j \quad (2.59)$$

$$= - \sum_{i,j} \frac{\partial}{\partial \mathbf{w}_h^T} \pi_k(t_j; \mathbf{w}) \mathbf{v}_j \quad (2.60)$$

$$= - \sum_{i,j} \pi_k(t_j; \mathbf{w}) [\delta_{kh} - \pi_h(t_j; \mathbf{w})] \mathbf{v}_j^T \cdot \mathbf{v}_j. \quad (2.61)$$

L'algorithme IRLS est adapté à la régression logistique multi-classes et permet d'estimer un maximum global (définition A.5). En pratique, cet algorithme itératif maximise la fonction Q_1 en moins d'une dizaine d'itérations. Par ailleurs, la fonction Q_2 peut être maximisée analytiquement par rapport aux paramètres $(\beta_k)_{1 \leq k \leq K}$ et $(\sigma_k^2)_{1 \leq k \leq K}$. La mise à jour de ces $2K$ paramètres restants est décrite par la suite.

Mise à jour de la matrice β

L'annulation de la dérivée partielle de la fonction Q_2 par rapport à β_k (équation 2.62) conduit à résoudre un problème des *moindres carrés pondérés* par $\tau_{ijk}^{(q)}$ (équation 2.63). L'estimateur qui annule $\frac{\partial}{\partial \beta_k} Q_2$ est donné dans l'équation 2.64 et s'écrit :

$$\beta_k^{(q+1)} = \arg \max_{\beta_k} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \log \mathcal{N}(x_{ij}; \beta_k^T \mathbf{T}_j, \sigma_k^2) \quad (2.62)$$

$$= \arg \min_{\beta_k} \frac{1}{\sigma_k^2} \sum_{i,j} \tau_{ijk}^{(q)} \cdot (x_{ij} - \beta_k^T \mathbf{T}_j)^2 \quad (2.63)$$

$$= \left(\sum_{i,j} \tau_{ijk}^{(q)} \cdot \mathbf{T}_j \cdot \mathbf{T}_j^T \right)^{-1} \cdot \left(\sum_{i,j} \tau_{ijk}^{(q)} \cdot \mathbf{T}_j \cdot x_{ij} \right). \quad (2.64)$$

Mise à jour du vecteur σ^2

Enfin, l'annulation de la dérivée partielle de la fonction Q_2 par rapport à la variance σ_k (équation 2.65) conduit à l'estimation de la variance d'une loi normale multivariée et pondérée par les probabilités a posteriori (équation 2.66). L'estimateur qui annule la dérivée partielle $\frac{\partial}{\partial \Sigma_k} Q_2$ est donné dans l'équation 2.67 et s'écrit :

$$\sigma_k^{2(q+1)} = \arg \max_{\sigma_k^2} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \log \mathcal{N}(x_{ij}; \beta_k^{T(q+1)} \mathbf{T}_j, \sigma_k^2) \quad (2.65)$$

$$= \arg \min_{\sigma_k^2} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \left[2 \log \sigma_k + \frac{1}{\sigma_k^2} (x_{ij} - \beta_k^{T(q+1)} \mathbf{T}_j)^2 \right] \quad (2.66)$$

$$= \frac{1}{\sum_{i,j} \tau_{ijk}^{(q)}} \cdot \sum_{i,j} \tau_{ijk}^{(q)} \cdot [x_{ij} - \beta_k^{T(q+1)} \mathbf{T}_j]^2. \quad (2.67)$$

L'algorithme EM répète les étapes E et M jusqu'à convergence de la log-vraisemblance, donnée dans l'équation 2.49. Le pseudo-code 2.2 décrit le modèle RHLP dans une formulation batch (hors-ligne) en mode non-supervisé. L'extension multidimensionnelle de ce modèle sera décrite en mode non-supervisé et semi-supervisé dans la section 5.2, ainsi que sa formulation matricielle. La complexité dans le pire des cas du modèle RHLP univarié (Samé et al., 2011) est de $\mathcal{O}(I_{EM} I_{IRLS} t m K^3 r^3)$, où I_{EM} et I_{IRLS} correspondent aux nombres d'itérations des algorithmes EM et IRLS. L'algorithme EM converge vers un optimum local et la qualité de l'estimateur $\hat{\theta}$ dépend de son initialisation (voir sous-section 2.2.3). L'algorithme est lancé plusieurs fois avec une initialisation aléatoire puis on conserve le paramètre (après convergence) avec la plus grande vraisemblance.

Algorithme 2.2 : Pseudo-code du modèle RHLP univarié.

Entrées : Matrice $\mathbf{x}$ des t courbes univariées de longueur m , nombre K de segments et degré r des polynômes pour chaque courbe

```

// Initialisation aléatoire de  $\theta^{(0)}$  si non fourni
1  $q \leftarrow 0$ 
2 Répéter
    // Étape E : calcul de la loi a posteriori (éq. 2.53)
3 Pour  $(i, j, k) \in \{1, \dots, t\} \times \{1, \dots, m\} \times \{1, \dots, K\}$  faire
4     
$$\tau_{ijk}^{(q)} \leftarrow \frac{\pi_k(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_k^{T(q)} \mathbf{T}_j, \sigma_k^{2(q)})}{\sum_{h=1}^K \pi_h(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_h^{T(q)} \mathbf{T}_j, \sigma_h^{2(q)})} \quad (2.68)$$

    // Étape M : mise à jour des paramètres
    // IRLS : paramètre de la logistique (éq. 2.55)
5     
$$\mathbf{w}^{(q+1)} \leftarrow \arg \max_{\mathbf{w}} \sum_{j=1}^m \sum_{k=1}^K \left( \sum_{i=1}^t \tau_{ijk}^{(q)} \right) \cdot \log \pi_k(t_j; \mathbf{w}) \quad (2.69)$$

6 Pour  $k \in \{1, \dots, K\}$  faire
7     // Matrices des coefficients polynomiaux (éq. 2.64)
    
$$\beta_k^{(q+1)} \leftarrow \left[ \sum_{j=1}^m \left( \sum_{i=1}^t \tau_{ijk}^{(q)} \right) \mathbf{T}_j \mathbf{T}_j^T \right]^{-1} \cdot \left[ \sum_{j=1}^m \mathbf{T}_j \left( \sum_{i=1}^t \tau_{ijk}^{(q)} \cdot x_{ij} \right) \right] \quad (2.70)$$

8     // Matrices de variance-covariance (éq. 2.67)
    
$$\sigma_k^{2(q+1)} \leftarrow \left( \sum_{i=1}^t \sum_{j=1}^m \tau_{ijk}^{(q)} \right)^{-1} \cdot \sum_{i=1}^t \sum_{j=1}^m \tau_{ijk}^{(q)} (x_{ij} - \beta_k^{T(q+1)} \mathbf{T}_j)^2 \quad (2.71)$$

9      $q \leftarrow q + 1$ 
10 Jusqu'à convergence;
Sortie : Paramètre  $\hat{\theta}$  estimé au sens du maximum de vraisemblance.

```

2.3.3 Cas d'étude réel

A titre d'illustration, on présente un cas d'étude descriptif sur le cours du pétrole de janvier 1974 à octobre 2012. Six séries temporelles (prix mensuels en \$) ont été obtenues à partir de l'Agence d'Information sur l'Energie (*Energy Information Administration* ou *EIA*) et ces six courbes sont représentées dans la figure 2.8. Chaque courbe est longue de 466 points et ces données seront analysées en échelle logarithme décimal (Zeileis et al., 2003).

Afin de choisir un modèle m réalisant un compromis entre simplicité et précision (dilemme biais/variance), les critères BIC et ICL de sélection de modèles sont utilisés (voir la sous-section 2.2.2). Pour calculer ces deux critères, il est nécessaire de fixer ν_m , la dimension du paramètre $\hat{\theta}_m$. Dans le cas du modèle de régression RHLP monodimensionnel, cette dimension s'écrit :

$$\nu_m = \underbrace{2(K-1)}_w + \underbrace{(r+1)K}_\beta + \underbrace{K}_{\sigma^2}. \quad (2.72)$$

La figure 2.9 représente les critères BIC, ICL et le temps CPU après cent lancements (avec initialisation aléatoire) du modèle RHLP sur les six courbes disponibles. Plusieurs configurations conjointes du nombre des segments $2 \leq K \leq 12$ et de l'ordre des polynômes $0 \leq r \leq 6$ ont été mises en compétition dans un mode non supervisé. Un ordre $r = 3$ a été mis en évidence par le critère BIC, le critère ICL semble favoriser les modèles complexes de cette étude. Et on obtient un nombre $K = 8$ de segments en maximisant les critères BIC et ICL en fonction d'un ordre 3. Le modèle à huit segments ($K = 8$) d'ordre trois ($r = 3$) a donc été retenu pour cette étude. Il est à noter que le temps CPU augmente avec le nombre de segments car la dimension du paramètre du modèle est factorisable par K , le nombre de segments (équation 2.72).

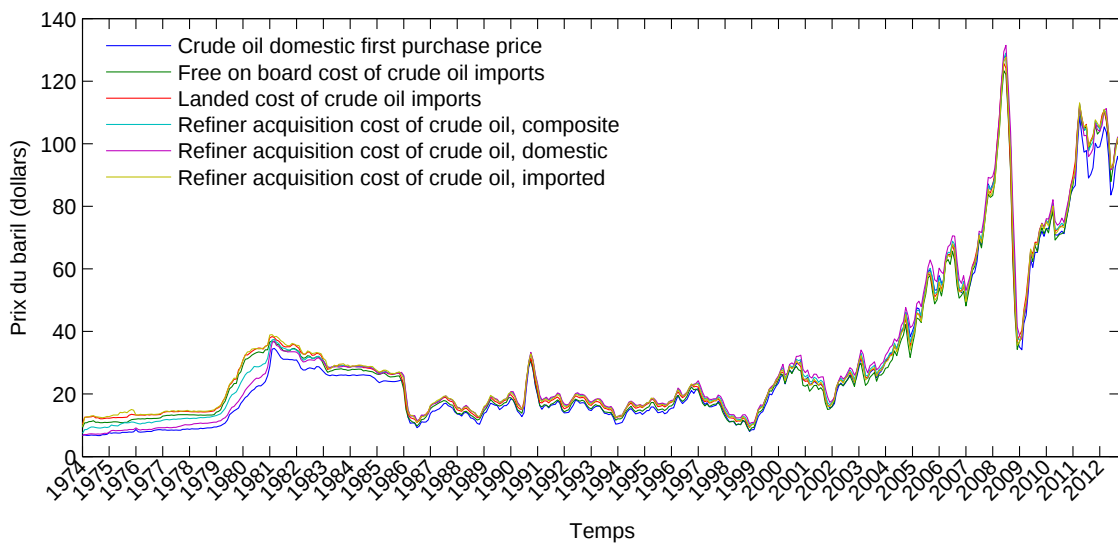


FIGURE 2.8 – Cours du pétrole de janvier 1974 à octobre 2012.

Sources : U.S. Energy Information Administration.

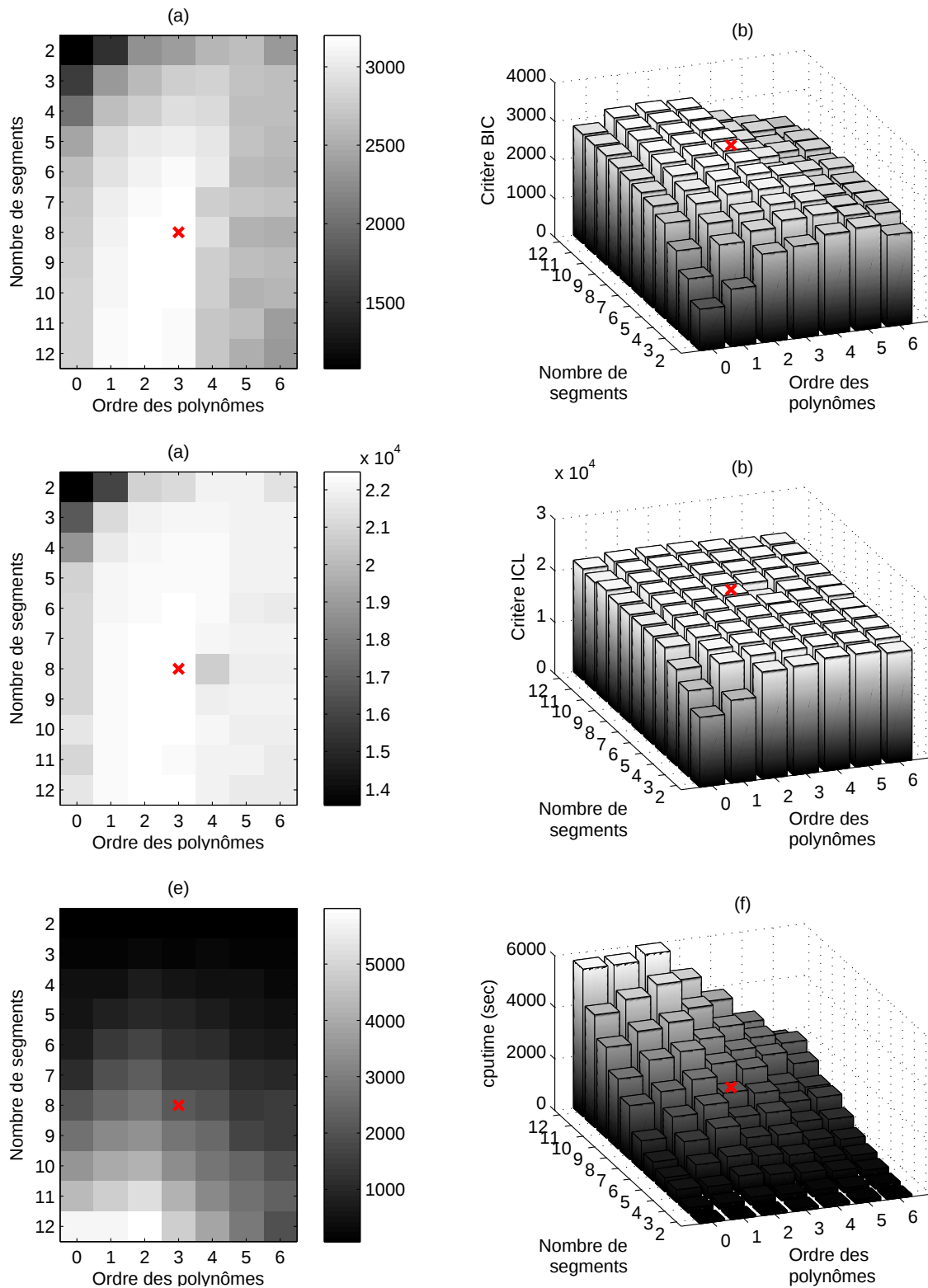


FIGURE 2.9 – Sélection du modèle de régression univarié à partir des critères BIC (a,b) et ICL (c,d), et temps CPU (e,f).

Le modèle choisi (x) est composé de 8 segments dont chaque polynôme est d'ordre 3.

La figure 2.10 présente la segmentation automatique obtenue par le modèle RHLF composé de huit segments contigus de degré trois. Le prix du pétrole est volatile et son évolution dépend de facteurs économiques et géopolitiques complexes (Amenc et al., 2008). Les sept points de changement détectés peuvent être brièvement expliqués par les événements suivants :

Juin 1979, *Second choc pétrolier* : L'été 1979 marque le début de la révolution iranienne avec le départ du Shah qui sera suivi de la guerre Iran-Irak.

Mars 1986, *Contre-choc pétrolier* : Une baisse brutale du prix du pétrole (jusqu'à 10\$) est expliquée par une surproduction de pétrole.

Janvier 1991, *Opération tempête du désert* : Cette opération américaine a mis fin à la seconde guerre du Golfe.

Octobre 1995 : La production de pétrole américain dépasse son importation.

Février 1999 : Baisse brutale du cours du pétrole due à la crise financière asiatique.

Mars 2002, « *Troisième choc pétrolier* » : Début d'un nouveau choc pétrolier à cause d'une forte spéculation financière selon l'O.P.E.P.

Novembre 2008, *Crise financière* : La crise bancaire et financière de l'automne 2008 provoque un « krach » financier et notamment une baisse brutale du cours du pétrole.

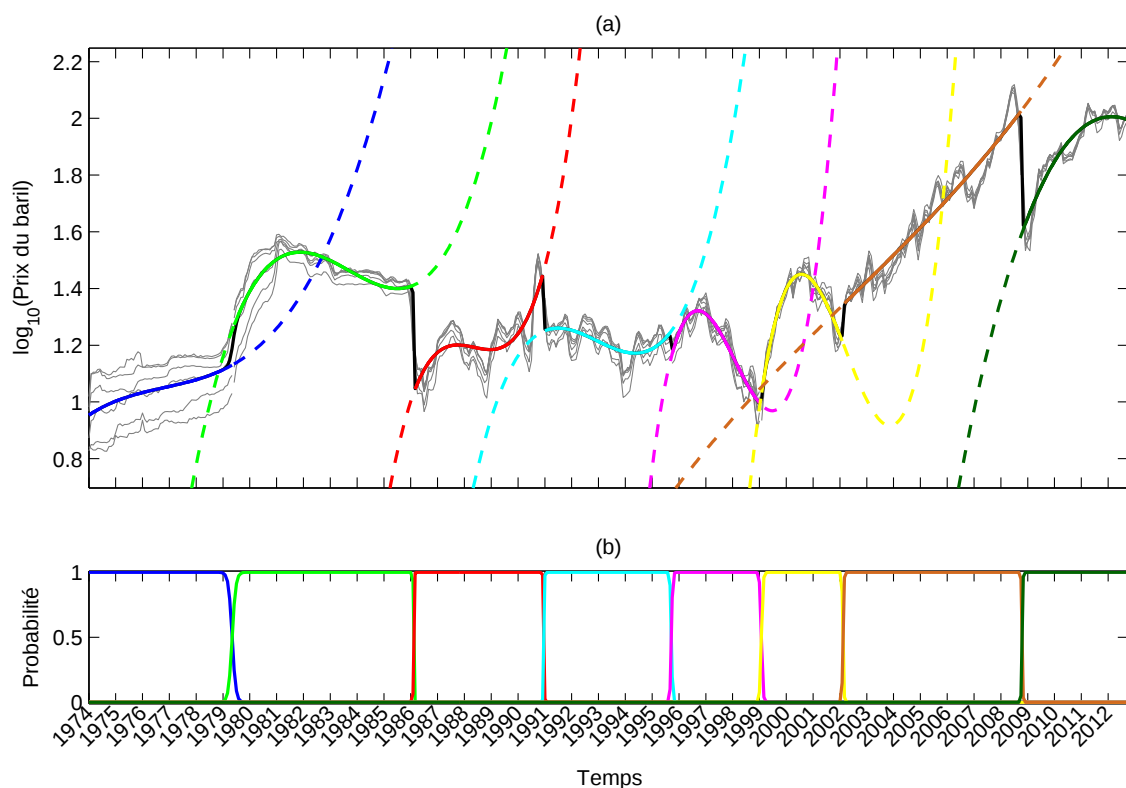


FIGURE 2.10 – Segmentation des courbes en huit morceaux de degré trois.
Représentation des huit polynômes qui composent la courbe moyenne (a).
Probabilités logistiques avec transitions entre segments, en fonction du temps (b).

2.4 Conclusion

Après une synthèse des principales politiques de maintenance d'un système industriel, ce chapitre présente la direction de nos travaux dans le cadre de la maintenance conditionnelle. Une maintenance efficace s'appuie sur une procédure régulière de contrôle non destructif des systèmes critiques, elle-même basée sur un processus de diagnostic. L'approche à base de reconnaissance des formes a été retenue pour nos travaux. Ne nécessitant pas de modèle physique a priori du système, cette approche permet d'extraire de la connaissance d'une base historisée de données mesurées sur le système à surveiller. Les paradigmes associés à l'apprentissage d'un modèle statistique qui est estimé à partir de la base de données, sont passés en revue dans ce chapitre. En particulier, on décrit les modèles de mélange qui permettent de modéliser les données peu supervisées dans des groupes homogènes.

Nous avons ensuite présenté le diagnostic alimenté par des données fonctionnelles. Ce choix est motivé par la forme des données collectées sur les systèmes de porte d'autobus. Un mélange de régression adapté à la modélisation de ce type de données (également appelées courbes) est détaillé. On démontre notamment la flexibilité de ce modèle RHLP basé sur une transformation logistique. Ce modèle est illustré par un cas d'étude sur le cours du pétrole de 1974 à 2012 et met en évidence quelques points de changement significatifs. Le vecteur paramètre issu de cette modélisation résume simplement une courbe ou un ensemble de courbes et nous aborderons au chapitre 5 son utilisation à des fins de détection ou d'identification. L'extension multidimensionnelle de ce modèle de régression sera présentée à cette occasion.

Le prochain chapitre présente les principaux outils nécessaires à la mise en œuvre de la première étape de diagnostic d'un système : la détection, ou comment détecter le plus tôt possible que le système évolue vers un comportement anormal. Le cadre de la détection séquentielle à base d'hypothèses est particulièrement détaillé. Il est montré que certains détecteurs vérifient des critères d'optimalité sous certaines conditions. On met en évidence des problèmes communs entre plusieurs communautés statistiques et notamment en reconnaissance des formes. Enfin, on décrit quelques approches récentes pour la détection de changements dans une séquence de données fonctionnelles.

Chapitre 3

Détection statistique de défaillances

Sommaire

3.1	Formalisation du problème de la détection de points de changement	54
3.1.1	Quelques généralités sur les tests d'hypothèses	55
3.1.2	Détection séquentielle de changements sous forme de test d'hypothèses	56
3.2	Algorithmes classiques de détection à base de test d'hypothèses	58
3.2.1	Quelques tests pour le problème de détection séquentielle	58
3.2.2	Résultats classiques d'optimalité d'un détecteur	64
3.2.3	Contrainte de fausses alarmes et estimation du seuil de détection	67
3.2.4	Évaluation d'un détecteur	70
3.3	Cartes de contrôle et autres tests séquentiels	72
3.3.1	Principales cartes de contrôle multivariées	72
3.3.2	Quelques extensions à base de test séquentiel	78
3.4	Détection à base de reconnaissance des formes	80
3.4.1	Détection d'anomalies et classification mono-classe	80
3.4.2	Détection de changements dans des modèles markoviens	82
3.4.3	Détection de changements par la sélection de modèles	84
3.4.4	Détection de changements dans une séquence de données fonctionnelles	85
3.5	Conclusion	87

3.1 Formalisation du problème de la détection de points de changement

L'analyse des points de changement, ou *change-point analysis*, dans un processus stochastique est une problématique importante et répandue dans plusieurs domaines statistiques. Elle est également connue sous les dénominations : segmentation, changement de structure/de régime, etc. Ce chapitre est dédié au problème de la détection de points de changement dans une série temporelle collectée séquentiellement ou non. L'hypothèse de base de cette problématique est que les propriétés statistiques du processus sont différentes avant et après l'instant du point de changement, ce qui se traduit par une modification dans la distribution des données observées. Le changement soudain dans la dynamique de la séquence d'observations peut concerner par exemple la moyenne, la variance, ou d'autres statistiques, et signale la plupart du temps un changement d'état du système observé. Il est à noter que nous allons principalement considérer le problème de la détection d'un seul point de changement dans ce chapitre. En effet dans un contexte séquentiel (décision dès l'arrivée d'une nouvelle donnée), le détecteur pourra être réinitialisé après chaque nouvelle détection. [Killick et al. \(2012\)](#) ont proposé une plateforme qui regroupe de nombreuses publications et logiciels liés à ce type de problème ([Martin et Doncarli, 1998](#)). Les chapitres 4 et 5 du mémoire abordent ce problème d'un point de vue pratique, respectivement pour la surveillance des systèmes de freinage et de porte.

Dans le contexte de données collectées séquentiellement, l'objectif est de détecter tout changement le plus tôt possible (faible délai de détection) tout en garantissant un minimum de fausse détection (faible taux de fausses alarmes). Plusieurs articles et ouvrages sont consacrés à ce sujet ([Isermann, 1984](#); [Basseville, 1988](#); [Basseville et Nikiforov, 1993](#); [Lai, 1995](#); [Basseville et Cordier, 1996](#); [Lai, 2001](#)) dont un nouvel ouvrage qui sera prochainement publié par [Tartakovsky et al. \(Livre en préparation\)](#). On peut également citer les articles de [Tartakovsky et Moustakides \(2010\)](#) et [Polunchenko et Tartakovsky \(2012\)](#) qui dressent eux aussi un état de l'art des méthodes de détection séquentielle de points de changement.

Ce chapitre vise à dresser un panorama de différentes techniques issues de la théorie des tests d'hypothèses et de l'apprentissage statistique (chapitre 2.2.2). On cherche à établir un parallèle entre ces approches souvent issues de communautés différentes. Après quelques généralités sur les tests d'hypothèses dans la sous-section 3.1.1, on commence par introduire les deux cadres de la détection séquentielle à base des tests d'hypothèses dans la sous-section 3.1.2. Dans la section 3.2, on présente quelques tests séquentiels de référence ainsi que les principaux critères d'optimalité d'un détecteur. Puis, on présente une stratégie d'estimation du seuil de détection, pierre angulaire des tests séquentiels, basée sur le taux théorique de fausses alarmes. La section 3.3 présente quelques algorithmes dérivés des tests séquentiels, dont les principales cartes de contrôle multivariées. Un état de l'art des stratégies d'estimation des points de changement par des cartes de

contrôle est notamment dressé par [Amiri et Allahyari \(2012\)](#). Enfin, la section 3.4 décrit quelques approches de type reconnaissance des formes (RdF) pour le problème de détection de nouveautés, de segmentation à l'aide de modèles markoviens et de sélection de modèles. Enfin, quelques méthodes de détection appliquées à des données fonctionnelles sont présentées dans la section 3.4.4.

3.1.1 Quelques généralités sur les tests d'hypothèses

A chaque instant t , on suppose qu'un vecteur $x \in \mathbb{R}^d$ est observé, les coordonnées de ce vecteur pouvant être une vitesse, des coordonnées GPS, une mesure de pression,... Ce vecteur est considéré comme la réalisation d'une variable aléatoire X . En outre, on suppose que x est distribuée suivant une loi $\mathcal{P}_\theta$, où $\theta \in \Theta$ est un vecteur paramètre identifiable¹. On suppose également que la loi (ou distribution) $\mathcal{P}_\theta$ admet une densité de probabilité, une espérance théorique et une variance, notées respectivement $p(x; \theta)$, $\mathbb{E}(x; \theta)$ et $\text{Var}(x; \theta)$.

On commence par définir les concepts d'hypothèse et de test d'hypothèses dans les définitions 3.1 et 3.2.

Définition 3.1 (Hypothèse statistique $\mathcal{H}$)

Pour tout entier $k \geq 0$, une hypothèse $\mathcal{H}_k = \{\theta \in \Theta_k\}$ est un sous-ensemble de Θ . On dit que l'hypothèse $\mathcal{H}_k$ est simple lorsque l'ensemble Θ_k contient un seul élément, et l'hypothèse est dite composite sinon. L'hypothèse $\mathcal{H}_k$ est vraie si $x \sim \mathcal{P}_\theta$ où $\theta \in \Theta_k$.

Définition 3.2 (Test d'hypothèses)

Un test δ à K hypothèses est une application surjective de la forme suivante :

$$\delta : \begin{array}{ccc} \mathbb{R}^d & \rightarrow & \{\mathcal{H}_0, \dots, \mathcal{H}_K\} \\ x & \mapsto & \mathcal{H}_k \end{array} . \quad (3.1)$$

On notera que dans ce cas, Θ est partitionné en un ensemble fini de $K + 1$ sous-ensembles disjoints :

$$\Theta = \bigcup_{k=0}^K \Theta_k, \quad \text{avec } \Theta_k \cap \Theta_l = \emptyset \text{ et } 0 \leq k < l \leq K. \quad (3.2)$$

Un test d'hypothèses est donc une règle de décision qui définit également une partition sur $\mathbb{R}^d$ de $K + 1$ parties disjointes. Si l'hypothèse $\mathcal{H}_k$ est vraie, alors sa partie antécédente dans $\mathbb{R}^d$ est appelée région d'acceptation.

Dans le cas d'hypothèses simples, la qualité d'un test par ses risques est spécifiée par la définition 3.3.

1. Le paramètre θ est identifiable si l'application $\theta \mapsto \mathcal{P}_\theta$ est injective. En d'autres termes, si on a l'équivalence : $\mathcal{P}_{\theta_1} \neq \mathcal{P}_{\theta_2}$ si et seulement si $\theta_1 \neq \theta_2$.

Définition 3.3 (Risques et Probabilité de fausse alarme)

Pour un test δ , le risque de $(k + 1)^{\text{e}}$ espèce s'écrit :

$$\alpha_k(\delta) = P(\delta(x) \neq \mathcal{H}_k \mid \mathcal{H}_k \text{ vraie}), \quad (3.3)$$

Le risque de 1^{re} espèce est aussi appelé probabilité de fausse alarme. Dans le cas d'un test binaire, cette quantité désigne la probabilité $P(\mathcal{H}_1 \mid \mathcal{H}_0)$.

3.1.2 Détection séquentielle de changements sous forme de test d'hypothèses

La détection séquentielle suppose qu'on observe une séquence (ou suite) de données $x_1, \dots, x_t$. Supposons que les observations $x_1, \dots, x_{t_c-1}$ suivent la distribution $\mathcal{P}_{\theta_0}$ et que $x_{t_c}, \dots, x_t$ suivent la distribution $\mathcal{P}_{\theta_1}$. On appelle t_c , le *temps de changement* (ou *instant de rupture*) à partir duquel la loi des observations change de $\mathcal{P}_{\theta_0}$ à $\mathcal{P}_{\theta_1}$. Il est souvent raisonnable de supposer que le processus est stationnaire avant et après t_c (Davis et Franke, 2008). D'autre part, dans nos travaux, les données sont supposées indépendantes et identiquement distribuées (i.i.d.) avant et après changement. Il est toutefois possible d'adapter les détecteurs à des données dépendantes (Lai, 1998).

L'objectif du problème de détection, défini par l'équation 3.4, est de proposer un test binaire capable de détecter un changement brusque.

$$\left\{ \begin{array}{ll} (\mathcal{H}_0) & x_i \stackrel{\text{iid}}{\sim} \mathcal{P}_{\theta_0} \quad \forall i = 1, \dots, t \\ (\mathcal{H}_1) & \begin{array}{ll} x_i \stackrel{\text{iid}}{\sim} \mathcal{P}_{\theta_0} & \forall i = 1, \dots, t_c - 1 \\ x_i \stackrel{\text{iid}}{\sim} \mathcal{P}_{\theta_1} & \forall i = t_c, \dots, t \end{array} \end{array} \right. , \quad (3.4)$$

En d'autres termes, il faut choisir entre deux hypothèses : $\mathcal{H}_0$ représente l'hypothèse nulle (pas de changement) et $\mathcal{H}_1$ est l'hypothèse alternative (avec changement), composite notamment à cause du paramètre inconnu t_c (variable aléatoire à valeurs dans $\mathbb{N}^*$).

On distingue généralement deux cadres pour la détection séquentielle de changement (Basseville et Nikiforov, 1993) :

Détection hors-ligne : A partir d'une base de mesures collectées $x_1, \dots, x_t$ de taille finie, on cherche à détecter un changement brusque. L'objectif est donc de construire un test dont la sortie sera l'hypothèse $\mathcal{H}_0$ ou $\mathcal{H}_1$ (voir équation 3.4), en s'appuyant sur l'ensemble des mesures :

$$\delta(x_1, \dots, x_t) \in \{\mathcal{H}_0, \mathcal{H}_1\}. \quad (3.5)$$

Détection en-ligne : Les mesures sont collectées séquentiellement, au fur et à mesure du temps, et la taille de l'échantillon collecté n'est pas connue a priori. L'objectif est de construire un test capable de détecter une rupture le plus rapidement possible en s'appuyant sur l'ensemble des mesures disponibles à cet instant. Il s'agit de trouver successivement les applications δ telles que :

$$\begin{aligned} \delta(x_1) &\in \begin{cases} \mathcal{H}_0 : t_c > 1 \\ \mathcal{H}_1 : t_c = 1 \end{cases} \\ &\vdots \\ \delta(x_1, \dots, x_t) &\in \begin{cases} \mathcal{H}_0 : t_c > t \\ \mathcal{H}_1 : 1 \leq t_c \leq t \end{cases} \\ &\vdots \end{aligned} \quad (3.6)$$

Pour nos applications, nous avons choisi une approche de détection en-ligne car nous souhaitons détecter des anomalies le plus tôt possible (prise de décision dès l'acquisition de nouvelles données).

L'algorithme de détection signale un temps d'arrêt t_a , appelé *temps d'alarme*, qui ne correspond pas toujours au temps t_c . Si une rupture est détectée avant un changement brusque ($t_a < t_c$, $\mathcal{H}_1$ validée à tort), on parle de *fausses alarmes*. Sinon, la rupture est signalée avec un *retard à la détection* de valeur $\tau = t_a - t_c$ et illustré par la figure 3.1.

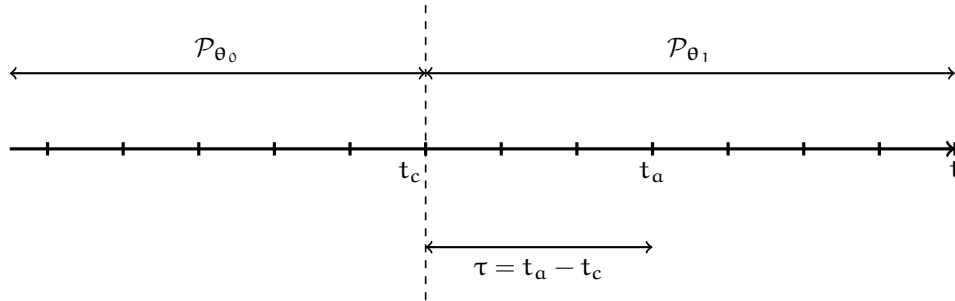


FIGURE 3.1 – Changement brusque à l'instant t_c détecté au temps t_a avec retard.

Définition 3.4 (Probabilités liées à la qualité de la détection d'un test binaire)

- Probabilité de fausse alarme (ou niveau de signification du test) : $\alpha = \alpha_0(\delta)$;
- Niveau du test : $1 - \alpha_0(\delta)$;
- Probabilité de non détection (pour un instant t fixé) : $\alpha_1(\delta)$;
- Probabilité de bonne détection (ou puissance du test) : $\beta = 1 - \alpha_1(\delta)$.

	$\hat{\mathcal{H}}_0$	$\hat{\mathcal{H}}_1$
$\mathcal{H}_0$	$1 - \alpha_0$	α_0
$\mathcal{H}_1$	α_1	$1 - \alpha_1$

3.2 Algorithmes classiques de détection à base de test d'hypothèses

Cette section a pour objet de passer en revue les principales méthodes de détection séquentielle basées sur un test d'hypothèses et leurs garanties d'optimalité. La plupart de ces méthodes ont été examinées par [Basseville et Nikiforov \(1993\)](#), par [Lai \(2001\)](#), et on s'appuie également sur les travaux de [Fillatre \(2011\)](#). Récemment, des états de l'art ont été dressés sur ce type de problèmes et méthodes par [Polunchenko et Tartakovsky \(2012\)](#), et dans un contexte bayésien par [Tartakovsky et Moustakides \(2010\)](#).

La section commence par la présentation de quelques méthodes classiques de détection en ligne dans la sous-section 3.2.1. Les différences entre ces méthodes résident dans la formulation des statistiques de test, dans les hypothèses prises sur la nature des données et la connaissance disponible du processus observé. Puis, on présente les résultats classiques d'optimalité sur ces détecteurs ainsi qu'une stratégie pour l'estimation du seuil de détection, respectivement dans les sous-sections 3.2.2 et 3.2.3. Enfin, on décrit quelques indicateurs en terme de performance de détection dans la sous-section 3.2.4.

3.2.1 Quelques tests pour le problème de détection séquentielle

On se place dans le cadre d'un test binaire, illustré par l'équation 3.4, afin de choisir une décision entre une hypothèse de changement $\mathcal{H}_1$ au temps t_c , et une hypothèse sans changement $\mathcal{H}_0$. Lorsque l'hypothèse $\mathcal{H}_1$ est acceptée, les méthodes de détection signalent usuellement l'instant de rupture t_c à partir de la règle d'arrêt suivante :

$$t_a = \inf \{t \geq 1 : g_t \geq h\}, \quad (3.7)$$

où t_a est le *temps d'alarme* signalé par le détecteur, g_t est la *statistique de test* à l'instant t et h est un seuil fixé en fonction d'une contrainte liée à la fréquence de fausses alarmes.

Les méthodes séquentielles de détection à base de tests s'appuient le plus souvent sur le *rapport de vraisemblance* (Likelihood Ratio, LR), une quantité couramment utilisée en statistique. En particulier, le logarithme de ce rapport de vraisemblance est employé par la plupart des procédures de détection présentées jusqu'à la fin de cette sous-section. Introduisons cette quantité à partir des deux densités paramétriques supposées connues $p(\cdot; \theta_0)$ et $p(\cdot; \theta_1)$, respectivement avant et après changement, sur une observation mesurée à l'instant t :

$$S_t = \log \frac{p(x_t; \theta_1)}{p(x_t; \theta_0)}, \quad (3.8)$$

où θ_0 et θ_1 représentent les paramètres des deux densités différentes ($\theta_0 \neq \theta_1$). Il est à noter que l'espérance de cette quantité change de signe à mesure que les observations suivent la loi avant ou après le temps de changement t_c :

$$\mathbb{E}(S_t | t < t_c; \theta_0) < 0 \quad \text{et} \quad \mathbb{E}(S_t | t \geq t_c; \theta_1) > 0. \quad (3.9)$$

Test de Shewhart

La procédure de [Shewhart \(1931\)](#) est une des premières procédures de détection en ligne à avoir été proposée. Elle est basée sur le partitionnement de l'échantillon séquentiellement observé en fenêtres locales de taille W (fixée par l'utilisateur). Ce type de test ne permet de prendre une décision qu'après la réception de W nouvelles observations. Formellement, la statistique du test de Shewhart s'écrit :

$$g_t = \sum_{i=t-W+1}^t S_i, \quad (3.10)$$

où S_i est défini dans l'équation 3.8.

L'hypothèse $\mathcal{H}_1$ est acceptée lorsque la statistique g_t dépasse un seuil h fixé. La procédure de Shewhart traite séquentiellement le flux d'observations en fenêtres de taille W et le temps d'alarme t_a est donc signalé par la règle suivante :

$$t_a = \inf \{t = W, 2W, \dots : g_t \geq h\}. \quad (3.11)$$

Cette procédure est également présentée dans ([Basseville et Nikiforov, 1993](#), section 2.1), et appelée test à Taille d'Echantillon Fixe (TEF) par [Fillatre \(2011\)](#). Il est précisé que l'algorithme TEF est deux fois moins rapide que l'algorithme du CUSUM, présenté par la suite. En pratique, cette méthode peut être assez sensible aux variations du processus observé et admettre une probabilité de fausses alarmes relativement importante.

A titre d'illustration, étudions le débit moyen annuel de la rivière du Nil de 1872 à 1970, décrit par la figure 3.3(a). La série change de moyenne après 1898, date de la construction du premier barrage à Assouan. La figure 3.2 présente deux statistiques de Shewhart avec des fenêtres de taille 2 et 5. Augmenter cette taille W diminue le nombre de fausses alarmes mais espace les prises de décision. Pour $W = 2$, on observe une fausse alarme (1889) et quatre non détections. Pour $W = 5$, il n'y a plus de fausse alarme et de non détection mais le temps d'alarme est $t_a = 1901$, soit un retard à la détection de $\tau = 3$.

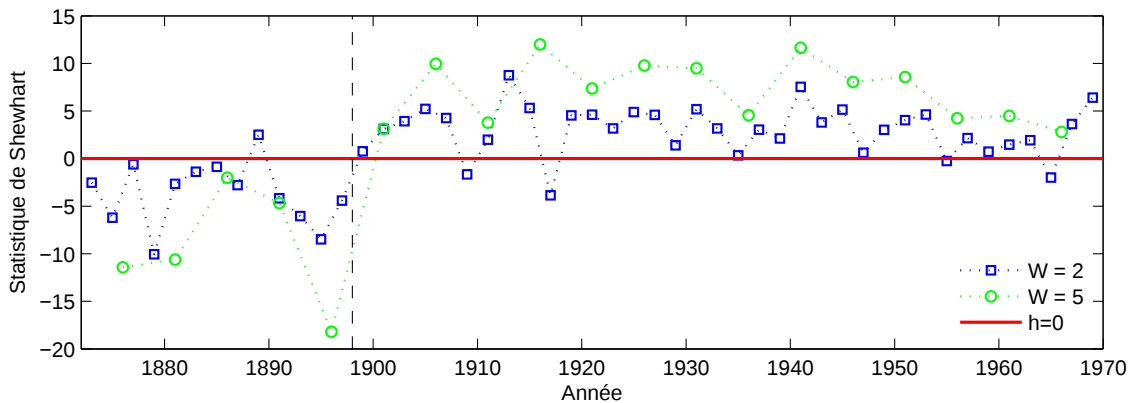


FIGURE 3.2 – Procédure de Shewhart appliquée au suivi des débits moyens annuels de la rivière du Nil de 1872 à 1970 (voir figure 3.3(a)).

Test du CUSUM

Le test du CUSUM, initié par [Page \(1954\)](#), est la première méthode de détection séquentielle avec des propriétés d'optimalité. Dans le cas de données indépendantes, [Lorden \(1971\)](#) a montré que cette méthode est asymptotiquement optimale (voir critère 3.1). L'optimalité non asymptotique a été établie par [Moustakides \(1986\)](#). Cette procédure est à l'origine de nombreuses stratégies de détection et reste fréquemment employée de nos jours. Dans le chapitre 5, nous introduisons une stratégie séquentielle de détection de type CUSUM. La règle du CUSUM classique s'appuie sur la somme cumulée S_{pt} suivante :

$$\begin{aligned} S_{pt} &= \log \frac{p(x_1, \dots, x_t | t_c = p; \theta_0, \theta_1)}{p(x_1, \dots, x_t; \theta_0)} \\ &= \log \frac{p(x_1, \dots, x_{p-1}; \theta_0) \cdot p(x_p, \dots, x_t; \theta_1)}{p(x_1, \dots, x_{p-1}; \theta_0) \cdot p(x_p, \dots, x_t; \theta_0)} \\ &= \sum_{i=p}^t \log \frac{p(x_i; \theta_1)}{p(x_i; \theta_0)}, \end{aligned} \quad (3.12)$$

où le temps de changement t_c est considéré comme une variable aléatoire, les paramètres θ_0 et θ_1 sont supposés connus, et les observations $x_1, \dots, x_t$ sont indépendantes.

La procédure du CUSUM utilise la règle d'arrêt définie par l'équation 3.7 avec la statistique de test suivante :

$$g_t = \max_{1 \leq p \leq t} S_{pt}. \quad (3.13)$$

Il est à noter que la procédure du CUSUM peut être interprétée comme une utilisation répétée du test séquentiel du rapport de vraisemblance (SPRT), introduit par Wald, pionnier dans l'analyse des problèmes de détection séquentielle ([Basseville et Nikiforov, 1993](#), section 2.2). D'autre part, la statistique du CUSUM admet une formulation récursive, ce qui peut se révéler particulièrement utile dans le cadre de la détection en-ligne.

$$g_t = (g_{t-1} + S_t)^+, \quad (3.14)$$

où $x^+ = \max\{0, x\}$ et S_t est défini dans l'équation 3.8. Les deux formulations (équations 3.13 et 3.14) sont équivalentes dès lors que la statistique est positive : $g_t \geq 0, \forall t$.

Reprenons l'exemple de la figure 3.3(a) représentant le débit moyen annuel de la rivière du Nil de 1872 à 1970. On suppose que ce débit annuel suit une loi normale avant 1898 et une autre après 1898. Les isocontours et paramètres de ces deux densités sont représentés par la figure 3.3(b). Nous avons calculé les statistiques de test de manière récursive (équation 3.14) et, non récursive (équation 3.13), illustrées par la figure 3.3(c). Les deux formulations ont donné lieu à la détection d'un même changement en 1900. Enfin, il est à noter que la version récursive accélère le temps de calcul d'un facteur 50 sur cet exemple, selon nos expérimentations.

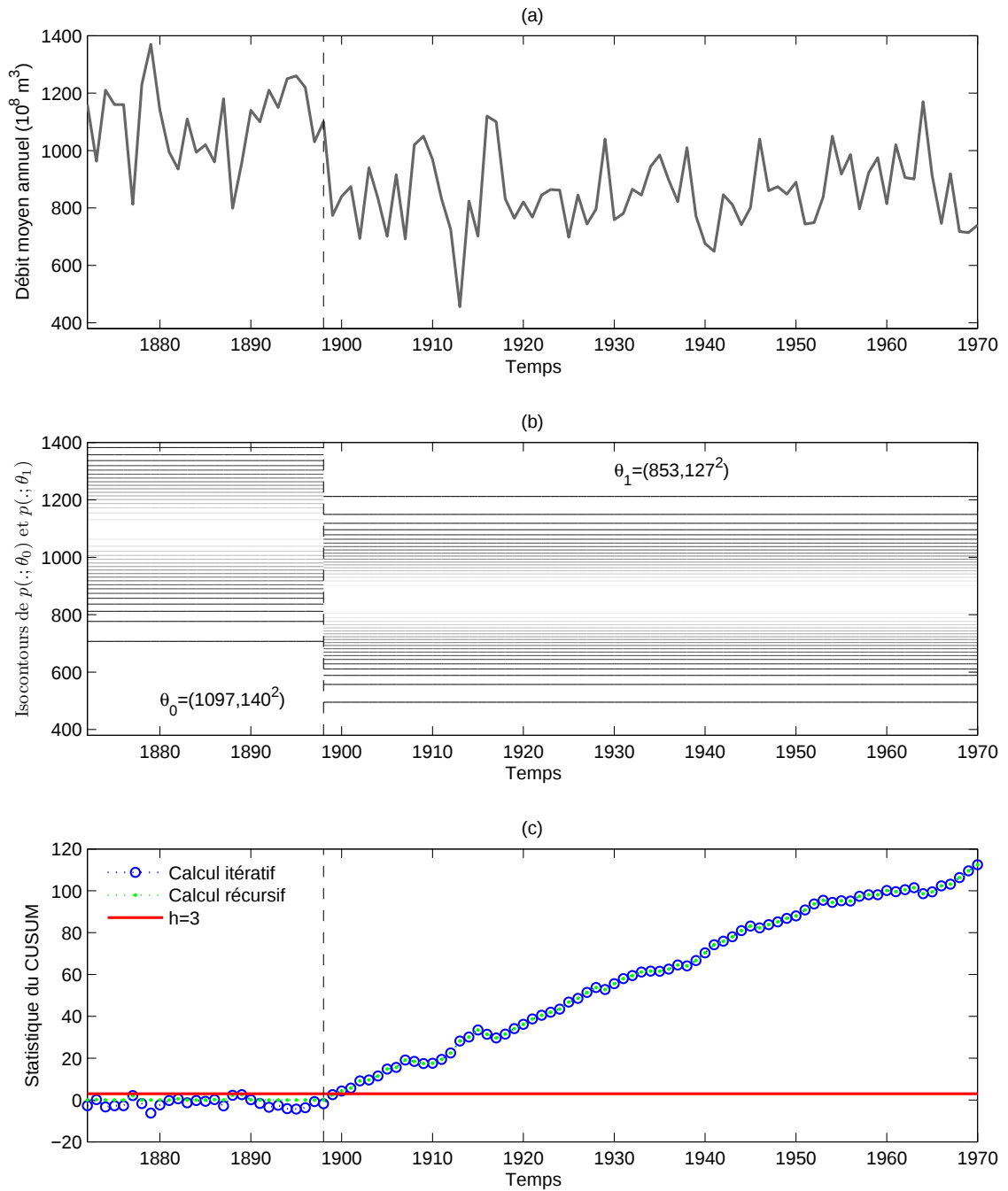


FIGURE 3.3 – Procédure du CUSUM appliquée au suivi des débits moyens annuels de la rivière du Nil de 1872 à 1970 avec un changement en 1898 (a). Isocontours des densités normales avant et après changement (b). Statistiques du CUSUM avec un calcul itératif et récursif. Le temps d'alarme est signalé en 1900 (c).

Il existe plusieurs autres formulations de tests séquentiels basées sur le rapport de vraisemblance, dont l'optimalité a parfois pu être établie. Nous présentons ici quelques-uns de ces tests.

Test du GLR

Les hypothèses du CUSUM ne correspondent pas toujours aux situations rencontrées dans des applications réelles. En particulier, la densité $p(\cdot; \theta_1)$ est rarement connue car il est difficile de connaître à l'avance la forme exacte de la distribution après changement. Le test du *rapport de vraisemblance généralisé* (*Generalized Likelihood Ratio, GLR*) proposé par [Lorden \(1971\)](#) peut être envisagé lorsque le paramètre θ_1 est inconnu mais appartient à un certain espace Θ . Ce paramètre peut alors être estimé par le principe du maximum de vraisemblance (MV). Le temps d'alarme est toujours signalé selon l'équation 3.7 et la statistique du GLR est définie par :

$$\begin{aligned} g_t &= \max_{1 < p \leq t} \sup_{\theta_1 \in \Theta} \log \frac{p(x_1, \dots, x_t | t_c = p; \theta_0, \theta_1)}{p(x_1, \dots, x_t; \theta_0)} \\ &= \max_{1 < p \leq t} \sum_{i=p}^t \log \frac{p(x_i; \hat{\theta}_1)}{p(x_i; \theta_0)}, \end{aligned} \quad (3.15)$$

où les observations $x_1, \dots, x_t$ sont indépendantes, le paramètre θ_0 est supposé connu, et le paramètre $\hat{\theta}_1$ est estimé par maximisation de la vraisemblance $p(x_p, \dots, x_t | \theta)$. Une procédure de type GLR est notamment proposée par [Siegmund et Venkatraman \(1995\)](#) au sens du critère 3.1 de Lorden pour un changement de moyenne dans une séquence de variables normales. Enfin, il est à noter que la statistique du GLR n'a pas de formulation récursive contrairement à la celle du CUSUM (équation 3.14).

Test du GLRFM

Aussi bien le test du GLR que celui du CUSUM cherchent un point de changement entre les temps 1 et t (équations 3.13 et 3.15). Ces formulations impliquent donc de calculer chaque statistique sur une infinité d'observations à mesure que l'instant t croît. Concernant une procédure de type GLR, l'estimation du paramètre de la densité $p(\cdot | \theta_1)$ est souvent coûteuse en temps de calcul et peut s'avérer, en pratique, difficilement implémentable. Une première alternative a été proposée par [Willsky et Jones \(1976\)](#) afin de réduire la complexité computationnelle de l'étape 3.15 en maintenant un temps de calcul constant à chaque temps t . L'idée est de limiter la recherche du point de changement à une fenêtre de taille W fixée, sans démarrer la recherche depuis la première observation. La statistique du *GLR à fenêtre limitée* (*GLRFM*) est définie par :

$$g_t = \max_{t-W < p \leq t} \sup_{\theta_1 \in \Theta} \sum_{i=p}^t \log \frac{p(x_i | \theta_1)}{p(x_i | \theta_0)}. \quad (3.16)$$

Quelques résultats d'optimalité de cette règle ont été obtenus par [Lai \(1998\)](#) sous certaines hypothèses, pour des données indépendantes ou dépendantes.

Par ailleurs, d'autres solutions ont été proposées afin de réduire la complexité computationnelle d'une procédure de type GLR. L'*approche asymptotique locale*, décrite par [Basseville et Nikiforov \(1993\)](#), pages 126-129), permet de remplacer la tâche de maximisation de la vraisemblance par des approches locales ; l'idée est ici de surveiller des résidus à partir du gradient de la log-vraisemblance (fonction score). D'autre part, [Nikiforov \(2000, 2003\)](#) propose une statistique avec une formulation récursive et l'optimalité de la procédure est établie à partir d'un critère lié au retard moyen à la détection (équation 3.27).

Test de SR

Supposons à nouveau que les données i.i.d. $(x_t)_{t \geq 1}$ observées séquentiellement suivent une densité connue : $p(\cdot; \theta_0)$ quand $t < t_c$ et suivent une autre densité $p(\cdot; \theta_1)$ quand $t \geq t_c$. [Shiryaev \(1963\)](#) propose une procédure de détection dans un cadre bayésien. Ainsi, l'instant de changement t_c est une variable aléatoire qui suit une loi géométrique *a priori* : $P(t_c = t) = q(1 - q)^{n-1}$. Quand $q \rightarrow 0$, Shiryaev et [Roberts \(1966\)](#) proposent la même statistique de test :

$$g_t = \lim_{q \rightarrow 0} \sum_{p=1}^t \prod_{i=p}^t \frac{p(x_i; \theta_1)}{(1 - q) \cdot p(x_i; \theta_0)} = \sum_{p=1}^t \prod_{i=p}^t \frac{p(x_i; \theta_1)}{p(x_i; \theta_0)}. \quad (3.17)$$

Cette mesure est appelée la *statistique du test de Shiryaev-Roberts (SR)* et permet de signaler un changement avec la règle d'arrêt donnée en équation 3.7. Comme pour la règle du CUSUM, cette statistique admet une formulation récursive :

$$g_t = (g_{t-1} + 1) \frac{p(x_t | \theta_1)}{p(x_t | \theta_0)} \quad \text{avec } t \geq 1 \text{ et } g_0 = 0. \quad (3.18)$$

Le critère proposé récemment par [Tartakovsky et Veeravalli \(2005\)](#) établit l'optimalité asymptotique d'une procédure de type SR, dans une approche bayésienne sous certaines hypothèses sur la distribution a priori (indépendante des observations). Il est à noter que ce même critère est utilisé pour montrer l'optimalité de tests de type CUSUM et de type SR dans ([Tartakovsky et Moustakides, 2010](#)).

Comme pour le test du CUSUM, il est possible de modifier la procédure de SR lorsque le paramètre θ_1 est inconnu mais appartient à un espace Θ . [Pollak \(1987\)](#) propose ainsi d'étendre la règle de Shiryaev-Roberts (équation 3.17) et d'utiliser la statistique :

$$g_t = \sum_{p=1}^t \int_{\Theta} \left(\prod_{i=p}^t \frac{p(x_i | \theta_1)}{p(x_i | \theta_0)} \right) \pi(\theta_1) d\theta_1, \quad (3.19)$$

où $\pi(\theta_1)$ est une distribution de probabilité sur un intervalle ouvert de Θ .

3.2.2 Résultats classiques d'optimalité d'un détecteur

La performance d'un détecteur est généralement évaluée par deux types d'indicateurs (définition 3.4) : une mesure du délai de détection après un changement et une mesure liée à la fréquence des fausses alarmes. Le premier résultat d'optimalité a été montré par Shiryaev (1961) pour une approche de détection en-ligne, en s'appuyant sur une formulation bayésienne du problème considérant que l'instant de rupture t_c suit une loi a priori. Le second critère d'optimalité, à la base des méthodes de détection que nous allons utiliser, a été proposé par Lorden (1971) pour un problème de détection en-ligne. Il repose sur un test de type *minimax*, défini en 3.5. Dans cette partie, on s'appuie notamment sur les formalisations de (Fillatre, 2011, sous-section 1.4.2).

Définition 3.5 (Test de type minimax)

Un test $\bar{\delta}$ est dit *minimax* s'il est de la forme :

$$\bar{\delta} = \arg \min_{\delta} \left\{ \max_k \alpha_k(\delta) \right\}. \quad (3.20)$$

Le temps t_c étant la réalisation d'une variable aléatoire, on définit la distribution $\mathcal{P}_{(t_c)}$ sur la séquence d'observations $(x_1, x_2, \dots)$ telle qu'un changement brusque existe à l'instant t_c , c'est-à-dire que $(x_1, \dots, x_{t_c-1}) \sim \mathcal{P}_{\theta_0}$ et $(x_{t_c}, x_{t_c+1}, \dots) \sim \mathcal{P}_{\theta_1}$. On suppose également que la loi $\mathcal{P}_{(t_c)}$ admet une espérance théorique $\mathbb{E}_{(t_c)}$ issue de la probabilité $\mathcal{P}_{(t_c)}$. Remarquons enfin que $\mathcal{P}_{(\infty)} = \mathcal{P}_{\theta_0}$ et $\mathbb{E}_{(\infty)} = \mathbb{E}(\cdot; \theta_0)$.

Définition 3.6 (Fonction ARL)

La fonction d'ARL (Average Run Length) calcule la durée moyenne entre deux fausses alarmes (ou la durée moyenne avant une fausse alarme) :

$$\text{ARL} : \begin{array}{l} \Theta \rightarrow \mathbb{R} \\ \theta \mapsto \mathbb{E}(t_a; \theta) \end{array}. \quad (3.21)$$

Cette fonction est capitale pour la construction de nombreuses méthodes séquentielles de détection, y compris pour l'algorithme proposé dans le chapitre 5.

Critère 3.1 (Critère d'optimalité de Lorden (1971))

On définit $\bar{\tau}_1$, le retard moyen à la détection comme l'espérance du retard à la détection conditionnellement à l'instant de rupture et aux observations $x_1, \dots, x_{t_c-1}$:

$$\bar{\tau}_1 = \mathbb{E}_{(t_c)}(t_a - t_c + 1 \mid t_a \geq t_c, x_1, \dots, x_{t_c-1}). \quad (3.22)$$

Le critère de Lorden est obtenu pour des données indépendantes et identiquement distribuées avant et après l'instant de rupture. Il consiste à minimiser le pire retard

moyen à la détection $\bar{\tau}_1^*$ sous la contrainte d'avoir fixé le temps moyen entre deux fausses alarmes :

$$\bar{\tau}_1^* = \sup_{t_c \geq 1} \text{ess sup } \mathbb{E}_{(t_c)}(t_a - t_c + 1 \mid t_a \geq t_c, x_1, \dots, x_{t_c-1}) \quad (3.23)$$

$$\text{s.c.} \quad \mathbb{E}(t_a; \theta_0) \geq \gamma > 0, \quad (3.24)$$

où ess sup est l'opérateur du supremum essentiel défini dans la section A.2, et la contrainte 3.24 impose que la fonction ARL sur θ_0 , définie en 3.6, soit supérieure ou égale à une constante γ .

Afin d'établir l'optimalité de ce critère dans un cadre asymptotique, Lorden s'appuie sur la borne inférieure du pire retard moyen :

$$\inf_{t_a} \{ \bar{\tau}_1^* / \mathbb{E}(t_a; \theta_0) \geq \gamma \} \approx \frac{\log \gamma}{\text{KL}(\theta_1, \theta_0)}, \quad \text{quand } \gamma \rightarrow \infty, \quad (3.25)$$

où $\text{KL}(\theta_1, \theta_0) = \int_{\mathbb{R}} p(x; \theta_1) \cdot \log \frac{p(x; \theta_1)}{p(x; \theta_0)} dx > 0$ est l'entropie relative ou la divergence de Kullback-Leibler entre les distributions $\mathcal{P}_{\theta_0}$ et $\mathcal{P}_{\theta_1}$.

L'optimalité du test du CUSUM (équation 3.13) a été prouvé à partir du critère 3.1 de Lorden (1971) dans le sens où il minimise le retard moyen à la détection défini par l'équation 3.22. Le seuil du test $h \approx \log \gamma$ minimise ainsi le pire retard moyen à la détection, quand $\gamma \rightarrow \infty$. Ce résultat s'appuie sur le changement de signe de l'espérance du rapport de log-vraisemblances après changement (sous $\mathcal{H}_1$) décrit dans l'équation 3.9. Moustakides (1986) et Ritov (1990) ont établi l'optimalité du critère 3.1 de Lorden dans un cadre non asymptotique (avec $\gamma \in \mathbb{N}^*$ fixé). Les critères suivants décrivent des extensions/généralisations du critère de Lorden.

Critère 3.2 (Critère d'optimalité de Pollak et Siegmund (1975))

On définit $\bar{\tau}_2$, le retard moyen à la détection comme l'espérance du retard à la détection conditionnellement à l'instant de rupture :

$$\bar{\tau}_2 = \mathbb{E}_{(t_c)}(t_a - t_c + 1 \mid t_a \geq t_c). \quad (3.26)$$

Le critère de Pollak et Siegmund, version modifiée du critère de Lorden, consiste à minimiser le retard moyen à la détection $\bar{\tau}_2$, sous la contrainte d'ARL (équation 3.24) :

$$\bar{\tau}_2^* = \sup_{t_c \geq 1} \mathbb{E}_{(t_c)}(t_a - t_c + 1 \mid t_a \geq t_c). \quad (3.27)$$

Critère 3.3 (Critère d'optimalité de Lai (1998))

Le critère d'optimalité, introduit par Lai, diffère des deux critères précédant dans le sens où il consiste à minimiser l'espérance du retard positif $\bar{\tau}_3$ (on note

$(x)^+ = \max(0, x)$), pour une séquence de données dépendantes ou indépendantes :

$$\bar{\tau}_3 = \mathbb{E}_{(t_c)}(t_a - t_c)^+ \quad (3.28)$$

$$\text{s.c.} \quad \sup_{j \geq 1} P(j \leq t_a \leq j + m_\alpha; \theta_0) \leq \alpha \quad (3.29)$$

$$\text{avec} \quad \begin{cases} \inf \frac{m_\alpha}{|\log \alpha|} > I^{-1} \\ \log m_\alpha = o(\log \alpha) \end{cases} \quad \text{quand } \alpha \rightarrow 0, \quad (3.30)$$

où α est la probabilité de fausses alarmes maximale pendant la durée m_α ($m_\alpha \in \mathbb{N}^*$) et la constante I désigne la divergence de Kullback-Leibler entre $\mathcal{P}_{\theta_0}$ et $\mathcal{P}_{\theta_1}$ dans le cas de données indépendantes. Dans le cas de données dépendantes, la constante I est définie comme la limite presque sûre (sous $\mathcal{P}_{(t_c)}$ et uniformément pour $t_c \geq 1$) :

$$\frac{1}{t} \sum_{i=t_c}^{t_c+t} \log \frac{p(\mathbf{x}_i | \mathbf{x}_0, \dots, \mathbf{x}_{i-1}; \theta_1)}{p(\mathbf{x}_i | \mathbf{x}_0, \dots, \mathbf{x}_{i-1}; \theta_0)} \xrightarrow[t \rightarrow \infty]{\text{p.s.}} I.$$

Si les deux conditions 3.29 et 3.30 sont vérifiées et également la condition suivante

$$\lim_{t \rightarrow \infty} \sup_{t_c \geq 1} P_{(t_c)} \left(\max_{j \leq t} \sum_{i=t_c}^{t_c+j} \log \frac{p(\mathbf{x}_i | \theta_1)}{p(\mathbf{x}_i | \theta_0)} \geq I \cdot (1 + \zeta) \cdot t \right) = 0, \quad \forall \zeta > 0, \quad (3.31)$$

alors le retard $\bar{\tau}_3$ (e. 3.28) vérifie l'inégalité suivante (Lai, 1998, Théorème 2) quand $\alpha \rightarrow 0$:

$$\mathbb{E}_{(t_c)}(t_a - t_c)^+ \geq \left(\frac{P(t_a - t_c | \theta_0)}{I} + o(1) \right) \cdot |\log \alpha|, \quad (3.32)$$

uniformément pour $t_c \geq 1$.

Outre les critères qui viennent d'être énoncés, on peut également citer le critère d'optimalité proposé par Moustakides (2008) pour un test de type CUSUM, qui unifie plusieurs critères précédents et constitue une version étendue du critère de Lorden dans un cadre assez général. Le critère proposé récemment par Tartakovsky et Veeravalli (2005) est utilisé pour établir l'optimalité d'une procédure de type SR dans un cadre bayésien. Il est à noter que ce même critère est utilisé pour montrer l'optimalité de tests de type CUSUM et de type SR dans (Tartakovsky et Moustakides, 2010), présentés dans la sous-section 3.2.1.

On ne cherche pas ici à établir une liste exhaustive de l'ensemble des critères proposés pour des problèmes de détection à base de tests d'hypothèses. Les critères d'optimalité diffèrent par leur définition du retard à la détection et de leur manière de contrôler les erreurs de détection de changement (Fillatre, 2011, chapitre 1). Il est à noter que de nombreux détecteurs s'appuient sur la contrainte d'ARL définie dans l'équation 3.24. La prochaine sous-section met en évidence une relation directe entre la fonction d'ARL et le taux de fausses alarmes.

3.2.3 Contrainte de fausses alarmes et estimation du seuil de détection

En détection séquentielle, on cherche à proposer des méthodes qui détectent un changement le plus rapidement possible. Toutefois si le détecteur est spécifié avec une trop forte sensibilité, la procédure donne lieu à de fréquentes fausses alarmes, sous l'hypothèse $\mathcal{H}_0$. À l'inverse, réduire au minimum le taux de fausses alarmes conduit la procédure à augmenter le délai avant détection voire à ne détecter aucun événement. L'objectif est donc de développer des tests qui réalisent un compromis entre ces deux extrêmes et ainsi, minimisent le *délai moyen de détection* sous une contrainte liée au *taux de fausses alarmes* (TFA) (Tartakovsky et Moustakides, 2010).

Historiquement, la mesure standard liée à la fréquence des fausses alarmes est l'ARL, introduite dans la définition 3.1. Cette mesure correspond à la durée moyenne entre deux fausses alarmes. Dans cette sous-section, on met en évidence un lien entre la fonction d'ARL et le taux théorique de fausses alarmes pour un détecteur. La relation présentée est reprise de Verdier (2007, chapitre 3) dont les travaux portent sur une procédure de type CUSUM (voir sous-section 3.2.1) avec seuil adaptatif appliqué à des données dépendantes. Ce lien est important car il nous permet d'étendre les propriétés d'optimalité obtenues avec une contrainte sur l'ARL, à une contrainte sur le taux de fausses alarmes. En particulier, la procédure de détection présentée dans le chapitre 5 s'appuie sur une contrainte du TFA. La contrainte d'ARL est vérifiée par l'inverse du taux théorique de fausses alarmes comme le montre la proposition suivante :

Proposition 3.1 (Lien entre ARL et TFA)

Dans le cas de données indépendantes et à partir d'un seuil h fixé, le temps moyen entre deux fausses alarmes γ est égal à l'inverse du taux de fausses alarmes α :

$$\gamma = 1/\alpha. \quad (3.33)$$

Preuve. Notons $\alpha = \alpha_0(\delta)$, la probabilité de fausse alarme d'un test δ , comme nous l'avons déjà introduit dans la définition 3.3. À partir d'un seuil h fixé, on écrit la quantité $\alpha \in]0, 1[$ comme la probabilité conditionnelle sous l'hypothèse $\mathcal{H}_0$:

$$\begin{aligned} P(g_1 \geq h | \theta_0) &= \alpha \\ P(g_t \geq h | g_{t-1} < h, \dots, g_1 < h; \theta_0) &= \alpha, \forall t \geq 2, \end{aligned} \quad (3.34)$$

où g_t est la statistique du test δ au temps t et θ_0 est le paramètre de la distribution avant changement. Le seuil h est à choisir de telle sorte que la contrainte d'ARL (équation 3.24) soit vérifiée : $\mathbb{E}(t_a; \theta_0) = \gamma$. Comme les observations sont supposées indépendantes, γ correspond au temps moyen avant la première fausse alarme mais également au temps moyen entre deux fausses alarmes. Cette durée γ est choisie par l'utilisateur afin de régler la sensibilité du détecteur.

En d'autres termes, le taux de fausses alarmes α est considéré comme constant pendant toute la procédure de surveillance. Si la statistique de test g_t dépasse le seuil h , un temps d'alarme t_a est signalé par le détecteur au temps t . A partir de l'équation 3.34 et en considérant t_a comme une variable aléatoire, on peut écrire la probabilité de t_a sous l'hypothèse $\mathcal{H}_0$:

$$\begin{aligned}
 P(t_a = t; \theta_0) &= P(g_t \geq h, g_{t-1} < h, \dots, g_1 < h; \theta_0) \\
 &= P(g_t \geq h | g_{t-1} < h, \dots, g_1 < h; \theta_0) \times P(g_{t-1} < h, \dots, g_1 < h; \theta_0) \\
 &= P(g_t \geq h | g_{t-1} < h, \dots, g_1 < h; \theta_0) \times P(g_1 < h; \theta_0) \\
 &\quad \times \prod_{i=2}^{t-1} P(g_i < h | g_{i-1} < h, \dots, g_1 < h; \theta_0) \\
 &= \alpha \times (1 - \alpha)^{t-1}.
 \end{aligned} \tag{3.35}$$

On reconnaît une loi géométrique de paramètre α dans l'équation 3.35. D'après la définition donnée en équation 3.6, la fonction d'ARL est l'espérance du temps t_a avant changement. On calcule alors l'ARL en fonction de la probabilité de fausse alarme :

$$\begin{aligned}
 \mathbb{E}(t_a; \theta_0) &= \lim_{t \rightarrow \infty} \sum_{i=0}^t i \times P(t_a = i; \theta_0) \\
 &= \lim_{t \rightarrow \infty} \sum_{i=0}^t i \times \alpha \times (1 - \alpha)^{i-1} \\
 &= \lim_{t \rightarrow \infty} \alpha \times \sum_{i=0}^t i \times (1 - \alpha)^{i-1} \\
 &= \alpha \times \lim_{t \rightarrow \infty} \sum_{i=1}^t i \times (1 - \alpha)^{i-1} \\
 &= \alpha \times \left(-\frac{d}{d\alpha} \lim_{t \rightarrow \infty} \sum_{i=1}^t (1 - \alpha)^i \right) \\
 &= \alpha \times \left(-\frac{d}{d\alpha} \frac{1 - \alpha}{1 - (1 - \alpha)} \right) \\
 &= \alpha \times \left(-\frac{d}{d\alpha} \left(\frac{1}{\alpha} - 1 \right) \right) \\
 &= \alpha \times \left(\frac{1}{\alpha^2} \right) \\
 &= \frac{1}{\alpha}.
 \end{aligned} \tag{3.37}$$

Le temps moyen entre deux fausses alarmes γ est donc égal à l'inverse du taux de fausses alarmes α :

$$\mathbb{E}(t_a; \theta_0) = \frac{1}{\alpha} = \gamma.$$

□

On retrouve cette relation dans (Mei, 2003, chapitre 2). En particulier, ce lien théorique direct entre l'ARL sur θ_0 et la probabilité de fausse alarme nous permet de fixer un seuil de détection à partir d'un taux α fourni par l'utilisateur de la manière suivante :

1. Fixer une probabilité théorique de fausses alarmes α ;
2. Estimer un seuil de détection h_α en fonction du taux de fausses alarmes.

Une méthode d'estimation du seuil h_α est proposée dans la sous-section 5.3.4, basée sur du rééchantillonnage et l'estimation d'un quantile. Cette approche suppose de disposer d'un modèle capable de simuler (rééchantillonner) des observations sous l'hypothèse $\mathcal{H}_0$. On utilise alors le détecteur pour calculer séquentiellement les statistiques de tests à partir des données simulées. Le seuil de détection est alors estimé comme le quantile d'ordre $(1 - \alpha)$ d'une distribution paramétrique a priori sur les statistiques de test.

Cette stratégie proposée pour l'estimation d'un seuil de détection peut être abordée selon deux approches statistiques concurrentes qui sont résumées ci-après.

La première approche repose sur une série de simulations (avant changement) afin d'approcher une valeur de seuil vérifiant une contrainte d'ARL. En effet, ce type d'approche postule qu'il est statistiquement possible d'approcher une telle valeur de seuil à partir d'un grand nombre d'observations générées sous $\mathcal{H}_0$ (Lai, 1995). Lorsque la fonction d'ARL n'est pas spécifiée, il est parfois possible de la calculer ou de l'approcher en utilisant les approximations de Wald ou de Siegmund (Basseville et Nikiforov, 1993, sous-section 5.2.2). Estimer la fonction ARL peut s'avérer complexe, en raison de la difficulté à la calculer précisément (équation 3.36). Il est fréquent d'utiliser des approximations à partir de techniques de type Monte-Carlo (Bishop, 2006, Chapitre 11). Le rééchantillonnage est coûteux en temps de calcul et donc, peut s'avérer inadapté pour des problèmes de détection séquentielle. En particulier lorsque les distributions avant et après changement sont inconnues, car dans ce cas, il est également nécessaire de recalculer ces distributions à chaque estimation de seuil.

La seconde approche suppose qu'il est possible de modéliser la distribution des statistiques de test par une loi paramétrique *a priori* afin de calculer un seuil de détection (Gustafsson et Palmqvist, 2000). Ce type d'approche nécessite uniquement d'estimer les paramètres de la loi et permet d'atteindre, sans simulation, des valeurs de seuils avec une ARL très grande. Le seuil est alors directement calculé comme un quantile de la loi paramétrique. Cette approche suppose une hypothèse assez forte sur la forme de la distribution des statistiques de test mais en supposant que l'échantillon soit de taille suffisante, le seuil peut être estimé directement pour un taux de fausses alarmes très faible.

3.2.4 Évaluation d'un détecteur

Dans cette sous-section, on présente quelques indicateurs de performances pour mesurer la qualité de modèles dans un problème de détection. Le principal indicateur de performance des modèles proposés, dans nos travaux, est la *courbe ROC*, ou *caractéristique opérationnelle de réception* (Krzanowski et Hand, 2009) qui représente la *sensibilité* en fonction de l'*antispécificité* après plusieurs tests :

Antispécificité : probabilité de fausses alarmes PFA (ou erreur de première espèce) :

$$\text{RFP} = \frac{\text{FP}}{\text{FP} + \text{VN}} = 1 - \frac{\text{VN}}{\text{VN} + \text{FP}} ; \quad (3.38)$$

Sensibilité : probabilité de bonne détection TVP (ratio de vrais positifs) :

$$\text{RVP} = \frac{\text{VP}}{\text{VP} + \text{FN}} ; \quad (3.39)$$

où FP est le nombre de faux positifs, VN est le nombre de vrais négatifs, VP est le nombre de vrais positifs et FN est le nombre de faux négatifs. En théorie de la détection, la probabilité de fausses alarmes est appelée niveau de signification du test ou faux positifs, et celle de bonne détection est appelée puissance du test ou vrais positifs (définition 3.4).

La figure 3.4 présente les courbes ROC de trois types de classifieurs (ou détecteurs). Le classifieur parfait atteindrait le point (0, 1) et la diagonale (courbe noire) correspond à un classifieur aléatoire. Une mesure de comparaison, très largement utilisée, est le critère *AUC* calculé comme l'aire sous la courbe ROC (Bradley, 1997). Le meilleur profil de modèle maximise l'AUC. Il est à noter que Hand (2010) a récemment montré qu'une imperfection réside dans la définition d'AUC et propose un nouveau critère alternatif.

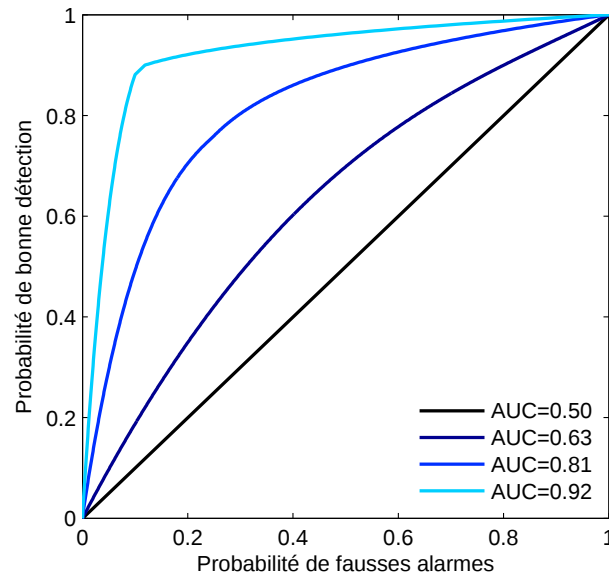


FIGURE 3.4 – Exemple de courbes ROC pour trois types de classifieurs. Le meilleur type de classifieur maximise l'aire sous la courbe (AUC).

Les autres indicateurs de performance des détecteurs séquentiels sont généralement liées aux mesures de retard à la détection et à la fréquence de fausses alarmes (Tartakovsky et Moustakides, 2010). Dans nos travaux, on utilise deux critères pour mesurer la performance des détecteurs :

- Probabilité de fausses alarmes (équation 3.3), ou *False Alarm Rate (FAR)* ;

$$\text{FAR}(t_a) = 1/\mathbb{E}(t_a | t_a < t_c). \quad (3.40)$$

- Temps moyen entre deux fausses alarmes (équation 3.36), ou *Average Run Length to false alarm (ARL)* :

$$\text{ARL}(t_a) = 1/\text{FAR}(t_a). \quad (3.41)$$

- Délai moyen de détection (équation 3.22), ou *Average Detection Delay (ADD)* :

$$\text{ADD}(t_a) = \mathbb{E}(t_a - t_c | t_a \geq t_c). \quad (3.42)$$

En pratique, ces quantités sont calculées de manière empirique et séquentiellement. Par exemple, l'ARL correspond à la moyenne des durées avant que le détecteur signale une alarme avant l'occurrence du vrai temps de changement t_c , le détecteur étant relancé après chaque fausse détection.

Dans le cadre de la théorie de détection de point de changement, la meilleure méthode de détection garantit le plus faible délai moyen de détection pour un taux de fausses alarmes fixé. Le choix du seuil de détection peut être effectué à partir de la courbe ROC ou de la courbe représentant l'ADD en fonction du FAR. Cette étape de calibrage consiste alors à choisir une méthode avec de faibles valeurs de FAR et d'ADD. Tartakovsky et al. (2005) utilisent cette représentation graphique pour montrer la performance de détecteurs appliqués à des problèmes d'attaque sur des réseaux d'ordinateurs distribués. Cet outil graphique sera également utilisé dans la section 5.4 pour évaluer la stratégie de détection séquentielle mise en œuvre pour la surveillance de portes dans un autobus, et détaillée dans le chapitre 5.

La prochaine section présente quelques algorithmes de détection séquentielle à base de test d'hypothèses sur des vecteurs multivariés. Ces méthodes, appelées *cartes de contrôle*, s'appuient sur les tests binaires séquentiels présentés dans la sous-section 3.2.1 et supposent certaines hypothèses particulières sur les distributions. La carte de Shewhart et celle de la moyenne mobile sont notamment présentées dans (Basseville et Nikiforov, 1993, section 2.1).

3.3 Cartes de contrôle et autres tests séquentiels

Cette section a pour objet de présenter quelques extensions des cartes de contrôle de Shewhart et de CUSUM, respectivement introduites par [Shewhart \(1931\)](#) et [Page \(1954\)](#). Pour cela, on s'appuie principalement sur les reviews de cartes multivariées de [MacGregor et Kourti \(1995\)](#), de [Lowry et Montgomery \(1995\)](#) et notamment celle de [Bersimis et al. \(2007\)](#). Certaines de ces méthodes sont employées pour le suivi de fonctionnement du système de freinage d'un autobus (voir le chapitre 4). Apparues dès le premier quart du 20^e siècle, les cartes de contrôle sont toujours largement employées pour la surveillance de processus industriels. Leurs relatives simplicités d'utilisation et d'implémentation expliquent leur popularité pour le contrôle qualité de nombreux produits industriels. La communauté de la *maîtrise statistique des procédés* (Statistical Process Control, SPC) a souvent fait l'objet de critiques et controverses vis-à-vis d'autres branches en statistiques « plus théoriques ». A ce propos, [Woodall \(2000\)](#) précise que « les cartes de contrôle sont étroitement liées aux tests d'hypothèses dans le cadre mathématique utilisé pour déterminer leurs performances statistiques. Les hypothèses associées ne sont pas requises pour les utiliser en pratique. » Ces cartes supposent que les distributions avant et après changement sont gaussiennes et la loi après changement est inconnue.

3.3.1 Principales cartes de contrôle multivariées

En SPC, les cartes de contrôle sont fréquemment utilisées pour détecter une variation anormale du centrage ou de la dispersion dans un signal supposé suivre une loi normale dite « sous-contrôle » $\mathcal{N}(\mu_0, \Sigma_0)$. Dans cette section, on présente des outils afin de détecter un changement de moyenne, également appelé décentrage du processus. La carte univariée de [Shewhart \(1931\)](#) teste simplement à chaque instant si une observation sort d'un intervalle de confiance [LCL, UCL] :

$$\begin{cases} \text{LCL} &= \mu_0 - \lambda \cdot \sigma_0 \\ \text{UCL} &= \mu_0 + \lambda \cdot \sigma_0 \end{cases}, \quad (3.43)$$

où LCL et UCL sont respectivement les limites inférieure et supérieure. σ_0 est l'écart-type et λ est usuellement fixé à 3 ([Lowry et Montgomery, 1995](#)). En dimension d , l'utilisation de cartes indépendantes sur chaque variable n'est pas adaptée. La figure 3.5 illustre l'intérêt de travailler en dimension deux pour le suivi d'une série bivariée. La carte bivariée mise en œuvre est une carte de χ^2 , qui sera présentée par la suite. Les limites de la carte univariée (vertes) et bivariée (bleu) ont été obtenues pour un taux de fausses alarmes $\alpha = 10^{-3}$. On distingue deux phases d'utilisation des cartes ([Woodall, 2000](#)) :

- *Phase I ou phase d'observation* : la densité sous-contrôle est inconnue (paramètres à estimer) et la carte est utilisée en mode hors-ligne à partir d'une base historisée ;
- *Phase II ou phase de surveillance* : la densité sous-contrôle est connue ou inconnue et la carte est utilisée en-ligne sur des données collectées séquentiellement.

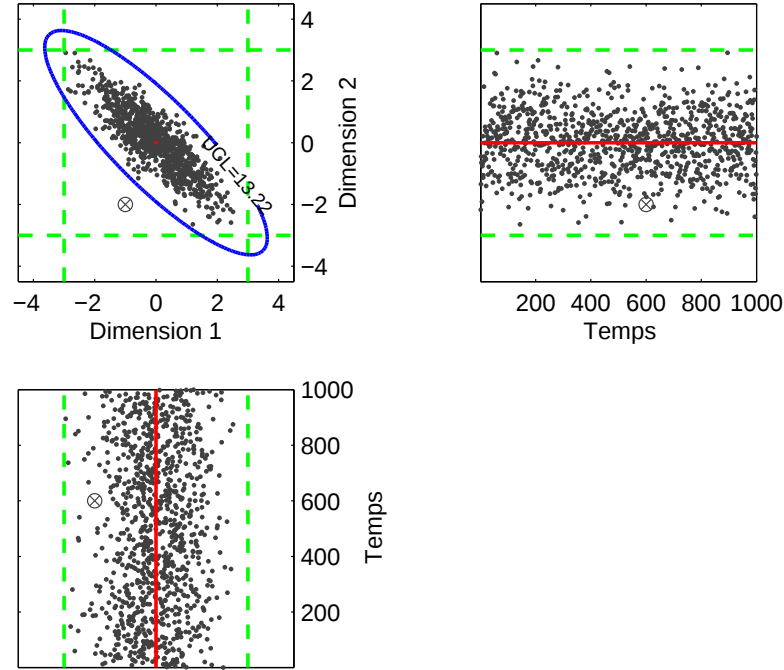


FIGURE 3.5 – Intérêt d’une approche bivariable pour des données en dimension deux.
Le point $\otimes$ sort de la région dite « sous contrôle » (ellipse en bleu)
mais reste dans l’intervalle de confiance sur chaque dimension.

On traite séquentiellement une suite $\mathbf{x}_1, \mathbf{x}_2, \dots$, où chaque observation est un vecteur de dimension d . En multidimensionnel, il n’y a pas de limite inférieure ($LCL = 0$). Le seuil de détection est la limite de contrôle supérieure UCL_α paramétrée par un taux théorique de fausses alarmes α . Pour chaque carte, une alarme est signalée lorsque la statistique de test dépasse la limite supérieure :

$$t_\alpha = \inf \{t = W, 2W, \dots : g_t \geq UCL_\alpha\}. \quad (3.44)$$

Avant d’exécuter une procédure de surveillance, il est nécessaire d’estimer le paramètre $\theta_0 = (\mu_0, \Sigma_0)$, lorsque celui-ci est inconnu, à partir d’un échantillon mesuré sur le système en bon fonctionnement. On définit les moments locaux d’ordre 1 et 2 pour chaque i^e sous-groupe de taille W :

$$\bar{\mathbf{x}}_i = \frac{1}{W} \sum_{j=0}^{W-1} \mathbf{x}_{i-j} \quad \text{et} \quad \bar{\mathbf{S}}_i = \frac{1}{W-1} \sum_{j=0}^{W-1} (\mathbf{x}_{i-j} - \bar{\mathbf{x}}_i)(\mathbf{x}_{i-j} - \bar{\mathbf{x}}_i)^T. \quad (3.45)$$

La matrice Σ_0 est estimée par la matrice de variance à échantillon groupé $\bar{\bar{\mathbf{S}}}_0$, ou *pooled sample variance matrix* (Thomaz et al., 2001), qui représente une matrice « moyenne » agrégée sur les n_t sous-groupes disponibles. D’autre part, le moment μ_0 du premier ordre est estimé par la moyenne empirique des $\bar{\mathbf{x}}_i$, notée $\bar{\bar{\mathbf{x}}}_0$. A partir d’un échantillon sous $\mathcal{H}_0$ de taille t_0 , on estime donc le paramètre $\hat{\theta}_0 = (\bar{\bar{\mathbf{x}}}_0, \bar{\bar{\mathbf{S}}}_0)$ par les approximations suivantes :

$$\bar{\bar{\mathbf{S}}}_0 = \frac{1}{n_{t_0}} \sum_{i=1}^{n_{t_0}} \bar{\mathbf{S}}_{i \cdot W} \quad \text{et} \quad \bar{\bar{\mathbf{x}}}_0 = \frac{1}{n_{t_0}} \sum_{i=1}^{n_{t_0}} \bar{\mathbf{x}}_{i \cdot W} \quad \text{avec} \quad n_t = \lfloor t/W \rfloor. \quad (3.46)$$

Cartes multivariées de Shewhart

Les cartes de type Shewhart sont conçues pour détecter des dérives de grande amplitude. Ce type de cartes s'avère sensible aux outliers (souvent dûs aux erreurs de mesure) et peuvent ainsi manquer de robustesse (Crosier, 1988; Bersimis et al., 2007).

Il existe plusieurs formulations d'une carte multivariée de Shewhart, et celles-ci diffèrent en fonction de la connaissance disponible sur la densité $p(\cdot|\theta_0)$ dite « sous contrôle », ou de la phase de contrôle. On parle notamment de carte χ^2 si le paramètre θ_0 est connu et de carte T^2 sinon. Toutes ces cartes signalent un changement à partir de la règle d'arrêt donnée en équation 3.44.

- *Contrôle sur des sous-groupes ($W > 1$) en phase II*

En phase II et lorsque le paramètre θ_0 est connu, la carte de contrôle est appelée *carte de χ^2* . En effet, la statistique g_t suit une loi de χ^2 à d degrés de liberté (Ryan, 2000). On obtient la statistique de test et la limite de contrôle suivantes :

$$\begin{cases} g_t &= (\bar{\mathbf{x}}_t - \boldsymbol{\mu}_0)^T \boldsymbol{\Sigma}_0^{-1} (\bar{\mathbf{x}}_t - \boldsymbol{\mu}_0) \\ \text{UCL}_\alpha &= \chi_d^2(\alpha) \end{cases}, \quad (3.47)$$

où $\chi_d^2(\alpha)$ est le quantile d'ordre $(1 - \alpha)$ de la loi de χ^2 avec d degrés de liberté et la statistique g_t correspond à la distance de Mahalanobis au carré entre le centroïde local $\bar{\mathbf{x}}_t$ et la moyenne $\boldsymbol{\mu}_0$ de la densité sous contrôle.

D'autre part, la carte de contrôle est appelée *carte de T^2* lorsque le paramètre θ_0 est inconnu. Ce paramètre est estimé via les approximations de l'équation 3.46.

$$\begin{cases} g_t &= (\bar{\mathbf{x}}_t - \bar{\bar{\mathbf{x}}}_0)^T \bar{\bar{\mathbf{S}}}_0^{-1} (\bar{\mathbf{x}}_t - \bar{\bar{\mathbf{x}}}_0) \\ \text{UCL}_\alpha &= c_1(d, W, n_t) \cdot F_{d, \zeta}(\alpha) \\ \text{avec } c_1(d, W, n_t) &= d(W+1)(n_t-1)/\zeta \quad \text{et} \quad \zeta = Wn_t - W - d + 1 \end{cases}, \quad (3.48)$$

où $F_{d, \zeta}(\alpha)$ est le quantile d'ordre $(1 - \alpha)$ de la loi de Fisher avec d et ζ degrés de liberté. Il est à noter que la quantité $g_t/c_1(d, W, n_t)$ suit la loi de Fisher $F_{d, \zeta}$ (Ryan, 2000).

- *Contrôle sur des sous-groupes ($W > 1$) en phase I*

En phase I et lorsque le paramètre θ_0 est inconnu, la carte est toujours de type T^2 . Toutefois, la limite de contrôle est légèrement modifiée :

$$\begin{cases} g_t &= (\bar{\mathbf{x}}_t - \bar{\bar{\mathbf{x}}}_0)^T \bar{\bar{\mathbf{S}}}_0^{-1} (\bar{\mathbf{x}}_t - \bar{\bar{\mathbf{x}}}_0) \\ \text{UCL}_\alpha &= c_2(d, W, n_t) \cdot F_{d, \zeta}(\alpha) \\ \text{avec } c_2(d, W, n_t) &= d(W-1)(n_t-1)/\zeta \end{cases}, \quad (3.49)$$

où la quantité $g_t/c_2(d, W, n_t)$ suit la loi de Fisher $F_{d, \zeta}$.

On distingue les cartes avec $W > 1$ et $W = 1$. Fixer la taille W à 1 permet une prise de décision rapide sur l'état du processus et cela dès la réception d'une nouvelle observation. Dans ce cas, la statistique de test ne permet plus de lisser les mesures par des fenêtres locales mais évalue la distance au centre μ_0 directement sur chaque observation. La diminution de W peut ainsi augmenter le nombre de fausses alarmes (voir la figure 3.2). On présente ce type de cartes en phase I et II.

- *Contrôle sur des individus ($W = 1$) en phase II*

En phase II et lorsque le paramètre θ_0 est connu, la carte est identique au cas $W > 1$. Il s'agit d'une carte de χ^2 définie par l'équation 3.47.

D'autre part lorsque le paramètre θ_0 est inconnu, ce paramètre est estimé empiriquement sur tout l'échantillon disponible :

$$\bar{S}_0 = \frac{1}{t_0} \sum_{i=1}^{t_0} (\mathbf{x}_i - \bar{\mathbf{x}}_0)(\mathbf{x}_i - \bar{\mathbf{x}}_0)^T \quad \text{avec} \quad \bar{\mathbf{x}}_0 = \frac{1}{t_0} \sum_{j=1}^{t_0} \mathbf{x}_j. \quad (3.50)$$

En phase II et lorsque le paramètre θ_0 est inconnu, la statistique est calculée à partir de l'estimation empirique $(\bar{\mathbf{x}}_0, \bar{S}_0)$. On parle de nouveau d'une carte de T^2 :

$$\begin{cases} g_t &= (\mathbf{x}_t - \bar{\mathbf{x}}_0)^T \bar{S}_0^{-1} (\mathbf{x}_t - \bar{\mathbf{x}}_0) \\ \text{UCL}_\alpha &= c_3(d, W, n_t) \cdot F_{d, W-d}(\alpha) \end{cases}, \quad (3.51)$$

avec $c_3(d, W, n_t) = \frac{d(W-1)(W+1)}{W(W-d)}$

où la quantité $g_t/c_3(d, W, n_t)$ suit la loi de Fisher $F_{d, W-d}$ (Lowry et Montgomery, 1995).

- *Contrôle sur des individus ($W = 1$) en phase I*

En phase I et lorsque le paramètre θ_0 est inconnu, la statistique de la carte de T^2 est à nouveau la distance de Mahalanobis empirique. En revanche, la limite de contrôle est un quantile de la loi Bêta :

$$\begin{cases} g_t &= (\mathbf{x}_t - \bar{\mathbf{x}}_0)^T \bar{S}_0^{-1} (\mathbf{x}_t - \bar{\mathbf{x}}_0) \\ \text{UCL}_\alpha &= c_4(W) \cdot B_{d/2, (W-d-1)/2}(\alpha) \end{cases}, \quad (3.52)$$

avec $c_4(W) = \frac{(W-1)^2}{W}$

où $B_{d/2, (W-d-1)/2}$ est le quantile d'ordre $(1 - \alpha)$ de la loi Bêta avec $d/2$ et $(W - d - 1)/2$ degrés de liberté. Il est à noter que la quantité $g_t/c_4(W)$ suit une loi Bêta $B_{d/2, (W-d-1)/2}$ (Lowry et Montgomery, 1995).

Enfin, la taille W de la fenêtre locale est généralement constante pendant toute la procédure (Bersimis et al., 2007). Une procédure est proposée par Aparisi (1996) afin de construire une carte de Shewhart avec des tailles adaptatives de fenêtres locales.

Cartes multivariées de type CUSUM

On présente ici quelques extensions multivariées de la carte du CUSUM introduite à l'origine par [Page \(1954\)](#). Les cartes de type Shewhart, présentées précédemment, peuvent détecter des changements de forte amplitude plus rapidement que les cartes de type CUSUM car celles-ci possèdent une certaine inertie (influence de points précédents). En revanche, les cartes de type CUSUM permettent de détecter des changements de plus faible amplitude, des dérives de 1 écart-type ou plus. Pour paraphraser ([Crosier, 1988](#)), « détecter une dérive de 5 écarts-types est presque trivial, alors que détecter une dérive de 0.05 écarts-types est presque impossible ».

On distingue quatre cartes multivariées de type CUSUM (somme cumulée). Pour chacune de ces cartes, la règle d'arrêt est toujours la même (équation 3.44). Les statistiques diffèrent d'une carte à l'autre et leurs distributions ne suivent pas de loi prédéfinie. Les limites de contrôle UCL sont obtenues par simulation de Monte-Carlo ou par approche à base de chaîne de Markov ([Crosier, 1988](#); [Pignatiello et Runger, 1990](#)) à partir d'un taux de fausses alarmes ou d'une ARL. Il en va de même pour le paramètre $k > 0$ de chaque carte. Il est à noter que le paramètre θ_0 peut être estimé s'il est inconnu (équation 3.46).

- *MCUSUM, ou CUSUM multivarié* ([Crosier, 1988](#))

Cette carte est la plus performante (meilleures propriétés d'ARL) des quatre et en particulier, elle est utilisée pour notre étude sur le système de freinage (chapitre 4). La statistique de test g_t , qui s'appuie sur un vecteur $\mathbf{s}_t$ défini par :

$$g_t = \sqrt{\mathbf{s}_t^T \boldsymbol{\Sigma}_0^{-1} \mathbf{s}_t}, \quad (3.53)$$

avec $\mathbf{s}_0$, le vecteur nul. On a :

$$\mathbf{s}_t = \begin{cases} 0 & \text{si } C_t \leq k_1 \\ (\mathbf{s}_{t-W} + \bar{\mathbf{x}}_t - \boldsymbol{\mu}_0) \cdot (1 - \frac{k_1}{C_t}) & \text{si } C_t > k_1 \end{cases} \quad (3.54)$$

$$\text{où } C_t = \sqrt{(\mathbf{s}_{t-W} + \bar{\mathbf{x}}_t - \boldsymbol{\mu}_0)^T \boldsymbol{\Sigma}_0^{-1} (\mathbf{s}_{t-W} + \bar{\mathbf{x}}_t - \boldsymbol{\mu}_0)}. \quad (3.55)$$

- *CUSUM de statistiques T, ou COT* ([Crosier, 1988](#))

Cette formulation est la plus directe en remplaçant la statistique des cartes de T^2 par sa somme cumulée. La statistique de cette carte s'appuie donc sur la somme cumulée d'un scalaire :

$$g_t = \max\{0, g_{t-W} + D_t - k_2\} \quad (3.56)$$

$$\text{avec } D_t = \sqrt{(\bar{\mathbf{x}}_t - \boldsymbol{\mu}_0)^T \boldsymbol{\Sigma}_0^{-1} (\bar{\mathbf{x}}_t - \boldsymbol{\mu}_0)} \text{ et } g_0 = 0. \quad (3.57)$$

La quantité D_t est simplement la distance de Mahalanobis entre la valeur moyenne de la fenêtre courante et la moyenne cible $\boldsymbol{\mu}_0$.

- *CUSUM multivarié 2, ou MC2* (Pignatiello et Runger, 1990)

Cette carte est très similaire à celle du COT. La seule différence réside dans le choix de calculer une distance D_t au carré dans la statistique de test :

$$g_t = \max \{0, g_{t-W} + D_t^2 - k_3\}, \quad (3.58)$$

où D_t^2 est défini par l'équation 3.57 et suit la loi de χ^2 à d degrés de liberté.

Selon Pignatiello et Runger (1990), la différence entre les cartes MC2 et COT peut être négligée. A l'inverse, Crosier (1988) explique qu'il y a équivalence entre les cartes de Shewhart avec une statistique T ou T^2 mais que ce n'est plus le cas pour une règle avec somme cumulée de cette statistique.

- *CUSUM multivarié 1, ou MC1* (Pignatiello et Runger, 1990)

La carte MC1 est utilisée pour la détection d'un changement de moyenne. La statistique g_t est calculée à partir d'une distance entre les sommes cumulées E_t des moyennes locales et la moyenne μ_0 de la densité sous contrôle :

$$g_t = \max \left\{ 0, \sqrt{E_t^T \Sigma_0^{-1} E_t} - k_4 \cdot m_t \right\}, \quad (3.59)$$

$$\text{avec } E_t = \sum_{i=0}^{m_t} (\bar{x}_{t-iW} - \mu_0) \quad \text{et} \quad m_t = \begin{cases} m_{t-W} + 1 & \text{si } g_{t-W} > 0 \\ 1 & \text{sinon} \end{cases}, \quad (3.60)$$

où m_t correspond au nombre de sous-groupes à l'instant t depuis le dernier renouvellement ($g_t = 0$) de la carte. L'ARL de cette carte peut être obtenue par simulation de Monte-Carlo et s'avère semblable à celle du CUSUM multivarié (Crosier, 1988).

Carte multivariée de type EWMA (Lowry et al., 1992)

La performance de la carte multivariée de type EWMA, appelée *carte multivariée pour moyennes mobiles à pondération exponentielle*, est comparable à la carte MCUSUM en terme d'ARL (MacGregor et Kourti, 1995; Bersimis et al., 2007). Ce type de carte est adapté pour le contrôle des individus ($W = 1$) et lorsque le processus est de type autorégressif (AR).

$$g_t = \mathbf{y}_t^T \Sigma_{\mathbf{y}_t}^{-1} \mathbf{y}_t, \quad (3.61)$$

$$\mathbf{y}_t = \lambda \bar{\mathbf{x}}_t + (I_d - \lambda) \mathbf{y}_{t-W} = \sum_{i=1}^{n_t} \lambda (I_d - \lambda)^{n_t-i} \bar{\mathbf{x}}_i, \quad (3.62)$$

où $\mathbf{y}_0 = \mu_0$, et, $\lambda = \text{diag}(\lambda_1, \dots, \lambda_d)$ avec $0 < \lambda_j \leq 1, \forall j = 1, \dots, d$. Lorsque $\lambda_j = 1$, on reconnaît la carte de T^2 . Plus λ_j est proche de 0, plus le processus est lissé et ainsi, la carte permet de mieux détecter des dérives de faible amplitude. Sans a priori sur les variables à surveiller, on pose $\lambda = \lambda_1 = \dots = \lambda_d$. Dans ce cas, on estime la matrice de covariance des $\mathbf{y}_t$ par :

$$\Sigma_{\mathbf{y}_t} = \frac{(1 - (1 - \lambda)^{2n_t}) \lambda}{2 - \lambda} \Sigma_0 \xrightarrow{t \rightarrow \infty} \frac{\lambda}{2 - \lambda} \Sigma_0. \quad (3.63)$$

3.3.2 Quelques extensions à base de test séquentiel

En SPC et en grande dimension ($\mathcal{O}(d) \geq 10$), les cartes de contrôles multivariées (sous-section 3.3.1) peuvent être inadaptées. Afin de réduire la dimension de l'espace de recherche, ces cartes multivariées peuvent être couplées à des méthodes de projection comme par exemple, l'Analyse en Composantes Principales (ACP), la Projection sur des Structures Latentes (PSL) (MacGregor et Kourti, 1995; Bersimis et al., 2007) ou l'Analyse en Composantes Indépendantes (ACI) (Albazzaz et Wang, 2004). A partir d'un échantillon sous $\mathcal{H}_0$, ces méthodes apprennent une transformation (combinaison linéaire des variables initiales) de manière à maximiser la covariance des données dans le nouvel espace. Après transformation, les cartes de contrôle sont employées pour détecter un changement sur les nouvelles variables mais celles-ci ne correspondent plus à aucune mesure physique contrairement aux variables de départ. Ce type d'approches avec projection est traité de manière plus détaillée dans les travaux de thèses de Verron (2007, sous-section 1.4.3). Dans le cadre de notre étude sur le systèmes de freinage (chapitre 4), seules deux variables d'intérêt ont pu être mises en évidence. Le suivi du système a donc été réalisé par des cartes de contrôle multivariées sans couplage avec une méthode de projection.

Comme le précise Woodall (2000), les hypothèses admises dans un contexte de SPC ne sont pas requises pour utiliser les cartes de contrôle en pratique. Ces hypothèses, jugées parfois restrictives, concernent notamment la stationnarité des processus et l'indépendance des données avant et après changement.

Le caractère stochastique du signal implique que l'hypothèse de stationnarité est difficile à admettre. La dynamique du processus à surveiller étant le plus souvent inconnue, il est difficile de concevoir une méthode capable de détecter un changement dont l'amplitude et la direction sont inconnues. En outre, il est usuellement supposé que les valeurs moyennes avant/après changement sont constantes alors qu'elles peuvent varier en fonction du temps (légers décalages après utilisation...). Pour répondre à ce type de problématiques, Tsung et Wang (2010) ont passé en revue plusieurs approches dont les cartes de contrôle adaptatives ayant la capacité de « mettre à jour » leurs paramètres en fonction de la situation, y compris à partir de mesures collectées en phase II. Des extensions adaptatives ont également été proposées pour les cartes de T^2 , CUSUM, EWMA et pour des cartes avec projection (ACP).

D'autre part, les observations sont communément supposées indépendantes. En pratique, les processus observés peuvent être dynamiques et ainsi, les valeurs de la série temporelle dépendent de ses valeurs aux temps précédents. A chaque instant t , il convient donc de considérer des densités conditionnelles $p(\mathbf{x}_t | \mathbf{x}_1, \dots, \mathbf{x}_{t-1}; \theta)$. Lai (1998) propose ainsi un cadre assez général pour un problème de détection en-ligne à partir de données dépendantes. Dans cette approche, de nouveaux critères de performance (extensions du critère 3.1 de Lorden) sont introduits afin de montrer l'optimalité de procédures de type GLR à fenêtres limitées. Verdier (2007) introduit une procédure de type CUSUM avec un seuil adaptatif (réestimé à chaque pas de temps). Cette règle est appliquée à une séquence

de données dépendantes et permet d'assurer un taux de fausses alarmes constant tout au long de la surveillance. En revanche, une procédure de type CUSUM suppose que les paramètres des densités $p(\cdot|\theta_0)$ et $p(\cdot|\theta_1)$ sont connus.

En pratique, on dispose de peu d'information sur les distributions du processus avant et après changement : le système est trop complexe pour le caractériser précisément, il n'y a pas d'expertise sur le procédé. La mise en place d'une procédure de type CUSUM avec garantie d'optimalité sur les données observées s'avère donc difficilement réalisable. Un test de type GLR (équation 3.15) permet en partie de répondre à cette problématique en estimant le paramètre θ_1 d'après changement, à chaque calcul de la statistique de test, par le principe du maximum de vraisemblance (MV).

Dans le cas où le paramètre θ_0 d'avant changement est inconnu, [Mei \(2006\)](#) introduit un cadre assez général qui suppose la distribution $p(\cdot;\theta_1)$ connue ou inconnue. Contrairement aux méthodes usuelles qui requièrent un taux de fausses alarmes fixé, Mei propose plusieurs méthodes optimales qui exigent de fixer un délai de détection et de minimiser le taux de fausses alarmes.

D'autre part, [Chatterjee et Qiu \(2009\)](#) propose une carte de type CUSUM quand la normalité des densités $p(\cdot|\theta_0)$ et $p(\cdot|\theta_1)$ n'est pas garantie et les paramètres θ_0 et θ_1 sont inconnus. Les limites de contrôle sont calibrées par rééchantillonnage, ou *bootstrap* ([Efron et Tibshirani, 1994](#)). Cette approche nécessite de disposer d'un grand nombre d'observations sous $\mathcal{H}_0$ (collectées en phase I) afin d'estimer correctement la forme de la distribution de la statistique. Cette stratégie d'estimation du seuil est en partie reprise dans notre étude expérimentale de la sous-section 5.4 pour la surveillance des portes.

En outre, l'optimalité asymptotique de plusieurs procédures de détection avec le couple (θ_0, θ_1) inconnu a été récemment obtenue par [Lai et Xing \(2010\)](#). Les tests sont de type SR (cadre bayésien) et GLR à fenêtres limitées (FM), et appliqués à des densités supposées exponentielles dont les paramètres $\hat{\theta}_0$ et $\hat{\theta}_1$ sont estimés par le principe du maximum de vraisemblance. Il est à noter que nous considérons une problématique semblable avec un test de type GLRFM dans notre étude sur le système de portes. Les distributions avant/après changement sont effectivement inconnues mais dans notre cas d'étude, ces densités sont estimées à partir d'un modèle de régression spécifique et adapté à une séquence de courbes, ou une séquence de *données fonctionnelles* (voir sous-section 3.4.4).

Enfin, une problématique de « décisions multiples » est souvent rencontrée dans des applications pratiques : parmi N populations synchrones, un changement survient à un instant sur l'une des populations surveillées et l'objectif est de signaler cet instant puis d'indiquer la bonne population. [Tartakovsky \(2008\)](#) propose des procédures de détection-isolation de type CUSUM et SR en combinant les approches à base de détection de changement ([Tartakovsky et al., Livre en préparation](#)) et test à hypothèses multiples ([Dragalin et al., 1999](#)). L'optimalité de ces tests multi-hypothèses de type CUSUM et SR a été obtenue au sens du critère 3.1 de Lorden.

3.4 Détection à base de reconnaissance des formes

Nous présentons dans cette dernière section quelques approches liées au problème de détection à partir d'outils (sous-section 3.4.2 et 3.4.1) et de structures de données (sous-section 3.4.4) communément étudiés en Reconnaissance des Formes (RdF, voir la section 2.2). Sans volonté d'exhaustivité, cette section a pour objet de dresser un panel de méthodes existantes partagées par plusieurs communautés statistiques malgré une terminologie différente : segmentation, détection de point de rupture, détection de changement de régime, détection d'anomalies.

3.4.1 Détection d'anomalies et classification mono-classe

Un panorama des méthodes de *détection d'anomalies* a récemment été dressé par [Chandola et al. \(2009\)](#). Cette problématique est définie comme le problème très général qui consiste à trouver des formes dans les données qui ne sont pas conformes à un comportement attendu. On utilise ici de manière analogue les problèmes de détection d'anomalies, de *classification mono-classe* ([Tax, 2001](#)), de *détection de nouveautés* ([Markou, 2003a,b](#)), et de *détection d'outliers* ([Hodge et Austin, 2004](#)). Il est à noter que la détection de nouveautés et la classification mono-classe supposent souvent l'émergence d'un nouveau mode de fonctionnement. L'hypothèse commune à toutes les approches citées dans cette partie est la suivante : « les échantillons significatifs (avant tout changement) sont regroupés ensemble et forment une région dense de l'espace de recherche ».

Les techniques à base de « voisinage » ([Chandola et al., 2009](#), section 5) estiment la densité du voisinage autour de chaque observation à partir d'une mesure de distance ou de similarité. Une anomalie est une donnée dont le voisinage est de densité trop faible. Parmi ces approches, on cite la méthode classique des *k-plus proches voisins* (*k-ppv*) : le degré d'anomalie d'une donnée est défini comme la distance cumulée à ses k plus proches voisins (parmi un ensemble d'observations fini). Puis, un seuil est défini sur le degré d'anomalie au dessus duquel une observation est déclarée anormale. D'autre part, [Breunig et al. \(2000\)](#) propose de définir le degré d'anomalie d'une donnée comme le ratio des densités locales moyennes de ses k plus proches voisins. On appelle cette mesure de similarité, le score *LOF* (*Local Outlier Factor*). Là encore, un seuil est défini sur le degré d'anomalie en dessous duquel une observation est déclarée anormale. Il est à noter que les méthodes basiques de *k-ppv* et de type *LOF* ont une complexité de $\mathcal{O}(N^2)$, où N est la taille de l'échantillon observé. Un inconvénient de ces méthodes est de devoir choisir une mesure de distance/similarité tout comme le k , ce qui n'est pas un choix trivial.

Ce type d'approche est fréquemment appliqué en maîtrise statistique des procédés (SPC) car toute anomalie peut être considérée comme une défaillance du système étudié. [Lee et al. \(2011\)](#) propose récemment une procédure de test dont la statistique est calculée à partir du score *LOF*. Cette approche est couplée avec une méthode de projection de type Analyse en Composantes Indépendantes (ACI).

A partir d'une base d'observations représentative d'un comportement attendu, l'objectif de la classification monoclasse (*one-class classification*) est d'apprendre une frontière autour de la région cible, de la *monoclasse*. Dans ce cas, l'idée est de délimiter une « hypersphère » autour des données, de décrire le support de la distribution inconnue, et cela sans chercher à minimiser son volume dans l'espace de départ. La définition de cette hypersphère s'appuie sur la description de vecteurs support (SVDD) introduite dans les travaux de thèse de [Tax \(2001\)](#). Le volume de cette sphère est minimisé dans l'espace des caractéristiques induit par une fonction à noyau (fonction vérifiant les conditions de Mercer) et représente ainsi une extension non supervisée des SVM ([Boser et al., 1992](#)). La fonction de décision obtenue retourne la valeur +1 lorsque le point observé se trouve dans la sphère et -1 partout ailleurs ([Schölkopf et al., 2001](#)). Il est à noter que cette méthode n'est pas entièrement non paramétrique car il est nécessaire de choisir une fonction à noyau et calibrer un paramètre $\nu \in]0, 1]$ ([Tax, 2001](#); [Schölkopf et al., 2001](#)). Cette étape de sélection de modèle est typiquement guidée par les connaissances disponibles sur l'application et sur les données observées.

Plusieurs études ont été récemment menées afin d'appliquer cette approche en SPC comme une alternative aux cartes de contrôle. On cite notamment les travaux de [Sun et Tsung \(2003\)](#) qui proposent des cartes à noyaux, ou *K chart*, à partir d'un noyau gaussien et d'un noyau polynomial. Dans cette approche, la limite de contrôle est calculée à partir des vecteurs support et la statistique de test est une distance à noyau entre le centre de l'hypersphère et la nouvelle observation. Les cartes à noyaux améliorent les détections comparées à des cartes T^2 lorsque les données sont non-gaussiennes. L'avantage des cartes de contrôle réside dans l'interprétation des limites de contrôle qui sont déterminées statistiquement à partir d'un taux de fausses alarmes ou d'une ARL. À l'inverse, dans une approche à noyau, la distribution de cette statistique est inconnue à cause de son caractère non paramétrique. Pour répondre à ce manque, [Sukchotrat et al. \(2009\)](#) proposent d'estimer les limites de contrôle par rééchantillonnage, ou *bootstrap* ([Efron et Tibshirani, 1994](#)). Les auteurs proposent également un détecteur de type k-plus proches voisins (kppv) à partir de la description SVDD. Cette seconde approche semble réduire considérablement les coûts de calcul.

Enfin, il est à noter que très peu de méthodes répondent simultanément aux problématiques de détection d'outliers et détection de changement. Toutefois, on cite les travaux de [Yamanishi et Takeuchi \(2002\)](#) qui cherchent à rapprocher ces deux problématiques en proposant une méthode en-ligne pour la surveillance d'une série temporelle non stationnaire. Le flux de données est modélisé par un processus auto-régressif (AR) qui est utilisé pour apprendre une densité de probabilité sur la séquence observée de manière incrémentale. L'idée est ensuite d'attribuer un degré d'anomalie à chaque nouvelle observation à partir du modèle appris.

3.4.2 Détection de changements dans des modèles markoviens

Dans cette sous-section, les observations ne sont plus indépendantes au sens strict. La séquence de données (observées ou non observées) est supposée vérifier la propriété de Markov : la probabilité qu'une de ces variables soit observée à un instant t dépend des variables aux temps précédents. Lorsque la structure du modèle est une simple chaîne et que chaque variable observée dépend uniquement de la variable précédente, on parle de *chaîne de Markov* à l'ordre 1 :

$$P(X_t = x_t | X_1 = x_1, \dots, X_{t-1} = x_{t-1}) = P(X_t = x_t | X_{t-1} = x_{t-1}). \quad (3.64)$$

Et la probabilité d'observer une instance de ce processus s'écrit :

$$P(\mathbf{x}) = P(x_1) \cdot \prod_{t=2}^T P(x_t | x_{t-1}). \quad (3.65)$$

Dans un modèle avec une chaîne de Markov cachée, ou *modèle de Markov caché (HMM)* (Rabiner, 1989) et illustré par la figure 3.6, chaque variable observée X_t à un instant t dépend uniquement d'une variable latente Z_t et chacune de ces variables Z_t dépend de la variable latente précédente. La propriété de Markov est donc vérifiée par la séquence d'états cachés Z . La probabilité jointe d'une séquence d'observations $\mathbf{x}$ et une séquence d'états $\mathbf{z}$ s'écrit :

$$P(\mathbf{x}, \mathbf{z}) = P(z_1)P(x_1|z_1) \cdot \prod_{t=2}^T P(z_t|z_{t-1})P(x_t|z_t). \quad (3.66)$$

Cette approche est largement employée pour l'analyse de points de changement à partir de données séquentielles. Un récent tutorial a été proposé par Luong et al. (2012) pour l'élaboration de HMMs appliqués au problème dit de point de changement sous la forme d'un problème de segmentation : l'objectif est d'estimer la meilleure partition de « segments » (disjoints) sur les états cachés et chaque transition entre segments correspond à un point de changement. On cite notamment les travaux de Baysse et al. (2012) qui pro-

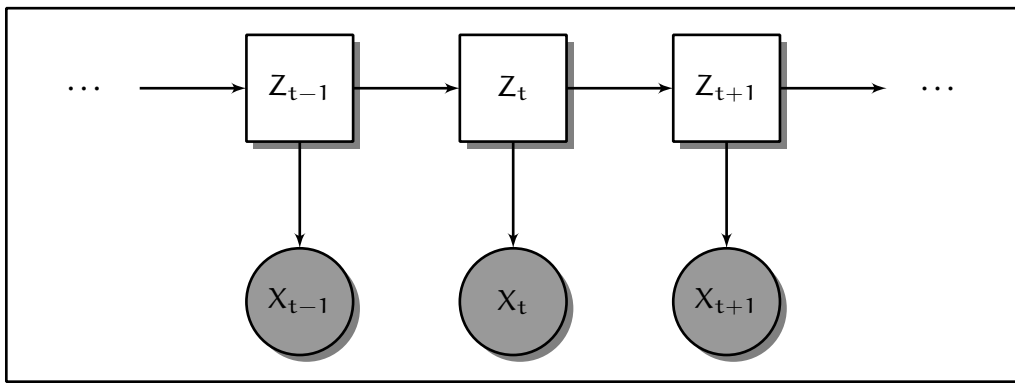


FIGURE 3.6 – Représentation graphique d'un modèle HMM.

posent d'utiliser un HMM comme outil d'aide à la maintenance pour la surveillance d'un équipement de Thalès optronique (utilisant à la fois l'optique et l'électronique).

De nombreuses extensions de type HMM ont été proposées pour ce type de modèle et le lecteur pourra se référer à l'état de l'art dressé dans les travaux de [Chamroukhi \(2010, section 2.5\)](#). On cite notamment l'extension du modèle HMM avec entrées-sorties, ou *Input-Output HMM (IOHMM)*, proposé par [Bengio et Frasconi \(1996\)](#). Dans cette approche, une troisième couche (séquence des observations « entrées », ou variable de « contexte ») $\mathbf{u} = (u_1, \dots, u_T)$ est ajoutée à la séquence des états $\mathbf{z}$ et à la séquence des sorties $\mathbf{x}$. A chaque instant t , les variables z_t (latente) et x_t (sortie) dépendent de u_t (entrée). Ce modèle génératif peut être appris par maximisation de la vraisemblance conditionnelle associée à la probabilité jointe suivante :

$$P(\mathbf{x}, \mathbf{z}|\mathbf{u}) = P(z_1|u_1)P(x_1|z_1, u_1) \cdot \prod_{t=2}^T P(z_t|z_{t-1}, u_t)P(x_t|z_t, u_t). \quad (3.67)$$

Ce modèle est plus « discriminant » qu'un HMM classique car l'analyse du point de changement est obtenue à partir de la loi conditionnelle a posteriori $P(\mathbf{z}|\mathbf{x}, \mathbf{u}) = P(\mathbf{z}|\mathbf{u})$ et non la loi marginale $P(\mathbf{z})$.

Les modèles markoviens de type HMM présentent l'inconvénient de supposer que les observations sont conditionnellement indépendantes entre elles d'une part, et de segmenter la séquence d'observations $\mathbf{x}$ à partir de la loi jointe $P(\mathbf{x}, \mathbf{z})$ d'autre part. [Do et Artières \(2006\)](#) suggèrent d'utiliser des modèles markoviens qui permettent d'estimer directement la loi conditionnelle $P(\mathbf{z}|\mathbf{x})$ mais ceux-ci perdent ainsi leur capacité à générer de nouvelles séquences d'observations. Les champs de Markov conditionnels, ou *Conditional Random Fields (CRF)*, répondent à cette problématique ([Lafferty et al., 2001](#)). Ce modèle est appris de manière discriminative en maximisant la probabilité conditionnelle suivante :

$$P(\mathbf{z}|\mathbf{x}; \lambda) = \frac{\exp(\lambda \cdot F(\mathbf{x}, \mathbf{z}))}{\sum_{\mathbf{z}^*} \exp(\lambda \cdot F(\mathbf{x}, \mathbf{z}^*))}, \quad (3.68)$$

où λ est un vecteur de poids, $F(\mathbf{x}, \mathbf{z})$ est le vecteur global de caractéristiques (choisies selon l'application) et la séquence d'états $\mathbf{z}$ vérifie la propriété de Markov suivante : $P(z_t|\{z_i\}_{i \neq t}, \mathbf{x}) = P(z_t|z_{t-1}, z_{t+1}, \mathbf{x})$ ([Sha et Pereira, 2003](#)).

Ce type de modèles nécessite généralement de disposer d'une base d'apprentissage labellisée (les observations sont étiquetées) car ces modèles sont appris dans un cadre supervisé ([Dietterich, 2002](#)). Plus généralement, ce type d'approches appartient à la famille des modèles markoviens avec changements de régimes, ou *Markov switching models* ([Frühwirth-Schnatter, 2006](#)).

3.4.3 Détection de changements par la sélection de modèles

En classification non supervisée (*clustering*), sélectionner le bon nombre de clusters dans un échantillon de données est toujours un problème ouvert (von Luxburg, 2010). On parle de *sélection de modèles* (sous-section 2.2.2) car l'identification des hyperparamètres (nombre de groupes en clustering, ordre d'un modèle de régression,...) dimensionne la structure du modèle (Schwarz, 1978) et précède l'apprentissage (estimation des paramètres) du modèle. On s'appuie généralement sur des critères issus de la théorie de l'information (AIC, BIC, ICL,...) qui privilégient l'adéquation du modèle à la distribution jointe des variables observées et celles de classe.

Dans cette sous-section, on associe la sélection de modèles non supervisée à la détection de changement structurel (point de rupture dans le modèle). Lorsqu'un système est en bon état de fonctionnement, un ou plusieurs clusters sont identifiés pour représenter correctement les données qu'il fournit. Lorsqu'un nouveau mode de fonctionnement défaillant apparaît, ces clusters ne suffisent plus à la représentation des données et les méthodes de sélection de modèle sont à même de l'indiquer.

Dans le cas où le nombre de segments K est inconnu, Yang et Kuo (2001) proposent une procédure de segmentation binaire capable de détecter plusieurs changements dans une séquence en s'appuyant sur le critère BIC. Lorsque K augmente, ce critère est de plus en plus difficile à calculer car le modèle devient de plus en plus complexe à évaluer. Les auteurs suggèrent de tester séquentiellement ($K = 1$ contre $K = 2$) afin de faciliter le calcul du critère. Cette approche séquentielle est similaire à la procédure que nous proposons dans la section 5.3. En effet, le détecteur proposé teste aussi deux hypothèses séquentiellement et permet de détecter plusieurs changements. A ce propos, Kim et al. (2009) introduisent une procédure de tests séquentiels basés sur des rapports de vraisemblance appliquée à un problème de segmentation univarié. La sélection des modèles obtenus est comparée à la sélection des critères de référence AIC, BIC et GCV (validation croisée généralisée). Enfin, Arlot et al. (2012) proposent récemment une méthode de segmentation basée sur la sélection de modèle, où le nombre de changements est inconnu. Il est montré expérimentalement sur des flux vidéos que cette approche permet de détecter des changements de distribution dans des moments d'ordre élevé, y compris avec les deux premiers moments (moyenne et variance) constants.

D'autre part, certains travaux ont permis d'étendre des outils de classification lorsque les données sont non stationnaires et les classes sont évolutives temporellement. On cite notamment les travaux de Boubacar (2006) qui proposent de partitionner dynamiquement les données à partir de règles d'adaptation récursives. Dans ce contexte, deux types de classifieur (modèles de mélange et méthodes à noyau) sont étudiés en terme de convergence et de complexité. Traoré (2010) étend une partie de ces travaux dans le cadre de la maintenance prévisionnelle (voir la sous-section 2.1.1) de systèmes dynamiques évolutifs. Un modèle polynomial de dégradation est utilisé par le processus de

pronostic afin de prédire la probabilité d'occurrence d'une défaillance y compris pour des dérives lentes de fonctionnement.

Dans notre étude sur les portes (chapitre 5), on observe des séquences de données fonctionnelles multivariées. Chacun de ces signaux est décomposable en segments contigus et décrit une phase opérationnelle de la manœuvre d'une porte en fonctionnement. Ces courbes sont modélisées par un modèle de régression avec changement de régime. Ce modèle génératif est dimensionné à partir des critères BIC et ICL, et cette étude est notamment décrite dans la sous-section 5.5.3.

La prochaine sous-section présente quelques stratégies de détection de changement(s) dans une séquence de données fonctionnelles. Cette problématique a émergé récemment car elle est présente dans de nombreux cas d'applications réelles.

3.4.4 Détection de changements dans une séquence de données fonctionnelles

Cette partie introduit quelques procédures de surveillance afin de détecter un (plusieurs) changement(s) dans un flux de données fonctionnelles (voir section 2.3). Dans la littérature, on parle souvent de *linear/non linear profile monitoring* pour désigner ce type de méthodes. En pratique, il devient de plus en plus répandu de surveiller un processus industriel en le caractérisant par des relations fonctionnelles entre une variable de type réponse et une ou plusieurs variables explicatives (Tsung et Wang, 2010). Cependant, la communauté en maîtrise statistique des procédés (SPC) ne s'est emparée de cette problématique que dans les années 2000. Jusque là, très peu de travaux de recherche ont été menés sur la surveillance statistique d'une séquence de courbes et l'adaptation des cartes de contrôle aux données fonctionnelles. Woodall et al. (2004) et Woodall (2007) dressent un état de l'art de ce type de méthodes en SPC tout en soulignant leur facilité d'interprétation et émettent quelques recommandations sur leur utilisation dans le cadre d'une surveillance en phases I et II (voir la sous-section 3.3.1). Les stratégies passées en revue, extraient à partir des courbes linéaires ou non linéaires, de simples pentes ou des paramètres plus élaborés (coefficients d'ondelettes et de splines). Récemment, l'ouvrage de Noorossana et al. (2012) présente un panorama des techniques et méthodes liées à la surveillance de signatures. On utilise ici de manière analogue les concepts de données fonctionnelles, données longitudinales, signatures, courbes ou trajectoires.

Ce problème a notamment été étudié dans une approche fréquentielle et au départ, dans un cadre non séquentiel. En particulier, on cite les travaux de Fan et Lin (1998) qui proposent d'utiliser un test d'hypothèses de manière à détecter un changement entre deux ensembles de courbes. Une transformée de Fourier discrète est appliquée sur chaque courbe afin de compresser les données sous la forme de coefficients indépendants et normaux dans le domaine fréquentiel. Les auteurs suggèrent alors d'utiliser un test (version adaptative du test de Neyman) afin de limiter le nombre de coefficients puis

de comparer la statistique de ce test à un seuil de détection. Dans une extension séquentielle et en phase II, un choix restreint de coefficients peut masquer la direction d'une dérive et ainsi, conduire à une non-détection. Plus récemment, [Jeong et al. \(2006\)](#) ont introduit une procédure de surveillance en-ligne basée sur une transformée en ondelettes en phase II et à partir d'une carte de type T^2 . Comme pour la transformée de Fourier évoquée précédemment, il est nécessaire de sélectionner un petit nombre de coefficients d'ondelettes avant d'exécuter la carte de contrôle. Les auteurs justifient le choix des coefficients sélectionnés en précisant que leur objectif est de surveiller localement certains changements (via certains coefficients) sur des trajectoires. Cette idée de localisation est reprise dans le chapitre 5 dans le sens où nous proposons une approche permettant de localiser les défauts par le biais d'une statistique dédiée.

Dans le cadre d'une surveillance en phase I, [Moguerza et al. \(2007\)](#) proposent une méthode pour la détection d'outliers dans une séquence de signatures non linéaires. Cette stratégie s'appuie sur une modélisation non paramétrique (régresseur régularisé de type SVM) des courbes et caractérise un « intervalle de confiance fonctionnel » où les deux limites de contrôle sont représentées par les quantiles α et $(1 - \alpha)$ des courbes lissées. Il est nécessaire d'estimer ces courbes limites sur un échantillon exhaustif car une anomalie est signalée dès qu'une signatures (lissée) dépasse une des courbes limites. Cette méthode s'avère assez sensible aux outliers et peut difficilement être envisageable pour une surveillance séquentielle en phase II. En revanche, l'idée de surveiller les courbes et non plus les paramètres d'une modélisation est reprise dans nos travaux sur les portes.

D'autre part, [Shiau et al. \(2009\)](#) suggèrent également de lisser les courbes comme prétraitement, avant d'exécuter une procédure de test. Afin de débruiter les signatures, les auteurs évoquent principalement les techniques de lissage par splines et B-splines (cubiques). La procédure de test en phase I consiste à choisir un sous-espace avec une ACP puis de tester les valeurs projetées avec une carte de type T^2 en phase II. La méthode est évaluée en terme d'ARL sur une séquence de courbes simulées avec bruits additifs.

Dans une approche de type RdF, [Pacella et Semeraro \(2011\)](#) proposent une procédure non supervisée à partir d'un réseau de neurones. Ce modèle est appris au préalable sur des courbes supposées sous contrôle afin d'apprendre un mode de bon fonctionnement. Cette approche s'avère très compétitive mais ici encore, la performance de ce modèle en phase II (dans un cadre séquentiel) dépend fortement de l'apprentissage du modèle.

Dans ce type de problématique, chaque observation est un échantillon fini d'individus collecté sur une fenêtre de temps. Il est possible de considérer la séquence de courbes comme une séquence composée de longs vecteurs univariés puis d'utiliser des détecteurs traditionnels comme les cartes de contrôle multivariées (sous-section 3.3.1) mais cette approche ne prend pas en compte l'aspect temporel de chaque courbe. Il est préférable de représenter ces courbes comme des valeurs discrétisées de fonctions et ainsi d'exploiter les relations dynamiques entre les différentes variables. Les techniques usuelles d'*analyse*

de données fonctionnelles (FDA) (Ramsay et Silverman, 2005; Ramsay et al., 2009) et en particulier, les extensions fonctionnelles des méthodes de projection (ACP, ICA) offrent un panel d'outils adaptés à ce type de données. A ce propos, Yu et al. (2012) proposent une nouvelle procédure de détection d'outliers basée sur une méthode d'ACP fonctionnelle (FPCA). Dans le nouvel espace qui explique 85% de la variance, la statistique de test calcule une distance entre les projections de chaque nouvelle courbe et les projections d'une courbe moyenne. Les auteurs précisent que le choix d'un tel sous-espace peut conduire à un test sous optimal.

Enfin, Aston et Kirch (2012) détaillent de nombreuses procédures de détection de changements dans une séquence de données fonctionnelles. Les propriétés de ces tests sont passées en revue dans le cadre d'un problème de détection à un seul changement (AMOC) et un problème de détection de changements épidémiques (le système retourne à son état nominal après une dérive de fonctionnement).

3.5 Conclusion

Ce chapitre introduit la problématique de détection de point de changement dans le cadre de la théorie des tests d'hypothèses (séquentiels ou non). Ce problème est rencontré dans plusieurs domaines statistiques et demeure incomplètement résolu dans de nombreux cas d'études. Toutefois, certains détecteurs de la littérature ont été proposés avec des propriétés d'optimalité dans des situations bien précises. Une liste non exhaustive de critères d'optimalité a ainsi été passée en revue et les tests de référence associés ont été présentés.

La surveillance en-ligne de processus industriels est généralement traitée en maîtrise statistique des procédés. Depuis la seconde moitié du XX^e siècle, les méthodes de type carte de contrôle sont largement employées pour le contrôle qualité. Leur rapidité de mise en œuvre, leur facilité d'interprétation et leurs garanties de performance statistique expliquent la popularité de ce type de méthodes. Les principales cartes de contrôles multivariées ont ainsi été présentées et quelques unes de leurs extensions adaptées à des situations souvent rencontrées en pratique (dépendance, processus non stationnaires,...).

Enfin, plusieurs parallèles ont été établis entre la détection de point de changement et d'autres problèmes statistiques (segmentation dans des modèles markoviens, sélection de modèles,...). Des travaux de recherche ont récemment abordé la détection séquentielle dans un flux de données fonctionnelles. Ce thème émerge fortement dans de nombreuses applications réelles, ce qui implique l'introduction d'une nouvelle classe d'outils adaptés. En particulier, une nouvelle stratégie séquentielle de détection et isolation dans une séquence de courbes est proposée dans le chapitre 5.

Chapitre 4

Détection d'anomalies à l'aide d'un modèle physique sur un système de freinage

Sommaire

4.1	Introduction	90
4.1.1	Différents systèmes de freinage	90
4.1.2	Description du système de freinage étudié	91
4.1.3	Démarche des travaux scientifiques	93
4.2	Modélisation dynamique du véhicule et des freins	94
4.2.1	Modélisation longitudinale du bus	94
4.2.2	Estimation de la pente	102
4.2.3	Représentation du fonctionnement normal du système de freinage	104
4.3	Stratégie séquentielle de détection d'anomalies	106
4.3.1	Description globale	106
4.3.2	Détection de défauts	107
4.4	Expérimentations sur données réelles	108
4.4.1	Acquisition des données de fonctionnement	108
4.4.2	Description des données réelles collectées	110
4.4.3	Dégradations simulées du système de freinage	113
4.4.4	Résultats expérimentaux et estimation du seuil de détection	116
4.5	Conclusion	119

4.1 Introduction

4.1.1 Différents systèmes de freinage

Historiquement, les premiers systèmes de freinage automobile étaient dérivés des systèmes employés sur les chariots hippomobiles, à savoir des freins à tambour externes placés uniquement sur les roues arrières (Puhn, 1985). Cela s'explique par la complexité d'installer des freins sur des essieux orientés et la peur de générer un tangage intempestif. Le frein interne à tambour est proposé par Renault en 1902. Puis, le freinage aux quatre roues basé sur un système hydraulique est introduit dans les années 1920. Les premiers freins à disque ont fait ensuite leur apparition fin 1940 début 1950. Le tableau 4.1 propose une comparaison entre les freins à tambour/à disque.

Frein	Tambour	Disque
Avantages	- Commande de frein à main plus simple	- Remplacement rapide des garnitures - Pas de déformation - Bon refroidissement
Inconvénients	- Mauvais refroidissement - Dilatation et déformation	- Matériel coûteux - Matériel bruyant

TABLE 4.1 – Comparaison entre tambours et disques de freins.

Au-delà de la popularité des systèmes ABS/ASR et de quelques raffinements, les systèmes de frein hydrauliques et pneumatiques ont assez peu changé depuis les années 1950 (Vlacic et al., 2001, chapitre 12). Un système de freinage est composé de plusieurs éléments (figure 1.6) et caractérisé par un circuit fermé, alimenté par un fluide pneumatique (faible pression $\sim 0-14$ bar) ou hydraulique (forte pression $\sim 0-30$ bar). Les systèmes de frein hydrauliques sont souvent utilisés dans des voitures particulières tandis que les systèmes pneumatiques sont utilisés dans des véhicules commerciaux tels que des autobus ou des camions (Limpert, 2011). Un système hydraulique peut rapidement perdre de son efficacité à mesure que le circuit surchauffe (dû au poids du véhicule et à un usage fréquent). Le tableau 4.2 présente une rapide comparaison entre les deux technologies :

Frein	Pneumatique	Hydraulique
Avantages	- Matériel peu coûteux - Peu de retour pédale	- Fluide incompressible - Force élevée et rapide
Inconvénients	- Fluide compressible - Matériel bruyant - Entretien fréquent	- Surchauffe - Matériel coûteux - Pollution en cas de fuite

TABLE 4.2 – Comparaison d'un système de frein pneumatique/hydraulique.

4.1.2 Description du système de freinage étudié

Le tableau 4.3 présente les commandes utilisées pour le freinage des autobus de l'étude :

Désignation	Commande	Activation	Fonctionnement
Ralentisseur	Électrique	Pédale de frein (appui de 0% à 18%)	Intégré à la boîte de vitesse
Frein principal	Pneumatique	Pédale de frein (appui de 18% à 100%)	Frottement des disques avants et arrières
Frein de secours / stationnement	Mécanique	Levier sur tableau de bord	Frottement des disques arrières
FIPO, frein de stationnement temporaire	Électrique	Interrupteur sur tableau de bord	Frottement des disques arrières

TABLE 4.3 – Quatre commandes de freinage pour un bus Citelis (Irisbus).

La boîte de vitesse intègre un ralentisseur hydrodynamique primaire. La mesure du couple de ce ralentisseur est a priori accessible sur le CAN J1939 (voir sous-section 1.3.3). Toutefois, cette information n'a pas été disponible pour nos travaux.

On s'intéresse au suivi de fonctionnement du frein principal (commande pneumatique); la figure 4.1 présente l'architecture du circuit pneumatique dont l'air comprimé est distribué à chaque élément par des conduites de frein. Ce type de circuit est essentiellement composé des six éléments suivants :

- *Compresseur* : produit l'air comprimé à une pression nominale de 14 bars et un débit de 700 L/min (situé dans le compartiment moteur);
- *Dessiccateur (avec valve de protection intégrée)* : élimine toute humidité et transmet l'air à une pression nominale de 12 bar;
- *Réservoir* : emmagasine l'air comprimé dans plusieurs réservoirs (un réservoir « avant » de 30 L et deux réservoirs « arrière » de 15 L chacun);
- *Robinet de frein* : pédale de frein utilisée afin de commander l'arrivée d'air comprimé (8 bars max) des réservoirs jusqu'aux cylindres (située au poste de conduite);
- *Cylindre* : transmet la force exercée par l'air comprimé à la timonerie mécanique. Contrairement aux cylindres simples (avant), les cylindres doubles (arrière) sont sollicités pour bloquer les roues à l'arrêt du bus;
- *Disque* : lorsque l'air comprimé exerce une pression sur le piston d'un cylindre, les garnitures viennent frotter sur le disque (figure 1.7(b)) entraînant l'arrêt du bus.

Il est à noter que le système ABS a pour objet de maintenir une bonne capacité de freinage en modulant la pression dans les cylindres de frein arrière via des capteurs électropneumatiques en évitant leur blocage. L'ABS est activé en fonction de la vitesse mesurée sur chaque roue.

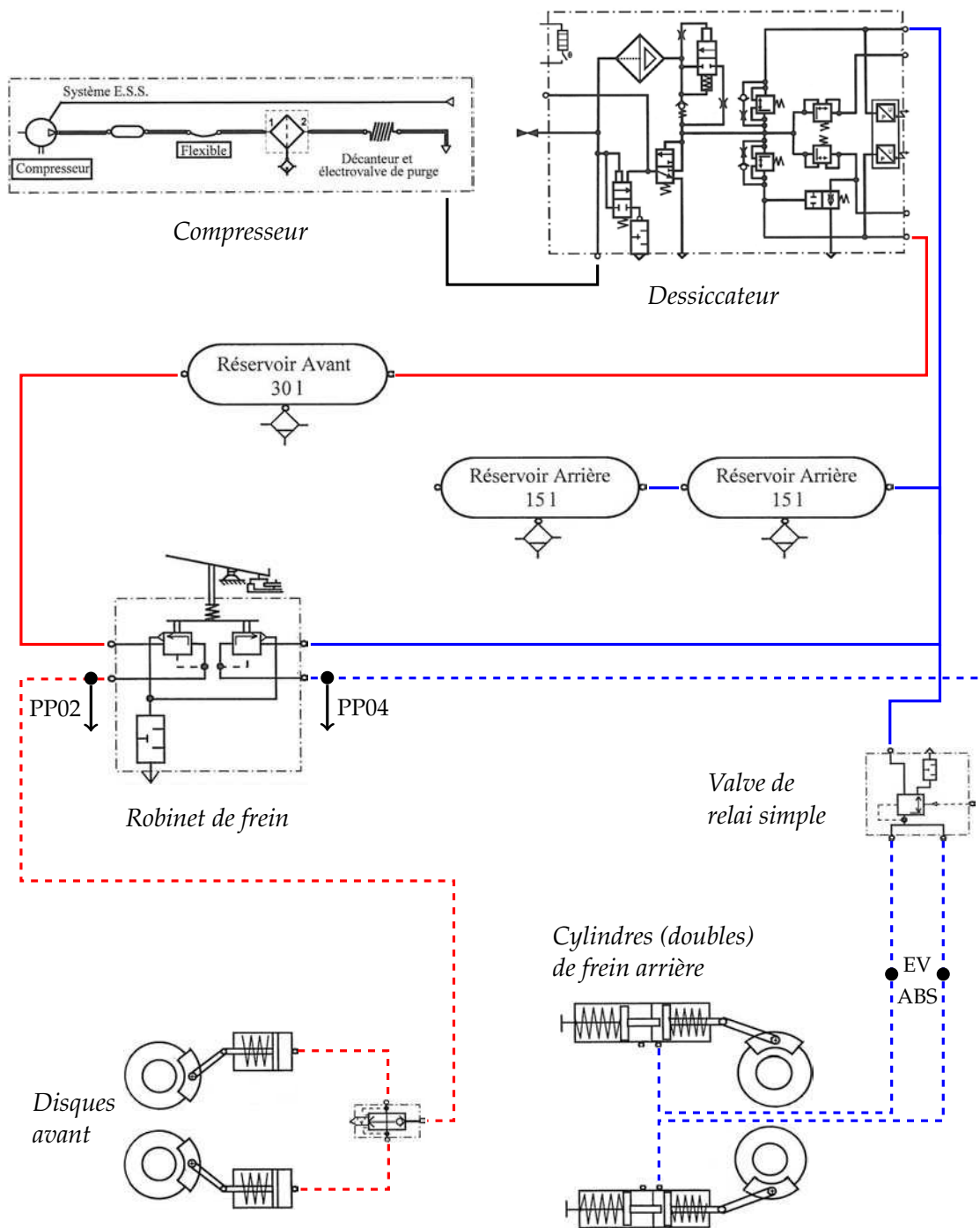


FIGURE 4.1 – Schéma de principe du système de freinage principal et circuit pneumatique d'un autobus Citelis (IRISBUS).
Ce circuit est composé de deux sous-circuits indépendants : **avant** et **arrière**.
Des mesures de pression sont collectées par des capteurs additionnels (en PP02 et PP04) installés en sortie du robinet de frein.

Outre les visites réglementaires, un système de freinage peut être évalué suivant deux niveaux d'inspections non destructives (Shaffer et Alexander, 1995a) :

Inspection « visuelle » : observation de l'épaisseur de la garniture des plaquettes, de l'usure et de la fuite des conduites de frein ;

Inspection « approfondie » : analyse de plusieurs indicateurs physiques dont la force et le couple de freinage, la décélération du véhicule,...

Le premier niveau est actuellement employé par le réseau STRAV de Transdev ; les signalements sont souvent remontés directement par les conducteurs. Cette stratégie souffre d'une certaine subjectivité et peut s'avérer facilement chronophage pour les équipes de maintenance. Des outils fournis par les constructeurs sont parfois disponibles pour aider à identifier la cause de certaines pannes, mais ceux-ci ne permettent pas de détecter une anomalie avant la dégradation du système.

4.1.3 Démarche des travaux scientifiques

L'objectif de ce chapitre est de présenter un nouvel outil d'aide à la maintenance dédié à la surveillance d'un système de freinage pneumatique à partir d'une analyse de données collectées via des capteurs embarqués. Il est à noter que cette approche peut aisément être déclinée pour un système hydraulique. Cette stratégie, jugée plus objective, correspond au second niveau d'inspection, dite « approfondie », et permettra à terme de diminuer les coûts de maintenance tout en augmentant les temps de disponibilité des véhicules.

Le développement de cet outil s'appuie sur une modélisation dynamique du véhicule, présentée dans la section 4.2. Cette approche nous permet d'extraire le couple de freinage à partir de données CAN. Le fonctionnement des freins est caractérisé par un modèle de régression entre ce couple de freinage et la pression de freinage délivrée dans les cylindres. Les variables extraites présentent l'avantage d'être physiquement interprétables. L'estimation de chaque variable physique fera l'objet d'une description détaillée. On décline alors une stratégie de détection séquentielle en trois étapes, dans la section 4.3, qui s'appuie sur le modèle de régression. Chaque plage de freinage est résumée par un couple de variables numériques surveillé séquentiellement. Les cartes de contrôle, décrites dans la section 3.3.1, représentent des méthodes de détection adaptées à notre problématique. Nous avons choisi d'implémenter les cartes multivariées de Shewhart et de CUSUM afin de détecter d'éventuelles anomalies sur le système de freinage. Dans la section 4.4, on présentera une étude expérimentale menée sur des jeux de données collectés dans des bus en exploitation dans le cadre du UC de Brunoy. Des dégradations simulées ont été ajoutées aux données réelles afin d'évaluer notre stratégie de détection. Les défauts choisis sont réalistes et ont permis de caractériser quelques dérives de fonctionnement des freins. Cette étude nous a permis de mettre en concurrence les deux cartes de contrôle et d'évaluer leurs performance en terme de détection à partir de la courbe ROC. Enfin, il est proposé une méthode pour estimer une limite de contrôle des détecteurs adaptée aux distributions des données étudiées.

4.2 Modélisation dynamique du véhicule et des freins

Cette section établit une relation entre deux indicateurs physiques : la pression de freinage P_b (mesurée à partir de capteurs additionnels en sortie du pédalier de frein – voir figure 4.1) et le couple de freinage global du véhicule T_b (estimation basée sur une modélisation dynamique du bus). Nous commençons par présenter un modèle longitudinal du bus (Nouvelière, 2002) alimenté par quelques mesures (données CAN) et constantes physiques. Puis, nous décrivons dans la sous-section 4.2.2 comment estimer la pente de la route, une mesure difficilement calculable pendant l'exploitation des bus sans connaissance a priori.

Les caractéristiques des véhicules pertinentes pour l'étude sont les suivantes :

- Véhicules : autobus IRISBUS Citelis Euro IV (standards et articulés)
- Moteur : diesel, d'un couple maximal de 1 100 Nm à 1 050 tr/min
- Boîte de vitesse : ZF 5HP502 auto (bus standard) / ZF 5HP504 auto (bus articulé)
- Géolocalisation : GPS avec longitude, latitude, altitude et boussole

4.2.1 Modélisation longitudinale du bus

Le tableau 4.4 introduit les notations des variables adoptées dans tout le chapitre 4 :

Description	Variable	Valeur/Unité
Accélération de la pesanteur	g	$9.80665 \text{ m} \cdot \text{s}^{-2}$
Accélération du véhicule	a	$\text{m} \cdot \text{s}^{-2}$
Coefficient de résistance au roulement	C_r	0.0065
Coefficient de résistance aérodynamique	C_a	2.8712775
Couple moteur	T_e	$\text{N} \cdot \text{m}$
Couple de freinage	T_b	$\text{N} \cdot \text{m}$
Force de traction	F_{trac}	N
Force de résistance au roulement	F_{drag}	N
Force de gravité	F_{grav}	N
Force de frottement aérodynamique	F_{aero}	N
Masse d'une roue	m_w	96.1 kg
Masse totale du véhicule	m	kg
Moment d'inertie du moteur	J_e	$1.47 \text{ kg} \cdot \text{m}^2$
Moment d'inertie total ramené sur l'axe moteur	J_{eq}	$\text{kg} \cdot \text{m}^2$
Pente de la route	θ	rad
Rapport de démultiplication totale	R_g	
Rayon d'une roue	r_w	0.4835 m
Vitesse de rotation d'une roue	ω_w	$\text{rad} \cdot \text{s}^{-1}$
Vitesse de rotation du moteur	ω_e	$\text{rad} \cdot \text{s}^{-1}$
Vitesse (longitudinale) du véhicule	v	$\text{m} \cdot \text{s}^{-1}$

TABLE 4.4 – Notations adoptées pour la modélisation longitudinale d'un autobus.

A partir des travaux de [McMahon et al. \(1992\)](#), [Hedrick et al. \(1997, section 5.1\)](#) et d'autres, on dérive une modélisation pour le contrôle longitudinal d'un autobus. Cette modélisation s'appuie sur le *Principe Fondamental de la Dynamique (PFD)* et notamment sur l'hypothèse que l'accélération longitudinale du véhicule peut être contrôlée directement. La non prise en compte d'un mouvement latéral peut sembler être un choix restrictif au départ mais le contrôle longitudinal du véhicule permet en pratique de réaliser un bon compromis entre simplicité et précision. Cette approche a été validée en condition réelle par [Song et Hedrick \(2004\)](#) et nous avons pu également vérifier quelques hypothèses simplificatrices de modélisation (voir figure 4.2).

On applique le PFD en translation sur le châssis (équation 4.1 et figure 4.3) et sur les roues avant (*front*) / arrière (*rear*) (équations 4.2, 4.3) du véhicule. La mise en équation du PFD dans une direction longitudinale est alors décrite par les trois équations suivantes :

$$\begin{cases} m_c \cdot a = F_{\text{reac},f} + F_{\text{reac},r} - F_{\text{aero}} - m_c \cdot g \cdot \sin(\theta) & (4.1) \\ m_{wf} \cdot a = F_{\text{trac},f} - F_{\text{reac},f} - m_{wf} \cdot g \cdot \sin(\theta) & (4.2) \\ m_{wr} \cdot a = F_{\text{trac},r} - F_{\text{reac},r} - m_{wr} \cdot g \cdot \sin(\theta) & (4.3) \end{cases}$$

où a est l'accélération longitudinale du véhicule, $F_{\text{trac}} = F_{\text{trac},f} + F_{\text{trac},r}$ est la force de traction au niveau des roues du véhicule, $F_{\text{reac}} = F_{\text{reac},f} + F_{\text{reac},r}$ est la force de réaction, θ est la pente de la route (sous-section 4.2.2), et la force de frottement aérodynamique (de la surface frontale du châssis dans l'air) est calculable par la relation :

$$F_{\text{aero}} = C_a \cdot v^2. \quad (4.4)$$

En sommant les équations 4.1, 4.2 et 4.3, on obtient l'équation du mouvement longitudinal du véhicule en translation :

$$m \cdot a = \sum F = F_{\text{trac}} - F_{\text{aero}} - F_{\text{grav}}, \quad (4.5)$$

où $m = m_c + m_{wf} + m_{wr}$ est la masse totale (châssis et roues) du véhicule, et la force de gravité du véhicule est calculable par la relation :

$$F_{\text{grav}} = m \cdot g \cdot \sin(\theta). \quad (4.6)$$

De manière analogue, il est possible d'appliquer le PFD en rotation au niveau des roues avant /arrière du véhicule :

$$\begin{cases} J_{wf} \cdot \frac{d}{dt} \omega_{wf} = -r_w \cdot F_{\text{trac},f} - r_w \cdot F_{\text{drag},f} - T_{bf} & (4.7) \\ \overline{J_{wr}} \cdot \frac{d}{dt} \omega_{wr} = -r_w \cdot F_{\text{trac},r} - r_w \cdot F_{\text{drag},r} - T_{br} + T_d & (4.8) \end{cases}$$

où J_{wf} est le moment d'inertie des roues avant, $\overline{J_{wr}}$ représente à la fois le moment d'inertie des roues arrières et le moment d'inertie du moteur (après transmission), $F_{\text{drag}} = F_{\text{drag},f} + F_{\text{drag},r}$ est la force de résistance au roulement, et T_d désigne le couple au niveau des roues motrices (roues arrières car les bus étudiés sont propulsés).

Afin de dériver une modélisation longitudinale du véhicule avec un compromis entre simplicité et précision, on admet les quatre hypothèses simplificatrices suivantes :

1. Les couples au niveau du moteur T_e et des roues motrices T_d sont liés par une relation linéaire sans retard :

$$T_d = \frac{T_e}{R_g}, \quad (4.9)$$

où R_g est le rapport de démultiplication total (du moteur aux roues) du véhicule.

2. Le mouvement latéral (trajectoire longitudinale, supposée en ligne droite) est négligé. On réduit alors l'accélération du véhicule selon l'équation suivante :

$$a = \frac{d}{dt}v. \quad (4.10)$$

3. L'embrayage est verrouillé (rapport de boîte constant) (Nouvelière, 2002, p.103). Les vitesses de rotation du moteur ω_e et des roues ω_w sont reliées linéairement :

$$\omega_w = R_g \cdot \omega_e. \quad (4.11)$$

Sans cette hypothèse, il faut considérer le frottement entre le moteur et l'arbre de transmission ; il faut alors introduire plusieurs états additionnels (Hedrick et al., 1997, p. 48) ce qui complexifie notablement la modélisation de la dynamique.

4. Le glissement entre pneumatiques et chaussée est négligé (Hedrick et al., 1997, p. 47) – validé expérimentalement par McMahan et al. (1992). La vitesse longitudinale v et la vitesse de rotation des roues ω_w sont reliées linéairement :

$$\begin{aligned} v &= r_w \cdot \omega_w \\ &= r_w \cdot R_g \cdot \omega_e. \end{aligned} \quad (4.12)$$

Nous avons validé expérimentalement cette égalité pour les rapports de boîte 2 à 5. La figure 4.2 met en évidence la bonne approximation de la vitesse longitudinale par l'équation 4.12. En effet, l'erreur quadratique moyenne (EQM) en vitesse est inférieure à 2.5 km/h (soit moins de 0.7 m/s) pour les rapports de boîte 2 à 5. En revanche, l'approximation de v n'est plus vérifiée dans le premier rapport de boîte (faible vitesse longitudinale du véhicule) à cause de l'embrayage et du convertisseur de couple, comme illustré par les figures 4.2(a) et 4.2(d). L'estimation de la vitesse par cette équation n'est pas non plus réaliste dans le cas d'une forte accélération ou lors d'un freinage d'urgence.

Enfin, il est important de noter que la décélération totale du véhicule est plus représentative que les couples à chaque roue pour la modélisation dynamique du véhicule en phase de freinage. C'est particulièrement vrai lorsque le freinage activé par une seule commande (Vlacic et al., 2001) qui est la pédale de frein dans notre cas d'étude.

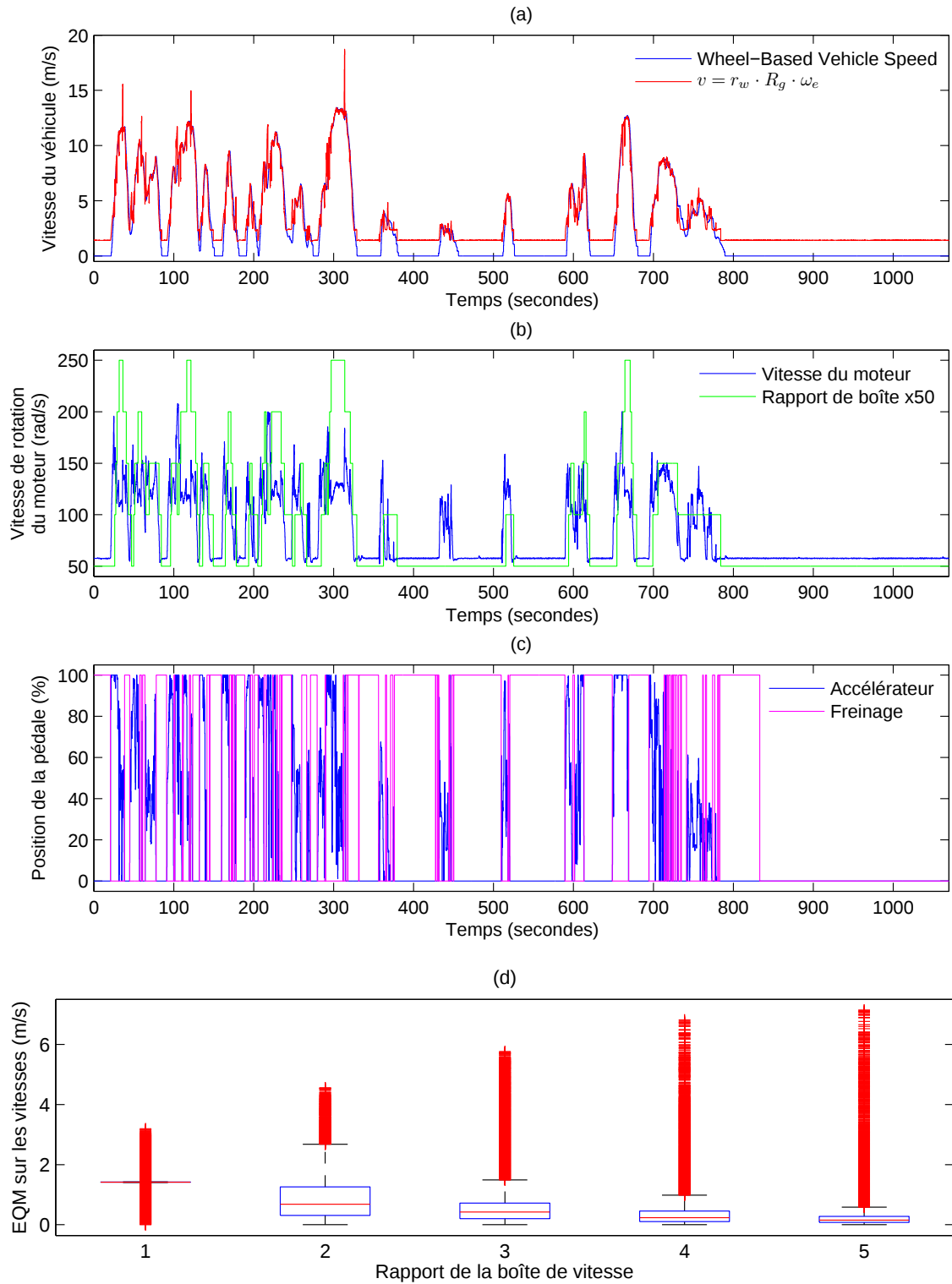


FIGURE 4.2 – Validation expérimentale des hypothèses 3 et 4 (équations 4.11 et 4.12) sur un échantillon de données extraites d'un bus standard en novembre 2012. Visualisation de la vitesse du véhicule mesurée / estimée (a), vitesse de rotation du moteur et rapport de boîte (b), positions des pédales de frein et d'accélérateur (c), et erreur quadratique moyenne entre vitesse mesurée/estimée en fonction du rapport de boîte engagé (d) - erreurs médianes : 1.41, 0.68, 0.42, 0.23, 0.15.

D'après les hypothèses 2 et 4, l'accélération angulaire est calculable à partir de l'accélération longitudinale : $\frac{d}{dt}\omega_w = \frac{a}{r_w}$. En utilisant cette égalité et les équations 4.7 et 4.8, on établit l'équation de la force de traction F_{trac} :

$$F_{\text{trac}} = -\frac{\overline{J_{wr}} + J_{wf}}{r_w^2}a + \frac{T_e/R_g - T_b}{r_w} - F_{\text{drag}}. \quad (4.13)$$

A partir des quatre hypothèses et des équations 4.5 et 4.13, on obtient l'équation longitudinale de la dynamique appliquée aux bus (Hedrick et al., 1997, p. 48), (Hedrick et Uchanski, 2001, p. 5), (Nouvelière, 2002, p.109) :

$$\left(m + \frac{\overline{J_{wr}} + J_{wf}}{r_w^2}\right)a = \frac{T_e/R_g - T_b}{r_w} - (F_{\text{drag}} + F_{\text{aero}} + F_{\text{grav}}), \quad (4.14)$$

où F_{drag} représente la force résistante de roulement et $\overline{J_{wr}}$ comprend les moments d'inertie du moteur (après transmission) et le moment d'inertie des roues arrières :

$$\begin{cases} F_{\text{drag}} &= m \cdot g \cdot C_r \cdot \cos(\theta) \\ \overline{J_{wr}} &= \frac{J_e}{R_g^2} + J_{wr} \end{cases}. \quad (4.15)$$

En introduisant le moment d'inertie totale sur l'axe moteur J_{eq} , on obtient la nouvelle équation longitudinale appliquée aux bus :

$$\frac{J_{eq}}{r_w} \cdot a = T_e - R_g \cdot (T_b + r_w \cdot (F_{\text{drag}} + F_{\text{aero}} + F_{\text{grav}})), \quad (4.16)$$

$$\text{avec } J_{eq} = [J_e + R_g^2 \cdot (m \cdot r_w^2 + J_{wr})]/R_g. \quad (4.17)$$

On en déduit l'expression du couple de freinage global du véhicule :

$$T_b = \frac{T_e - a \cdot J_{eq}/r_w}{R_g} - r_w \cdot (F_{\text{drag}} + F_{\text{aero}} + F_{\text{grav}}). \quad (4.18)$$

L'équation 4.18 met en évidence une relation entre le couple de freinage T_b et l'accélération longitudinale a . Expérimentalement, nous avons pu observer une forte corrélation entre ces deux variables. Le couple T_b représente une des clés de notre stratégie de détection (décrite en section 4.3).

En terme de bilan de moments ramenés sur l'axe moteur et à partir de l'équation 4.16, on obtient :

$$J_{eq} \cdot \frac{d}{dt}\omega_w = F_e - R_g \cdot (F_b + F_{\text{drag}} + F_{\text{aero}} + F_{\text{grav}}), \quad (4.19)$$

où les forces F_e et F_b désignent respectivement la force du moteur et celle de freinage. La figure 4.3 représente le bilan des forces exercées sur un autobus dans une modélisation longitudinale.

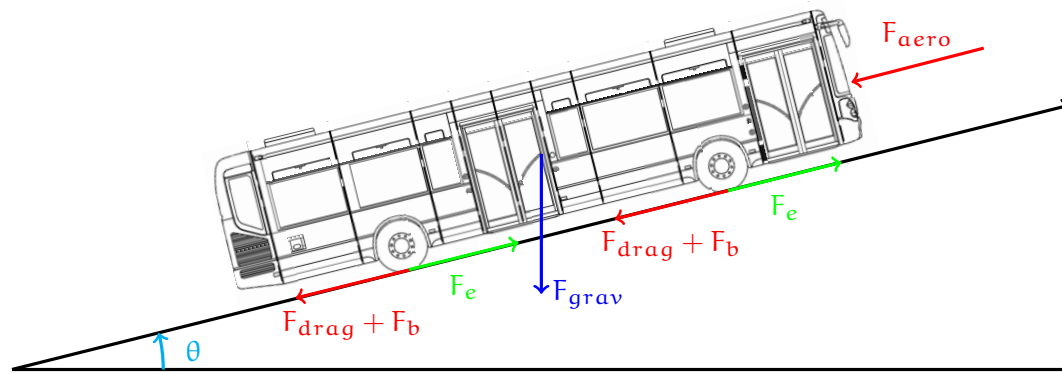


FIGURE 4.3 – Diagramme du corps libre appliqué à un bus.
Bilan des forces exercées sur un autobus (standard)
dans une modélisation longitudinale du véhicule.

Enfin, on détaille comment déterminer le rapport de démultiplication totale R_g et le moment d'inertie équivalent J_{eq} . D'un point de vue pratique, ces deux quantités sont calculées à partir de données CAN spécifiées dans le J1939 (voir sous-section 1.3.3).

Rapport de démultiplication totale R_g : ce rapport est le ratio de démultiplication du moteur jusqu'aux roues $R_g = \frac{\omega_w}{\omega_e}$ (seconde hypothèse, équation 4.11). À partir de la commande d'appui pédale (accélérateur/frein), une suite de démultiplications permet de transmettre un couple et une vitesse de rotation appropriés jusqu'aux roues du véhicule. La chaîne de traction (moteur, convertisseur de couple, boîte de vitesse et pont) fournit la puissance motrice au véhicule via des arbres de transmission en rotation. Song et Hedrick (2004) décrivent R_g comme le produit du *gear ratio in transmission* et du *final gear ratio* ; il s'agit respectivement du ratio de la boîte de vitesse et du ratio de pont. L'équation 4.20 met en relation ces deux ratios pour l'estimation du rapport R_g , illustré dans la figure 4.4 (bas).

$$R_g = \frac{1}{R_t \cdot R_p} . \quad (4.20)$$

Cette relation a été validée expérimentalement à partir des hypothèses trois et quatre, respectivement équations 4.11 et 4.12. L'estimation du rapport R_g , via ces deux équations, est illustré en figure 4.4 (haut). Cette estimation est grossière dans les petites vitesses de boîte à cause de l'embrayage et du convertisseur de couple.

Le ratio de la boîte de vitesse R_t est identifié par le SPN526 : *Transmission Actual Gear Ratio*, disponible sur le CAN J1939 via l'interface bus-FMS. Les ratios obtenus sont cohérents avec les valeurs théoriques préconisées par le constructeur dans le tableau 4.5 :

Type de boîte de vitesse	Type de bus de l'étude	Rapport de boîte				
		1	2	3	4	5
HP 502	Standard	3.43	2.01	1.42	1.00	0.83
HP 504	Articulé					

TABLE 4.5 – Ratios théoriques d'une boîte de vitesse ZF en fonction du rapport engagé.

D'autre part, le rapport de pont R_p (pont arrière, car propulsion) comprend le différentiel et des réducteurs intermédiaires. Le réducteur de la mécanique centrale est 21×67 et les réducteurs intermédiaires du pont dépendent du type de bus considéré : 20×36 pour un bus standard et 18×35 pour un articulé. On obtient donc les ratios de pont pour les bus Citelis :

$$\text{Bus standard : } R_p = \frac{67}{21} \times \frac{36}{20} . \quad (4.21)$$

$$\text{Bus articulé : } R_p = \frac{67}{21} \times \frac{35}{18} . \quad (4.22)$$

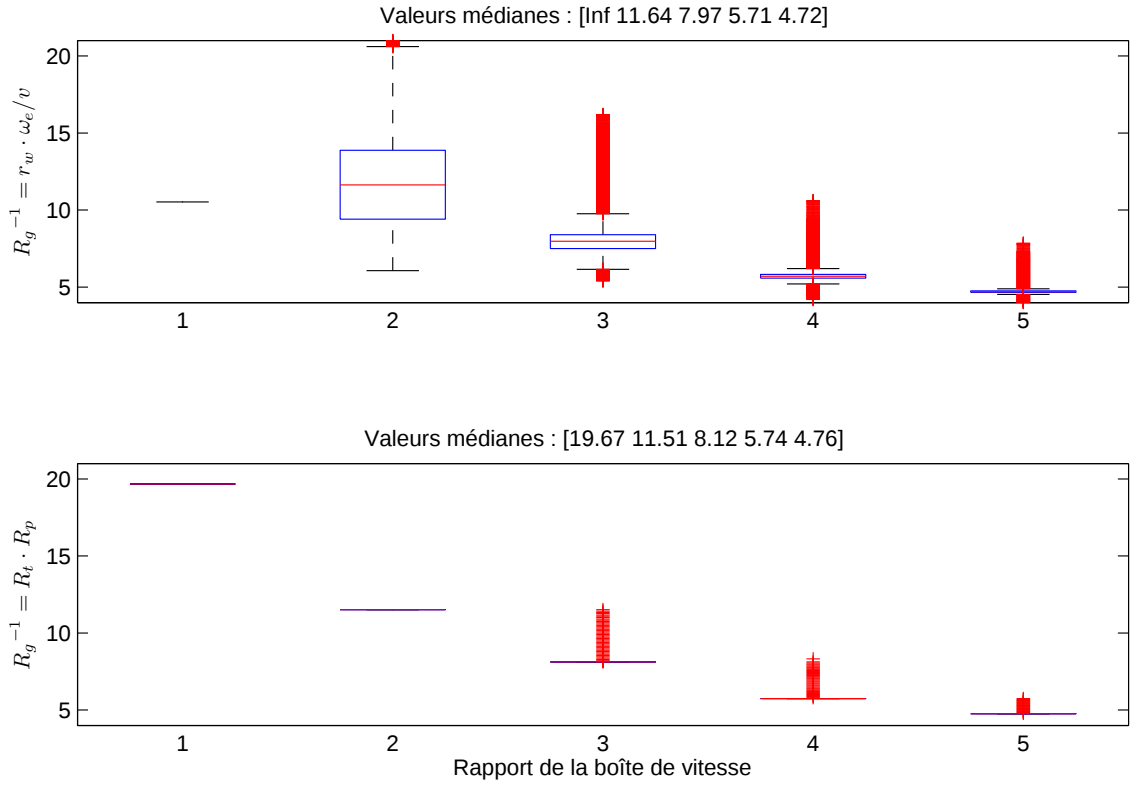


FIGURE 4.4 – Validation expérimentale du calcul du rapport de démultiplication totale, sur un échantillon de données extraites d’un bus standard en novembre 2012. Les rapports R_g calculés à partir des équations 4.11 et 4.12 (haut), et à partir de l’équation 4.20 (bas) sont analogues.

Moment d’inertie équivalent J_{eq} : ce moment d’inertie est l’inertie totale (au niveau véhicule) ramenée sur l’axe moteur. Sa valeur dépend du rapport de boîte engagé :

$$J_{eq} = \frac{J_e + R_g^2 \cdot (m \cdot r_w^2 + J_w)}{R_g} \quad (4.23)$$

où J_w est le moment d’inertie des roues (somme des moments d’inertie des roues avant, arrière et milieu dans le cas d’un bus articulé). Ces moments sont tous de la forme *masse de l’objet* $\times$ (*dimension caractéristique*)². Il est souvent plus pratique de caractériser les objets par leur rayon de giration k et dans le cas d’une roue de bus, on prend un rayon de giration $k = \frac{1}{\sqrt{2}} \cdot r_w$. D’autre part, on considère que les roues sont ramenées à une seule roue par essieu, et le bus est standard (2 essieux) ou articulé (3 essieux) ; on obtient donc un moment $J_w = m_w \cdot r_w^2$ pour les bus standards et $J_w = \frac{3}{2} \cdot m_w \cdot r_w^2$ pour les articulés. On parle alors d’un « modèle bicyclette » (Nouvelière, 2002, p.105). Ceci étant, le moment d’inertie des roues n’est pas significatif devant le produit $m \cdot r_w^2$ dans le calcul de J_{eq} .

4.2.2 Estimation de la pente

La pente θ est une caractéristique importante pour plusieurs problématiques liées au transport : minimisation de la consommation en carburant d'un poids lourd (Plasat, 2005), modélisation dynamique longitudinale d'un bus (Hedrick et al., 1997; Nouvelle, 2002), etc. La pente de la trajectoire suivie peut être mesurée par une instrumentation additionnelle comprenant par exemple une centrale inertielle munie d'un GPS. Il convient alors de modéliser le signal physiquement afin de le débruiter, d'éliminer les perturbations liées aux suspensions du véhicule et aux vibrations de la chaussée (Larsson, 2010). En outre, l'altitude mesurée par un GPS classique n'est pas une donnée fiable.

Dans le cadre de nos travaux, une approche à base de données a été proposée pour l'estimation des pentes d'un parcours donné. Son principal avantage est qu'elle ne nécessite pas l'installation d'équipements lourds. Les relevés de position (longitude, latitude) d'un simple GPS embarqué sont comparés à une base de référence du parcours dont chaque position est associée à une altitude et une pente. Une pente est estimée en lui affectant celle du point le plus proche au sens de la distance orthodromique (équation 4.26), évitant ainsi de construire une modélisation complexe du véhicule ; il est possible d'affiner hors ligne la résolution spatiale et le filtrage des signaux de la base de référence. Après mise en concurrence des données spatiales fournies par l'Institut Géographique National¹ (résolution de 50m×50m×1m) et Google² (résolution non précisée), nous avons choisi de construire notre base de référence à partir des données Google. Les pentes entre deux points x et y sont calculées à partir de leurs trois coordonnées spatiales :

$$\text{slope}_{\%}(x, y) = 100 \cdot \sin \left(\frac{\text{slope}_{\text{deg}}(x, y)}{180} \cdot \pi \right), \quad (4.24)$$

$$\text{où } \text{slope}_{\text{deg}}(x, y) = \arctan \left(\frac{\text{alti}(y) - \text{alti}(x)}{d_{\text{ortho}}(x, y)} \cdot \pi \right). \quad (4.25)$$

On utilise la distance orthodromique – distance à la surface du globe terrestre – entre les deux points x et y :

$$d_{\text{ortho}}(x, y) = R \cdot \arccos \left[\cos \left(\frac{x_2 \cdot \pi}{180} \right) \cdot \cos \left(\frac{y_2 \cdot \pi}{180} \right) \cdot \cos \left(\frac{y_1 - x_1}{180} \cdot \pi \right) + \sin \left(\frac{x_2 \cdot \pi}{180} \right) \cdot \sin \left(\frac{y_2 \cdot \pi}{180} \right) \right], \quad (4.26)$$

où R est le rayon de la terre, et les deux coordonnées représentent la longitude et la latitude en degrés décimaux WGS84 (système géodésique de référence pour le GPS).

La figure 4.5 illustre le parcours étudié dans le cadre de ces travaux de thèse : la ligne J (16 km) du réseau de la STRAV et quelques déviations. On remarque que le dénivelé atteint plus de 60m avec des pentes à plus de 7%. La base de référence a été construite en lissant à la fois les altitudes (filtre moyen) et les pentes (filtre médian).

1. IGN, espace professionnel : <http://professionnels.ign.fr/>.

2. The Google Elevation API : <https://developers.google.com/maps/documentation/elevation/>.

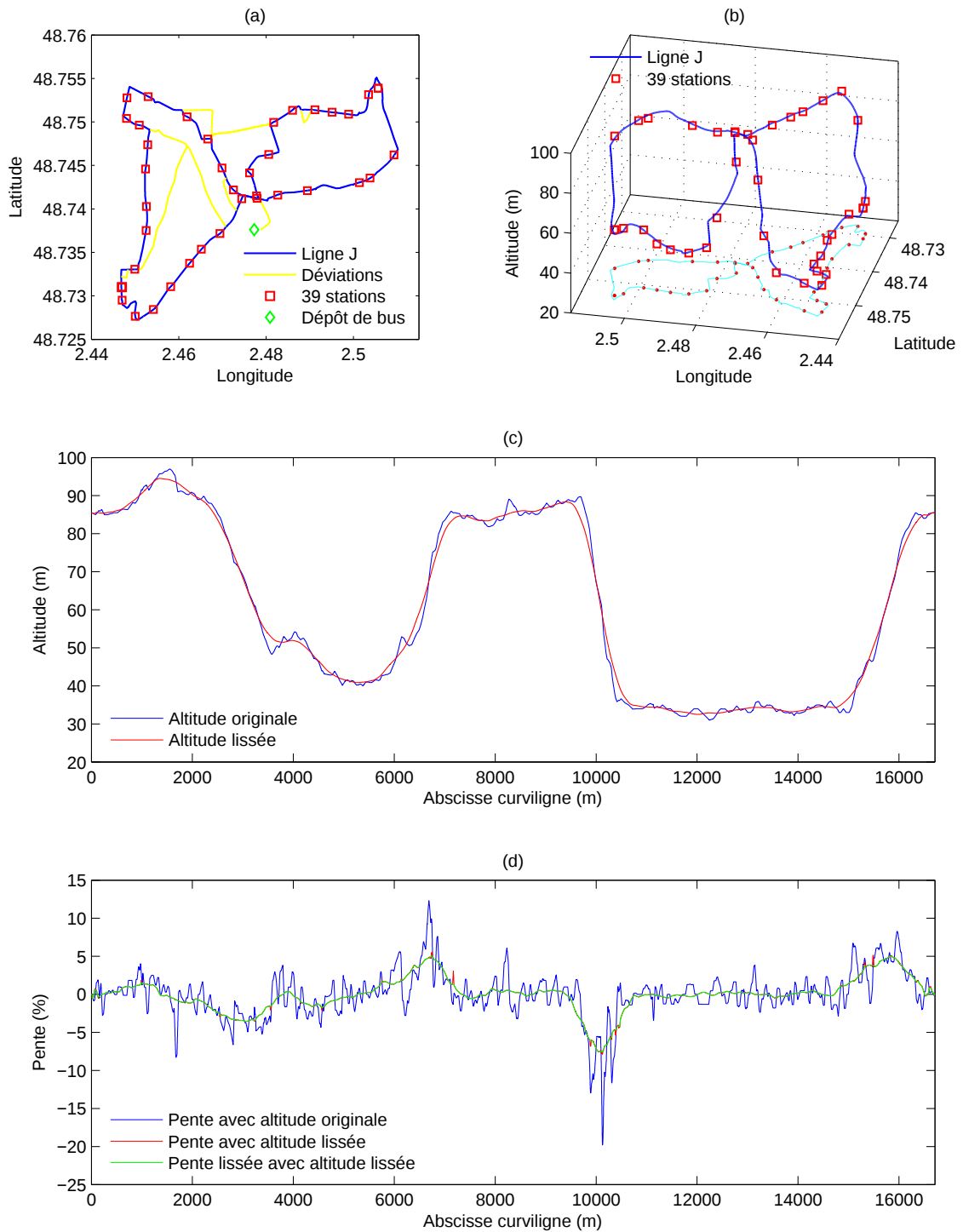


FIGURE 4.5 – Modélisation spatiale de la ligne J et estimation de la pente.

Parcours de la ligne J (16 km) et ses stations en 2D (a) – en 3D (b).

Altitude de la ligne J et lissage moyen de l'altitude par fenêtres glissantes (c).

Pente de la ligne J et lissage médian de la pente par fenêtres glissantes (d).

4.2.3 Représentation du fonctionnement normal du système de freinage

Il est possible de surveiller le système de freinage en s'appuyant sur la modélisation dynamique introduite dans la sous-section 4.2.1. Pour rappel, les couples générés par les freins et le moteur affectent directement la dynamique du véhicule à travers une interaction entre la route et les roues (équations 4.14 et 4.18). On cherche à suivre la relation entre la pression délivrée dans les cylindres et le couple de freinage. Il est souvent admis l'hypothèse que le couple de freinage et la pression des cylindres (au niveau des roues) sont linéairement reliés (équation 4.27). Cette relation linéaire est illustrée par la figure 4.6 et validée expérimentalement sur des freins à tambour par [Post et al. \(1975\)](#), [Hedrick et al. \(1997\)](#) et [Vlacic et al. \(2001\)](#) :

$$T_b = K_b \cdot P_b . \quad (4.27)$$

Physiquement, la pression P_b correspond à la force nécessaire pour presser les plaquettes de frein, fixées aux étriers, sur les disques de frein. Le système de contrôle longitudinal d'un véhicule se base souvent sur la variable K_b , appelée « gain de freinage » ou « efficacité du freinage » ([Hedrick et al., 1997](#)). Dans des conditions normales de fonctionnement, la plage de valeurs de K_b reste stable mais peut varier d'un véhicule à l'autre.

Il est important de noter que dans notre cas d'application, la décélération du véhicule atteint rarement des valeurs inférieures à $-3 \text{ m} \cdot \text{s}^{-2}$. A des niveaux de décélération si faibles, la pression de freinage des circuits avant et arrière sont très souvent similaires ([Vlacic et al., 2001](#)). Nous avons donc choisi d'estimer la pression du système de freinage P_b comme la moyenne des pressions avant et arrière ; le choix de cette pression moyenne s'explique également par la construction simplifiée du modèle dynamique et donc une estimation « en moyenne » du couple de freinage.

Remarque 4.1 (Couple du ralentisseur)

En phase de freinage, une partie du couple de freinage est produite par le ralentisseur de la boîte de vitesse (table 4.3). Il aurait donc été significatif de déterminer le gain de freinage exclusivement à partir du couple du système de freinage pneumatique T_{pn} et ainsi supposer la relation linéaire à partir de ce couple ([Song et Hedrick, 2004](#)) :

$$T_{pn} = T_b - T_{re} . \quad (4.28)$$

Toutefois, il n'a pas été permis de collecter la mesure du couple du ralentisseur (de la boîte de vitesse) T_{re} , disponible sur le CAN constructeur. Une amélioration non négligeable de nos résultats serait liée à l'acquisition de cette donnée.

La prochaine section présente la stratégie de détection proposée, qui est basée sur l'hypothèse de relation linéaire entre le couple et la pression de freinage (équation 4.27). On introduira un modèle de régression conçu pour capturer les imprécisions de notre modélisation de la dynamique du système de freinage. La détection sera menée séquentiellement sur deux indicateurs extraits du modèle de régression.

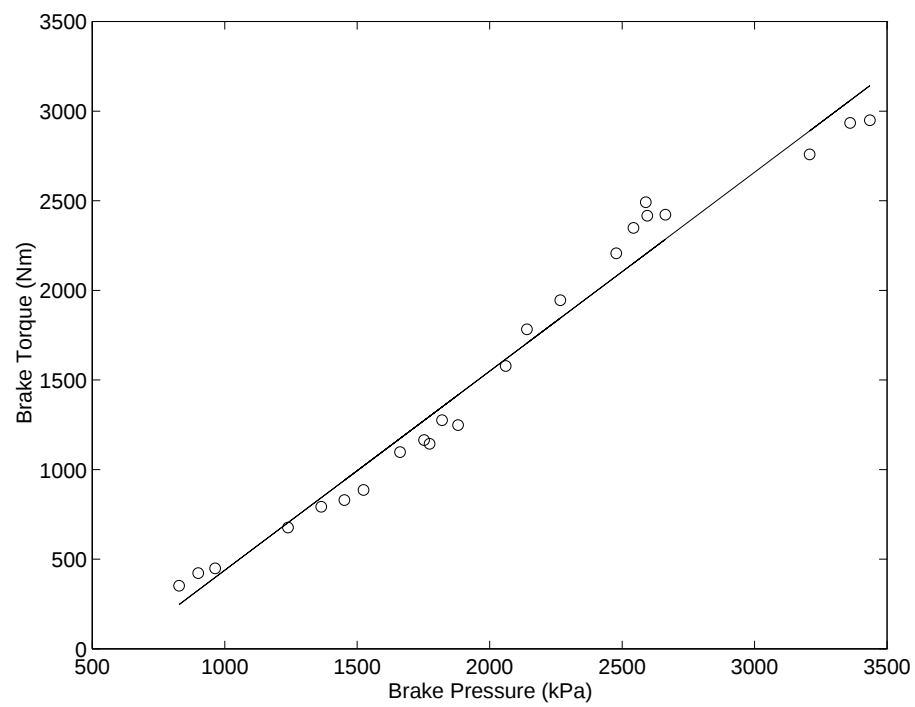


FIGURE 4.6 – Justification expérimentale de la relation linéaire entre pression et couple de freinage ([Hedrick et al., 1997](#)).

4.3 Stratégie séquentielle de détection d'anomalies

4.3.1 Description globale

Afin de capturer les imprécisions de la relation théorique entre couple et pression de freinage (équation 4.27) d'une part, et celles de la modélisation dynamique simplifiée d'autre part, on introduit le modèle de régression linéaire défini par l'équation :

$$T_b = K \cdot P_b + \epsilon, \quad (4.29)$$

où $\epsilon \sim \mathcal{N}(0, \sigma^2)$ est un bruit blanc additif (suivant une loi normale de moyenne nulle). Ce bruit est introduit pour capturer les incertitudes et imprécisions du modèle, et la variabilité potentielle du processus. La figure 4.7 résume ce modèle de régression :

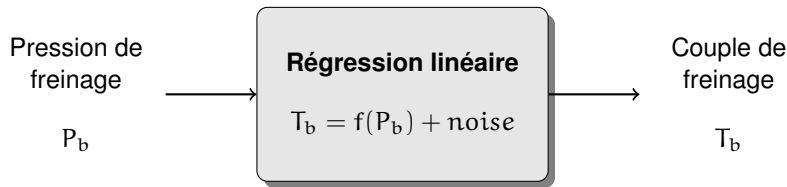


FIGURE 4.7 – Modèle de régression linéaire entre la pression et le couple de freinage.

Pour chaque plage de freinage, on construit un modèle de régression afin de calculer le gain K et l'erreur sur l'hypothèse de linéarité σ . Une variabilité significative de ces variables sera interprétée comme l'émergence d'une anomalie sur le système de freinage. Les couples (K, σ) seront collectés dynamiquement et considérés comme les entrées des méthodes de détection séquentielle décrites dans la sous-section 4.3.2. La figure 4.8 résume notre approche de surveillance globale du système de freinage par une stratégie séquentielle de détection en 3 étapes. Cette stratégie s'appuie à la fois sur la modélisation analytique du véhicule (section 4.2) et une méthode de détection à base de données.

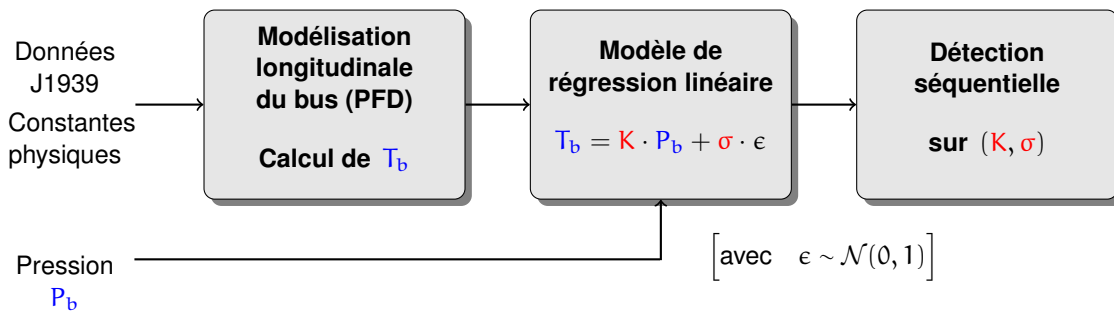


FIGURE 4.8 – Stratégie séquentielle de détection en 3 étapes pour la surveillance du système de freinage.

4.3.2 Détection de défauts

Notre objectif est de détecter tout changement dans le mode de fonctionnement du système de freinage. Dans ce but, nous avons construit une stratégie de surveillance de ce système dont l'ultime étape est la *détection de défauts* (Fault Detection). La détection d'anomalies consiste à trouver des patterns (formes) dans les données qui ne sont pas conformes aux formes nominales (Chandola et al., 2009). Afin d'établir une stratégie complète de diagnostic, il convient de localiser/identifier les formes dites « anormales » après la détection de défauts (voir sous-section 2.1.2).

Dans ces travaux, aucune signature de défaillance n'était disponible en pratique et les données collectées étaient a priori toutes en situation de bon fonctionnement. Ceci nous a conduit à introduire des défauts simulés (dégradation du couple de freinage et fuite d'air) afin d'évaluer nos algorithmes. Il n'a pas été possible de provoquer des défauts réels sur les freins en conditions opérationnelles pour des raisons évidentes de sécurité liées au maintien de l'exploitation des véhicules pendant nos travaux. L'objectif est ici de détecter le plus tôt possible un défaut tout en garantissant un minimum de fausses alarmes.

Les cartes de contrôle représentent des méthodes séquentielles adaptées à notre problématique permettant de suivre temporellement les variables K et σ . Un panel de ces méthodes séquentielles à tests d'hypothèses est dressé dans la section 3.3. Les formulations multivariées de ces cartes présentent l'avantage de prendre en compte les corrélations et dépendances entre variables. En outre, ces cartes se présentent sous la forme d'outils graphiques facilement interprétables et implémentables. Nous avons donc choisi de suivre l'état de fonctionnement du système de freinage à partir des cartes multivariées de Shewhart et de CUSUM, présentées dans la sous-section 3.3.1. Pour rappel, on se place dans le cadre de tests de détection en-ligne à deux hypothèses $\mathcal{H}_0$ (bon fonctionnement) et $\mathcal{H}_1$ (anomalie au temps t_c) – voir équation 3.6. Ces méthodes supposent l'hypothèse de normalité des distributions sous-jacentes. Nous assumons donc que le système de freinage peut être séquentiellement décrit par une variable aléatoire gaussienne de moyenne μ et de variance Σ sur les couples (K, σ) .

En pratique, les défauts qui seront introduits en sous-section 4.4.3, simulent un décentrage des distributions des couples (K, σ) . Les cartes de contrôle multivariées nous serviront donc à détecter les déviations de la moyenne de la loi supposée générer les couples. La performance des détecteurs sera mesurée par la courbe ROC et notamment par l'aire sous cette courbe, communément appelée AUC (voir sous-section 3.2.4). L'hypothèse de normalité n'étant pas tout à fait respectée (cf. figure 4.11) par les observations réelles (K, σ) , les seuils « théoriques » ne sont pas adaptés à nos détecteurs. Nous avons donc proposé d'estimer de nouvelles limites de contrôle basées sur un compromis entre taux de fausses alarmes et taux de bonne détection à partir de la courbe ROC. Ainsi, on pourra fixer un seuil de détection mieux adapté aux formes des couples (K, σ) .

4.4 Expérimentations sur données réelles

4.4.1 Acquisition des données de fonctionnement

Les données de fonctionnement du système de freinage sont collectées par un calculateur ATON, embarqué dans chaque autobus de la flotte étudiée (voir sous-section 1.3.2). Ces données sont échantillonnées à une fréquence d'environ 100 Hz et stockées dans des fichiers binaires compressés, initialisés au démarrage des véhicules. Dans le cadre de nos travaux, un outil de conversion, intitulé CANversionTool, a été développé en Java et MATLAB compilé afin de décoder facilement les fichiers enregistrés sur l'ATON ; de nombreuses configurations de données CAN et formats de sortie peuvent être sélectionnés. L'interface homme-machine (IHM) de cet outil est illustrée dans la figure 4.9, pendant la conversion de fichiers GZIP en fichiers MAT.

La figure 4.10 résume l'agrégation des données CAN pour la construction du modèle de régression, pierre angulaire de notre stratégie de surveillance avec a priori physique. D'une part, le couple de freinage T_b est déterminé à partir de neuf données CAN spécifiées dans le J1939 (SAE Technical Standards, 2006) et six constantes physiques (équation 4.18). Les variables intermédiaires sont encadrées par des losanges. D'autre part, la pression de freinage P_b est estimée en fonction de pressions mesurées par des capteurs additionnels via le CAN Intellibus (CAN constructeur propriétaire dédié au télédiagnostic de véhicules Irisbus).

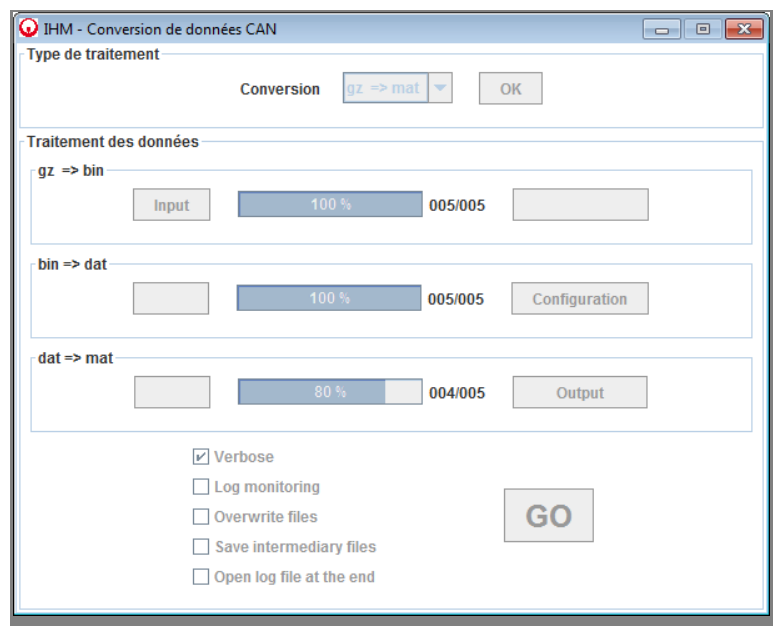


FIGURE 4.9 – CANversionTool : outil de conversion de données CAN adapté aux fichiers produits par le calculateur embarqué ATON.

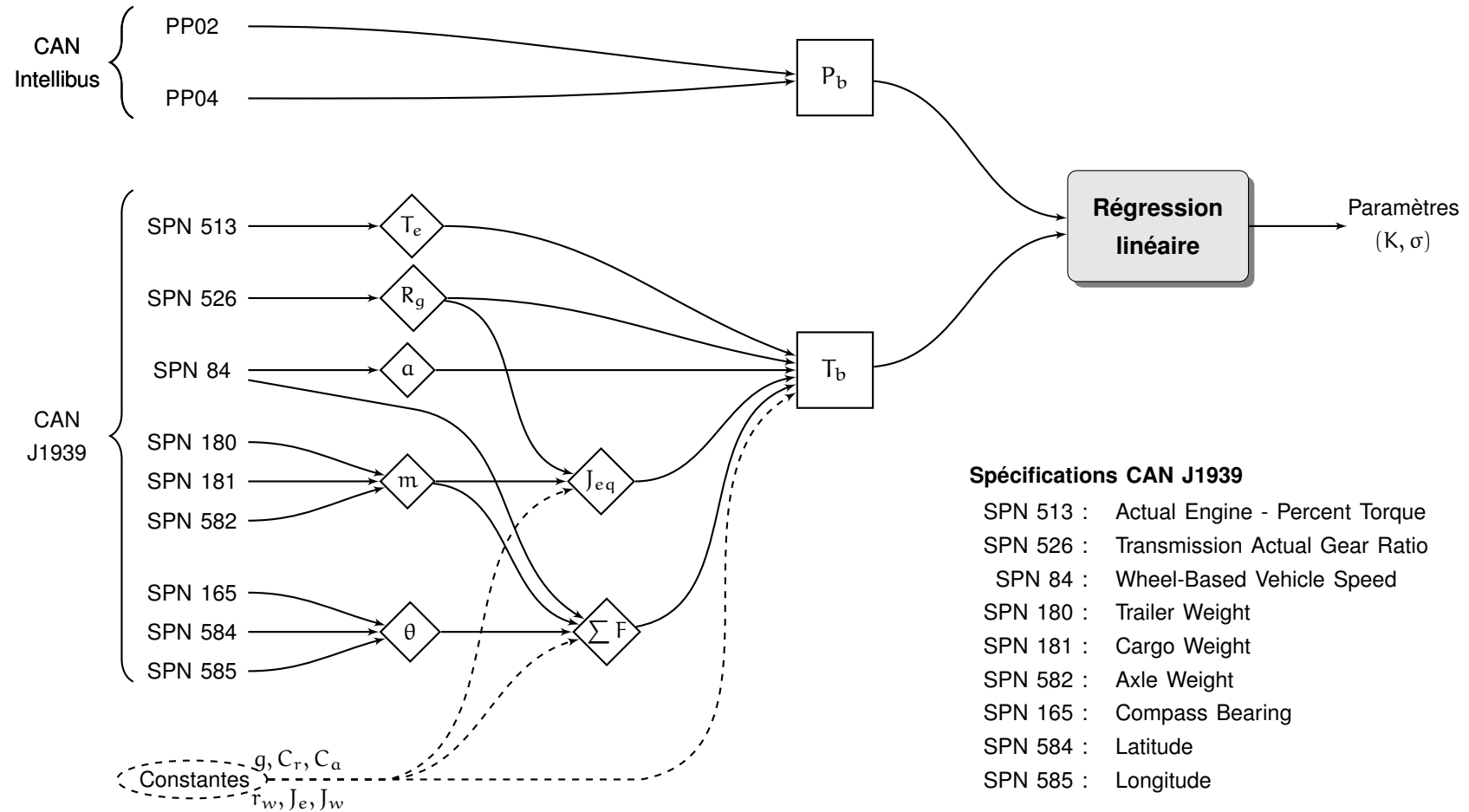


FIGURE 4.10 – Extraction du couple de paramètres (K, σ) à partir des données CAN collectées sur l'ATON.
 En phase de freinage, on détermine le couple de freinage T_b à partir de neuf données du CAN J1939 et de constantes (équation 4.18).
 La pression de freinage P_b est estimée en fonction des pressions mesurées par deux capteurs additionnels PP02 et PP04 (figure 4.1).

4.4.2 Description des données réelles collectées

Dans le cadre du UseCase de Brunoy, nous avons pu collecter l'ensemble des données préconisées - données CAN listées dans la figure 4.10 - sur quatre autobus en exploitation. La complexité de l'architecture liée à la collecte des données, décrite dans la figure 1.4, explique l'acquisition tardive des données sur une relativement courte période. Toutefois, nous avons constitué une base de **plages de freinage significatives** lorsqu'elles respectaient les conditions suivantes :

1. les plages de freinage sont sélectionnées en fonction d'une donnée binaire activée à l'appui de la pédale de frein, le SPN 597 - Brake Switch (voir figure 4.2) ;
2. pour chaque plage de freinage, le véhicule doit être localisé, au moins une fois, à moins de 15 mètres de la ligne J ou de ses déviations ;
3. pour chaque plage de freinage, on extrait les variables d'intérêt lorsque la vitesse du véhicule est supérieure à 2 km.h^{-1} et l'étendue est supérieure à 5 km.h^{-1} .

Cette procédure est utilisée pour extraire une séquence de couples (K, σ) qui représente le fonctionnement dynamique du système de freinage et chacun de ces couples est estimé à partir d'un modèle de régression linéaire. Le tableau 4.6 résume la base des données collectées (plages de freinage) sur les quatre véhicules en exploitation :

Bus	Type	Date de début		Date de fin		Plages de freinage
484	standard	2012/09/26	07-23-17	2012/11/06	01-28-02	5547
486	articulé	2012/09/26	13-10-46	2012/10/12	01-04-10	5060
488	articulé	2012/09/28	06-14-53	2012/11/04	22-30-12	3087
490	articulé	2012/09/28	15-35-51	2012/11/05	10-10-23	3299

TABLE 4.6 – Collection des plages de freinage pour chaque bus.

La figure 4.11(a) représente quelques plages de freinage du bus 484 en exploitation sur la ligne J et ses déviations (vert), du 26 septembre au 6 novembre 2012. On observe que certaines zones du parcours de référence sont rarement opérées en phase de freinage par l'autobus. Il s'agit de portions de routes rectilignes ou de déviations non exploitées par le véhicule. D'autre part, on observe que le bus circule souvent sur d'autres routes (rouge) que celles de la ligne J ce qui se traduit par un grand nombre de plages de freinage localisées en dehors du parcours de référence (bleu).

Les distributions des couples (K, σ) des bus 484, 486, 488 et 490 sont illustrées dans la figure 4.11(b). On observe que l'ordre de magnitude des couples diffère selon le type d'autobus : standard (12 m) ou articulé (18 m). Avec environ une tonne par mètre, un standard pèse 12 t et un articulé 18 t. Ceci explique des valeurs de gain K plus élevées (d'un facteur 3/2) pour les bus articulés par rapport à la distribution des couples du bus 484 (standard). On observe que la répartition des couples (K, σ) est similaire pour les bus

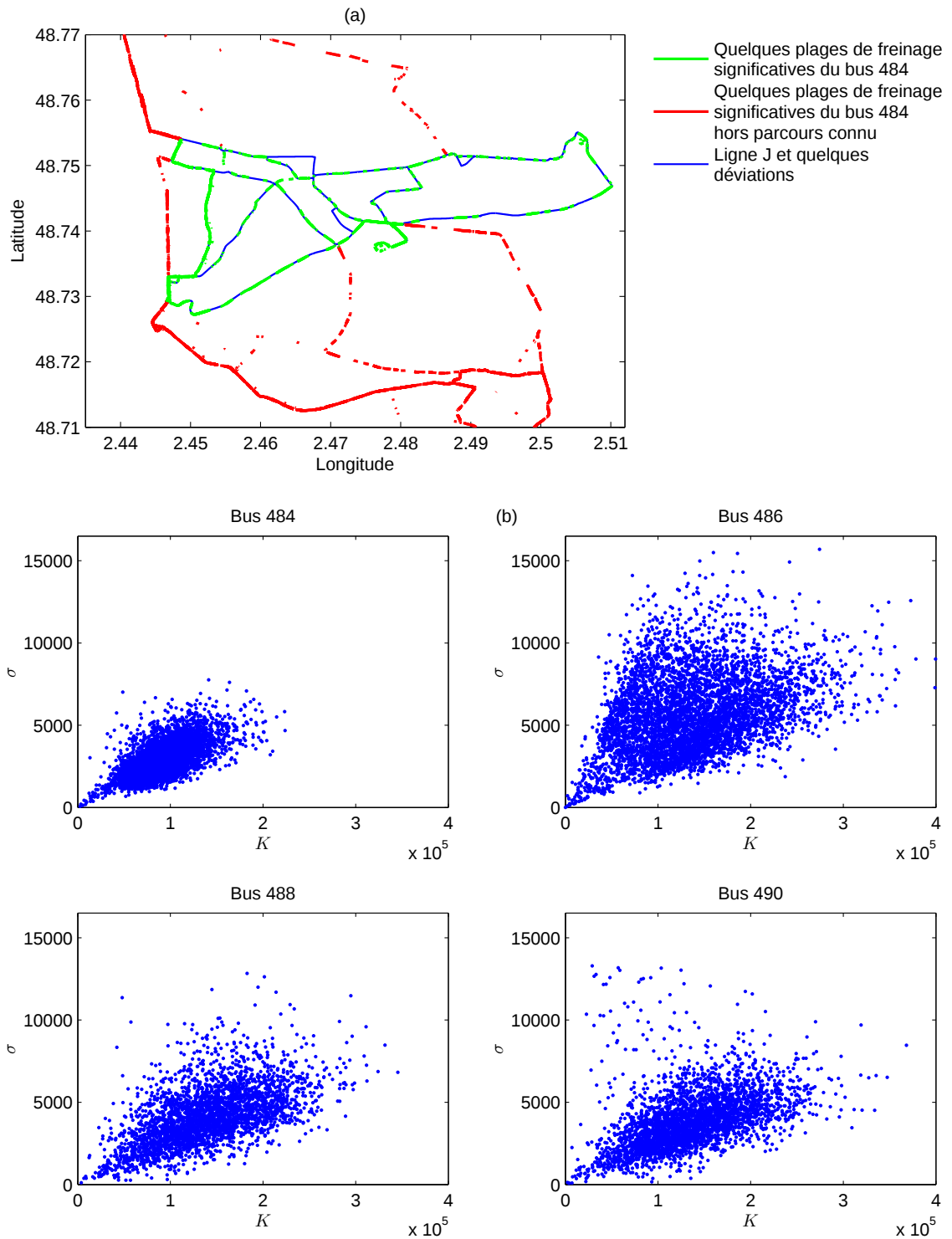


FIGURE 4.11 – Représentation spatiale de quelques plages de freinage significatives et distributions des couples (K, σ) représentant le fonctionnement des véhicules.

Plages de freinage du bus 484 sur le parcours de référence, pendant la période du 26 septembre au 6 novembre 2012 (a).
Distributions des couples (K, σ) pour les bus 484, 486, 488 et 490 dont le système de freinage est supposé en bon fonctionnement (b).

articulés 488 et 490. En revanche, la distribution des couples du bus 486 présente une erreur σ sur le modèle de régression globalement plus élevée que celle des autres bus. On attribue cet écart à des erreurs d'acquisition des données d'entrée et à la modélisation longitudinale car aucune dégradation liée au freinage n'a été signalée pendant cette période. Par conséquent, on choisit par la suite de travailler sur les couples (K, σ) du bus 484 qui sont en plus grand nombre.

Enfin, les distributions des couples semblent s'écarter légèrement d'une loi gaussienne (écart plus prononcé dans le cas du bus 486) - on observe notamment l'asymétrie des distributions. Toutefois, l'hypothèse de normalité sera assumée dans le cas de fenêtres locales (en moyenne) pour les tâches de détection avec les cartes de contrôle multivariées. L'étude sera menée dans les prochaines sous-sections.

Parmi la dizaine de données CAN collectées (voir figure 4.10), on décrit succinctement l'acquisition de certaines mesures :

- les *pressions de freinage* sont obtenues via des capteurs à résistance variable et correspondent aux pressions (0 – 16 bar) mesurées en sortie du pédalier de frein (voir figure 4.1). Ces capteurs additionnels ont été spécifiquement installés pour l'étude et sont localisés sur la platine pneumatique à l'avant du véhicule, à proximité du coffre des batteries ;
- la *charge « nette »* du véhicule est déduite de la somme pressions de suspension (relation linéaire entre masse et pression de suspension) mesurées sur chaque essieu ; des capteurs de pression spécifiques ont été installés pour les besoins du projet ;
- la *vitesse du véhicule* est mesurée via un capteur à réluctance variable : sur un disque solidaire du moyeu, un ensemble d'aimants est disposé sur la périphérie et une bobine est placée au voisinage immédiat. Ainsi, il suffit de compter le nombre d'impulsions créées par le passage d'un aimant devant la bobine par unité de temps. L'accélération est obtenue par dérivation de la vitesse en fonction du temps conformément à la seconde hypothèse de notre modélisation longitudinale (équation 4.10) ;
- les *coordonnées GPS* (position et gisement) sont enregistrées à partir d'un GPS intégré au second ordinateur embarqué, le VIDAC. Ces données sont transmises dans un format spécifié par le J1939 via un CAN reliant ce ordinateur et l'ATON (voir figure 1.4) ;
- les autres données (couple moteur et ratios de boîte) sont mesurées à partir de capteurs existants et transmises par le CAN constructeur via l'interface bus-FMS.

L'ensemble de ces données CAN sont horodatées et collectées par l'ATON sur des clés USB embarquées. Ces clés sont déchargées manuellement lors de visites régulières au dépôt de bus.

4.4.3 Dégradations simulées du système de freinage

Les systèmes de freinage pneumatique sont généralement employés dans les véhicules commerciaux comme des autobus ou des camions (Limpert, 2011) et la présence de défauts de freinage augmente le risque non négligeable d'accidents. Des défauts typiques du système de freinage concernent l'usure des garnitures de freins, les dérèglages des tiges de poussée en sortie des cylindres et des fuites dans le système pneumatique (Subramanian, 2006). Des approches à base de modèle physique ont été proposées par Karthikeyan et al. (2011) et Sonawane et al. (2011) afin de repérer ces défauts sur des freins à tambour avec cames en S (commandées par l'arbre à came, voir la figure 1.7(a)), système de freinage le plus fréquent dans les véhicules commerciaux en Inde.

Par ailleurs, Hedrick et al. (2001) ont proposé différentes catégorisations de dégradation aux conférences PATH³ en 2001 et 2002. Les auteurs présentent notamment deux types de dégradations aisément simulables et définies comme :

- *Défaut simple* : dégradation de 15% du couple de freinage

$$\widetilde{T}_b = (1 - 0.15) \cdot K_b \cdot P_b \quad (4.30)$$

- *Défaut multiple* : dégradation de 15% du couple et fuite de 10% de la pression de freinage dans le circuit pneumatique

$$\begin{cases} \widetilde{P}_b = (1 + 0.10) \cdot P_b \\ \widetilde{T}_b = (1 - 0.15) \cdot K_b \cdot \widetilde{P}_b \end{cases} \quad (4.31)$$

Afin d'évaluer notre stratégie de surveillance séquentielle, nous avons donc opté pour des dégradations plus ou moins abruptes du système de freinage à travers cinq situations en combinant les deux types de défauts proposés par Hedrick et al. (2001). L'équation 4.32 décrit une perte d'efficacité des freins, par exemple un dérèglage des tiges de poussée ou une usure des garnitures,... et l'équation 4.33 modélise une fuite dans le système pneumatique :

$$\begin{cases} \widetilde{T}_b = (1 + \gamma_1) \cdot T_b, \end{cases} \quad (4.32)$$

$$\begin{cases} \widetilde{P}_b = (1 + \gamma_2) \cdot P_b, \end{cases} \quad (4.33)$$

où (γ_1, γ_2) est une instance de défaut et, on a supposé l'égalité $T_b = K_b \cdot P_b$ (équation 4.27). On simule cinq situations de dégradation résumées par le tableau ci-après et la figure 4.12, où le processus est supposé stationnaire avant et après le changement du mode de fonctionnement du système. On distingue notamment les 1^{ère} et 5^{ème} situations qui représentent un changement abrupt de forte amplitude avec saut brutal du signal.

Situation	1	2	3	4	5
γ_1	-0.50	-0.25	-0.15	0	0
γ_2	0	0	0.15	0.25	0.50

TABLE 4.7 – Cinq situations de dégradations simulées pour le système de freinage.

3. Partners for Advanced Transportation Technology (PATH) : <http://www.path.berkeley.edu/>.

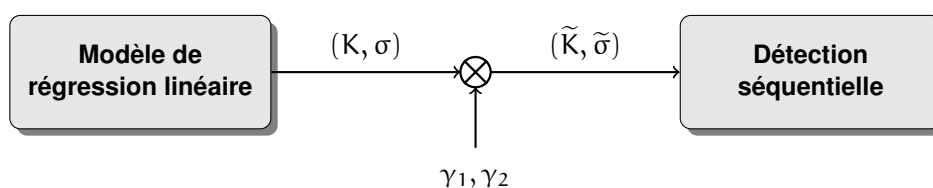


FIGURE 4.12 – Processus de dégradation du système de freinage.

Pour notre étude, les plages de bon et mauvais fonctionnement ont été réparties différemment parmi les plages de freinage relevées en exploitation pour chaque véhicule. La table 4.8 résume la base ainsi constituée :

Autobus	Nombre de plages de freinage en	
	bon fonctionnement	mauvais fonctionnement
484	3983	1564
486	3116	1944
488	2232	855
490	2121	1178

TABLE 4.8 – Répartition des plages de bon fonctionnement et celles dégradées.

Les bases ainsi constituées contiennent des couples (K, σ) identifiés à partir des enregistrements réels et utilisés tels quels ou bien modifiés grâce aux équations 4.32 et 4.33 pour simuler des dégradations réalistes et interprétables. Les tailles ont été fixées en fonction des données disponibles et non afin de favoriser le comportement d'une méthode de détection en particulier. D'une part, nous avons cherché à partitionner les bases obtenues avec environ les deux tiers de signaux en bon fonctionnement et le reste en mauvais fonctionnement. D'autre part, 50% des signaux en bon fonctionnement sont utilisés pour l'apprentissage des paramètres de la distribution avant changement. Plus le nombre de données disponibles est grand, meilleure est l'inférence des estimateurs. Il est à noter que nous avons proposé des situations de dégradation mono-défaut, avec un seul changement dans chaque séquence de plages de freinage. La simulation de plusieurs défauts successifs est envisageable avec notre méthode de dégradation et constitue une extension directe de notre stratégie de surveillance. Dans une situation multi-défauts, le détecteur devrait être mis à jour après chaque détection d'un défaut, à partir d'un nouvel échantillon de plages de freinage mesurées sur le système (supposé en bon fonctionnement). Une base de données réelles multi-défauts fera l'objet d'une analyse dans le chapitre 5, appliqué à l'étude des signaux de portes. L'évaluation d'un détecteur à deux hypothèses n'est pas évidente sur un jeu de données avec deux changements ou plus. C'est pourquoi notre stratégie de détection est évaluée sur des situations de dégradation mono-défaut.

La figure 4.13 illustre les cinq situations de dégradation simulée pour le bus 484 dans l'espace des variables (K, σ) et en fonction des plages de freinage collectées entre le 26 septembre 2012 et le 6 novembre 2012. On observe différentes directions de la dérive de fonctionnement après changement. D'une part, une diminution du couple T_b a pour effet

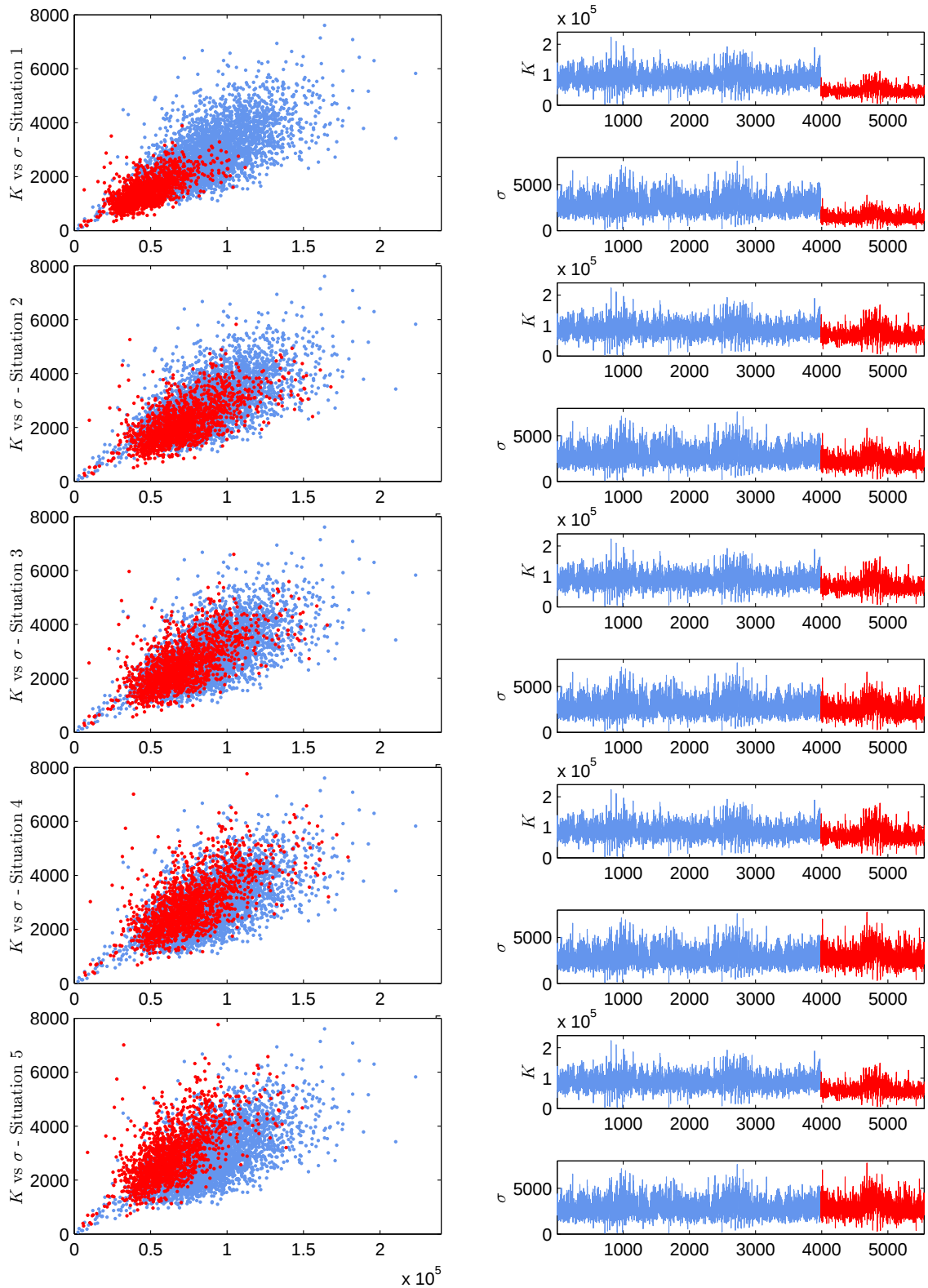


FIGURE 4.13 – Illustration des cinq situations de dégradation du système de freinage (d’une diminution du couple à une fuite de pression) pour le bus 484.

Pour chaque situation, les variables K et σ sont représentées en mode normal (bleu) et en mode dégradé (rouge) dans l’espace des couples (colonne de gauche) et en fonction des plages de freinage (colonne de droite).

de diminuer les valeurs des couples (K, σ) ; le saut est facilement visible dans la première situation. D'autre part, une augmentation de la pression P_b a pour effet de diminuer les gains de freinage mais n'affecte pas les valeurs de l'erreur sur le modèle de régression car l'erreur σ dépend principalement du couple T_b (équation 4.29) ; ce type de changement est aisément visible dans la cinquième situation. Les situations de dégradation 2, 3 et 4 sont plus difficilement détectables.

4.4.4 Résultats expérimentaux et estimation du seuil de détection

Comme décrit dans la sous-section 4.3.2, l'étape de détection de notre stratégie de surveillance est accomplie par une carte de contrôle. Nous avons mis en concurrence les cartes multivariées de Shewhart et de CUSUM, présentées dans la sous-section 3.3.1. Les variables en entrée sont les couples $\{(\widetilde{K}_t, \widetilde{\sigma}_t), t = 1, \dots, T\}$, où T est le nombre de plages de freinage disponibles pour chaque bus étudié (voir tableau 4.6). La population « sous contrôle », utilisée pour l'estimation des deux premiers moments de la loi génératrice, a été fixée à la moitié des plages de freinage en bon fonctionnement. Les deux cartes multivariées prennent en compte la dynamique des systèmes en moyennant les couples (K_t, σ_t) en fenêtres locales contiguës. La taille de ces fenêtres a expérimentalement été fixée à 10. Une taille trop petite augmente la sensibilité du détecteur et une taille trop élevée réduit radicalement le nombre d'échantillons à surveiller. Il est à noter que l'apprentissage des paramètres de la loi génératrice et les tests séquentiels requièrent un faible coût de calcul et peut être traité en temps réel pour une implémentation dans des calculateurs embarqués.

Les seuils de détection des deux cartes multivariées ont été statistiquement fixés, comme décrit dans la sous-section 3.3.1. Celui de la carte de Shewhart est le multiple réel d'un quantile de la loi de Fisher (équation 3.48). Crosier (1988, annexe A) décrit une procédure itérative pour estimer le seuil de détection de la carte du CUSUM. A partir d'un même jeu d'hyperparamètres (taille de fenêtre locale fixée $W = 10$, dimension $d = 2$ et un taux de fausses alarmes espéré $\alpha = 0.005$), les seuils théoriques de Shewhart et de CUSUM sont respectivement fixés à 2.00 et 5.50 avec l'hypothèse que la loi sous-jacente est normale. Sous l'hypothèse $\mathcal{H}_0$, la probabilité α d'observer la statistique de test dépasser la limite UCL_α est donc théoriquement de $1/200$. La figure 4.14 (colonne de gauche) présente les statistiques de test des deux cartes de contrôle et leurs seuils théoriques en fonction des fenêtres locales avec des plages en bon fonctionnement (bleu) et dégradées (rouge) du bus 484.

Présenté dans la sous-section 3.2.4, un critère de performance en détection est l'aire sous la courbe ROC (taux de faux positifs FPR VS taux de vrais positifs TPR), usuellement appelée AUC. La figure 4.14 (colonne de droite) présente les courbes ROC des deux détecteurs et les AUC associées, pour chaque situation de dégradation du bus 484 en faisant varier les seuils de détection. Le tableau 4.9 présente les mesures d'AUC des quatre bus pour les cinq situations de dégradation (équations 4.32 et 4.33). La carte multivariée du

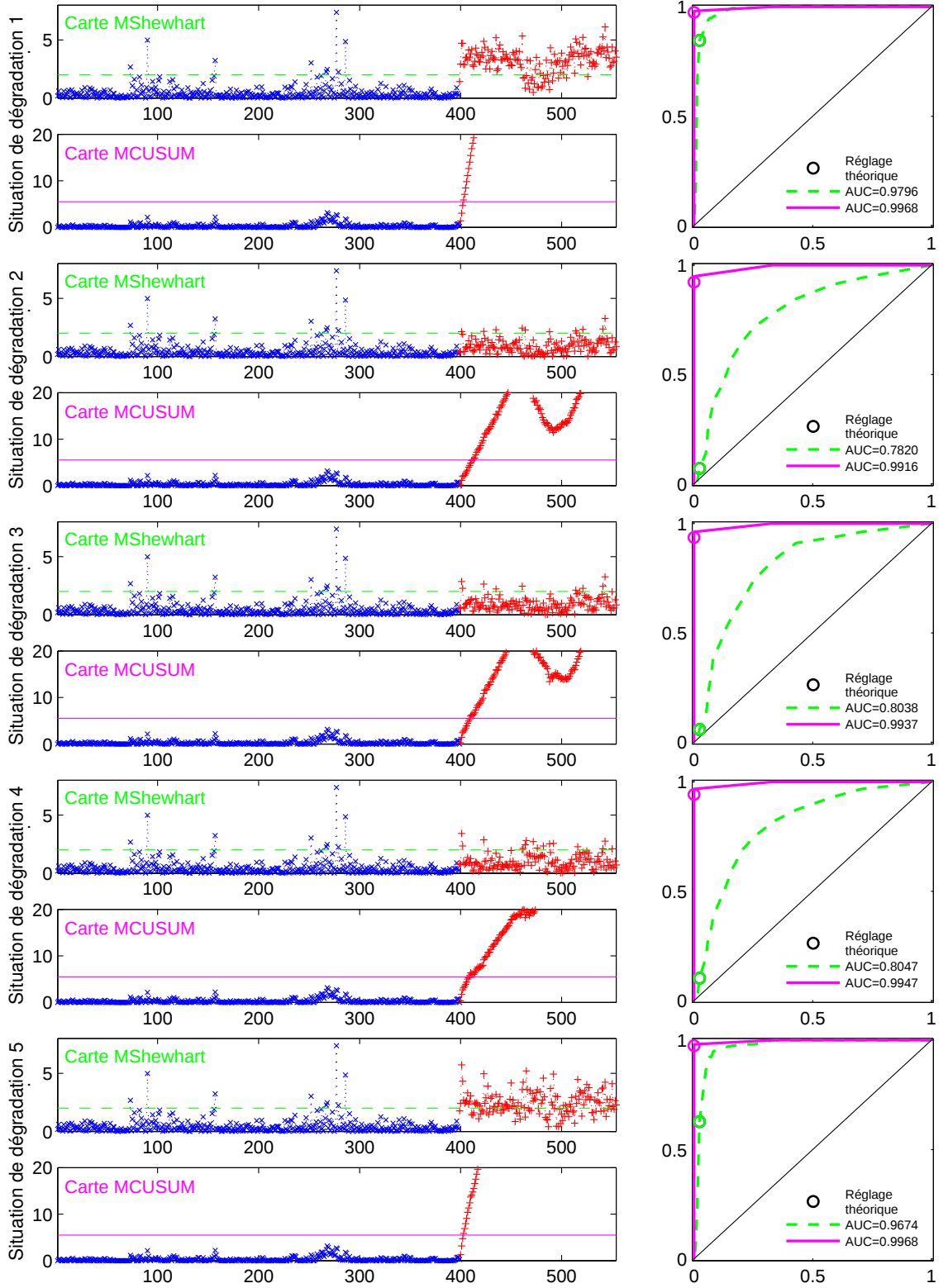


FIGURE 4.14 – Cartes multivariées de Shewhart et de CUSUM, et les courbes ROC (FPR vs TPR) pour les cinq situations de dégradation des plages de freinage du bus 484. Pour chaque situation, les cartes multivariées de **Shewhart** / **CUSUM** sont représentées avec leurs seuils théoriques UCL_{α} en fonction des fenêtres locales (colonne de gauche), et les courbes ROC de chaque détecteur (colonne de droite).

Carte Situation	Shewhart					CUSUM				
	1	2	3	4	5	1	2	3	4	5
bus 484	97.96	78.20	80.38	80.47	96.74	99.68	99.16	99.37	99.47	99.68
bus 486	92.26	63.37	56.43	48.41	66.32	99.56	98.30	97.39	93.31	98.22
bus 488	99.75	78.19	74.03	68.27	92.76	99.59	98.78	98.78	98.38	99.19
bus 490	96.77	58.42	56.25	59.36	87.14	99.41	96.28	94.33	95.31	99.61
moyenne	96.69	69.54	66.78	64.13	85.74	99.56	98.13	97.47	96.62	99.17

TABLE 4.9 – AUC (en %) des cartes de Shewhart (à gauche) et CUSUM (à droite).

Les cinq situations de dégradation sont listées dans le tableau 4.7.

CUSUM est nettement plus performante que celle de Shewart dans tous les cas étudiés. Les cartes de type CUSUM sont moins sensibles aux outliers et une dérive du vecteur moyen du processus est détectée plus rapidement que sur les cartes de Shewhart. On observe sur les courbes ROC du bus 484 que le seuil théorique du CUSUM (cercle magenta) offre un meilleur compromis en terme de détection que celui de Shewhart (cercle vert) ; la détection avec seuil théorique de la carte de Shewhart maintient un faible taux de fausses alarmes au prix d'un faible taux de vrais positifs. Le détecteur parfait est identifiable sur la courbe ROC au point (0, 1) et un détecteur complètement aléatoire est sur la diagonale, représentée par une ligne noire dans la figure 4.14 (colonne de droite). Les performances des deux détecteurs avec seuil théorique sont présentées dans le tableau 4.10 en termes de taux de vrais positifs et de faux positifs. La carte du CUSUM surpasse toujours celle de Shewhart en termes de vrais positifs. Toutefois, on observe qu'une anomalie a été signalée par le CUSUM avant l'introduction de notre dégradation simulée pour le bus 486, ce qui se traduit par un taux de fausses alarmes (FPR) à 20%.

Afin d'améliorer le compromis FPR et TPR, nous proposons de choisir un seuil de détection à partir de la courbe ROC : le détecteur avec seuil « optimisé » minimise la distance (euclidienne) entre la courbe ROC et le point (0, 1). Avec cette approche, nous obtenons une nouvelle limite de contrôle pour chaque détecteur : 0.62 pour la carte de Shewhart et 4.45 pour la carte du CUSUM. Le tableau 4.11 présente les performances avec seuil optimisé en terme de taux de vrais positifs et de faux positifs, sur les mêmes jeux de données. Les performances ont bien entendu été améliorées pour les deux cartes multivariées dans tous les cas étudiés. On note en particulier pour le bus 486, que le taux de bonne détection a été maintenu à 100% et le taux de fausses alarmes a été réduit à 15%.

Shewhart						CUSUM					
1	2	3	4	5		1	2	3	4	5	
Taux de vrais positifs (%)					FPR	Taux de vrais positifs (%)					FPR
84.62	7.05	5.77	10.26	62.82	2.26	97.44	91.67	92.95	93.59	97.44	0.00
77.95	11.79	9.23	11.28	17.44	9.00	100.00	100.00	100.00	100.00	100.00	20.90
71.76	3.53	2.35	1.18	12.94	0.00	95.29	75.29	65.88	42.35	90.59	0.00
14.53	0.00	0.00	0.85	5.13	0.00	92.31	0.00	0.00	0.00	88.03	0.00
62.21	5.59	4.34	5.89	24.58	2.82	96.26	66.74	64.71	58.99	94.01	5.23

TABLE 4.10 – Taux de vrais positifs et de faux positifs (en %) avec seuil théorique.

Shewhart						CUSUM					
1	2	3	4	5		1	2	3	4	5	
Taux de vrais positifs (%)					FPR	Taux de vrais positifs (%)					FPR
98.72	54.49	54.49	55.77	96.79	14.57	98.08	94.23	95.51	96.15	98.08	0.00
96.41	57.44	43.08	33.85	56.41	36.01	100.00	100.00	100.00	100.00	100.00	15.43
100.00	54.12	50.59	37.65	83.53	14.35	98.82	96.47	96.47	95.29	97.65	0.00
100.00	34.19	36.75	41.03	88.03	26.89	97.44	35.04	75.21	79.49	98.29	0.00
98.78	50.06	46.23	42.07	81.19	22.96	98.58	81.44	91.80	92.73	98.50	3.86

TABLE 4.11 – Taux de vrais positifs et de faux positifs (en %) avec nouveau optimisé.

4.5 Conclusion

Ce chapitre présente une nouvelle méthodologie de surveillance afin de détecter séquentiellement d'éventuelles anomalies sur un système de freinage pneumatique. Cette approche s'appuie sur la modélisation dynamique d'un autobus standard ou articulé. Cette modélisation longitudinale est proposée par plusieurs auteurs ([Song et Hedrick, 2004](#); [Nouvelière et Braci, 2007](#)) et nous avons pu valider expérimentalement un certain nombre d'hypothèses simplificatrices. La dynamique du véhicule pourrait sans doute être modélisée plus finement, mais notre choix de modélisation a été jugé pertinent au regard du type de données disponibles. Les signaux physiques sont collectés via des capteurs additionnels ou existants connectés à une architecture télématique dédiée. Ces données CAN sont utilisées afin d'extraire deux indicateurs de santé du système en utilisant conjointement le modèle dynamique et un modèle de régression.

Cette étude nous a permis de valider expérimentalement l'approche de détection séquentielle proposée. Des résultats en termes de détection ont été obtenus sur des jeux de données réelles avec dégradation simulée, car aucun jeu de données lié à une situation de dégradation réelle n'était disponible. Les modes de défaillance choisis correspondent à des défauts interprétables et réalistes sur les systèmes de freinage. Deux cartes multivariées ont été évaluées en terme de détection et la carte de type CUSUM a surpassé celle de Shewhart. Une méthode a été proposée afin d'estimer les seuils de détection de ces deux cartes et celle-ci permet d'améliorer quelque peu les résultats en détection.

Le prochain chapitre de ce manuscrit présente un autre outil d'aide à la maintenance dédié à la surveillance de portes pneumatiques avec une approche de diagnostic différente de celle proposée ici pour les freins.

Surveillance des portes d'accès à partir d'une séquence de courbes

Sommaire

5.1 Introduction	122
5.1.1 Différents systèmes de porte	122
5.1.2 Description du système de porte étudié	123
5.1.3 Démarche des travaux scientifiques	125
5.2 Un modèle de régression pour des courbes multivariées	126
5.2.1 Modèle de régression multivarié à processus caché	126
5.2.2 Cas d'étude simulé dans un cadre non supervisé	132
5.2.3 Apprentissage semi-supervisé des paramètres	136
5.2.4 Cas d'étude simulé dans un cadre semi-supervisé	138
5.3 Stratégie de détection séquentielle dans un flux de courbes	140
5.3.1 Rappel de quelques règles de détection en ligne	140
5.3.2 Approche globale : détection d'anomalies	142
5.3.3 Approche locale : localisation des défauts sur la signature	144
5.3.4 Estimation des seuils de détection et de localisation	146
5.4 Évaluation du détecteur sur données simulées réalistes	149
5.4.1 Travaux expérimentaux	149
5.4.2 Réglage de la structure du modèle de régression	150
5.4.3 Résultats expérimentaux	153
5.5 Expérimentations sur données réelles	156
5.5.1 Acquisition des données de fonctionnement	156
5.5.2 Description des données réelles collectées	156
5.5.3 A priori physique et sélection du modèle	159
5.5.4 Interprétation des défauts réels	161
5.5.5 Résultats en terme de détection, localisation et interprétation	163
5.6 Conclusion	168

5.1 Introduction

5.1.1 Différents systèmes de porte

Certains composants de systèmes de transport, dont les freins et les portes, requièrent un haut niveau de sécurité et de fiabilité. Le système de porte représente un facteur majeur impactant la disponibilité d'un véhicule et un facteur de risque pour les usagers. Chaque porte d'autobus, composée de deux battants, peut être commandée par un moteur ou une électrovalve, respectivement dans le cas d'une porte électrique ou pneumatique, et l'autobus ne peut démarrer si son verrouillage n'est pas acquis. La figure 5.1 présente les trois principaux modes d'ouverture/fermeture des portes d'un autobus.



FIGURE 5.1 – Les principaux systèmes de porte louvoyante : intérieure (haut), extérieure (milieu), ou coulissante (bas).

En Europe, la directive 2001/85/CE assure que : « Toute porte de service commandée et son système de commande doivent être conçus de façon qu'un passager ne puisse être blessé par la porte ou coincé dans la porte quand elle se ferme. » Plusieurs technologies peuvent être employées comme protections et on cite notamment deux systèmes qui ont pour action d'ouvrir les portes après sollicitation. Les *bords sensibles* pneumatiques, illustrés par la figure 5.3, sont constitués d'un profilé en caoutchouc intégrant un ruban de contact de sécurité. Un *tapis sensible* est intégré au sol devant une porte : le poids d'une personne provoque le contact entre les surfaces conductrices à l'intérieur du tapis. Enfin, un seuillage sur le temps maximum d'ouverture/fermeture (porte pneumatique) ou courant maximum du moteur (porte électrique) peut être également envisagé.

5.1.2 Description du système de porte étudié

Les autobus de l'étude sont équipés de portes de type louvoyante intérieure (voir figure 5.1). Des bords sensibles et un tapis électrique équipent chaque porte surveillée. Toutefois, les informations associées à ces équipements de sécurité n'ont pas été disponibles pour cette étude.

Notre stratégie de surveillance des portes s'est concentré sur le suivi de leur fonctionnement opérationnel pendant des mouvements d'ouverture et de fermeture. Les périodes de maintien des battants en position ouverte/fermée sont considérées comme non pertinentes pour le diagnostic des portes. Le système de porte pneumatique est décrit par la figure 5.3 et ce système est essentiellement composé des cinq éléments suivants :

- *Réservoir* : emmagasine l'air comprimé dans des servitudes (de 15 à 20 litres) et alimente en continu les portes du véhicule ;
- *Bloc électrovalve* : reçoit la commande électrique du tableau de bord puis distribue l'air aux deux vérins (un pour chaque battant de porte). Des interrupteurs permettent de distribuer l'air dans une chambre ou l'autre, autour du piston.

- *Vérins pneumatiques* : transforment l'énergie de l'air sous pression en énergie mécanique. Un vérin double-effet, illustré par la figure 5.2, comprend deux chambres à air autour d'un piston et permet de déplacer l'ensemble tige-piston dans les deux sens. L'effort en poussant (sortie de la tige) permet d'ouvrir un battant et, l'effort en tirant (rentrée de la tige) permet de fermer un battant ;

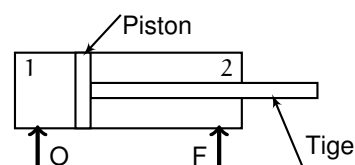


FIGURE 5.2 – Vérin double-effet.

- *Capteurs de pression* : mesurent la pression d'air (entre 0 et 9 bar) dans les chambres autour de chaque piston. Deux capteurs additionnels ont été installés sur chaque vérin, au niveau des orifices O et F (figure 5.2) et sont reliés au calculateur VIDAC. On dénombre donc quatre mesures de pression par porte surveillée :

Pression OAV : pression dans la chambre 1 du vérin avant. Lors d'une admission d'air, la sortie de la tige permet d'ouvrir le battant avant ;

Pression FAV : pression dans la chambre 2 du vérin avant. Lors d'une admission d'air, la rentrée de la tige permet de fermer le battant avant ;

Pression FAR : pression dans la chambre 2 du vérin arrière. Lors d'une admission d'air, la rentrée de la tige permet de fermer le battant arrière ;

Pression OAR : pression dans la chambre 1 du vérin arrière. Lors d'une admission d'air, la sortie de la tige permet d'ouvrir le battant arrière ;

- *Potentiomètres de porte* : mesurent les positions des deux battants de porte. Ces capteurs à résistance variable, de série, sont notamment utilisés pour déterminer un état ouvert et un état fermé de la porte. L'information est disponible sur le CAN constructeur.

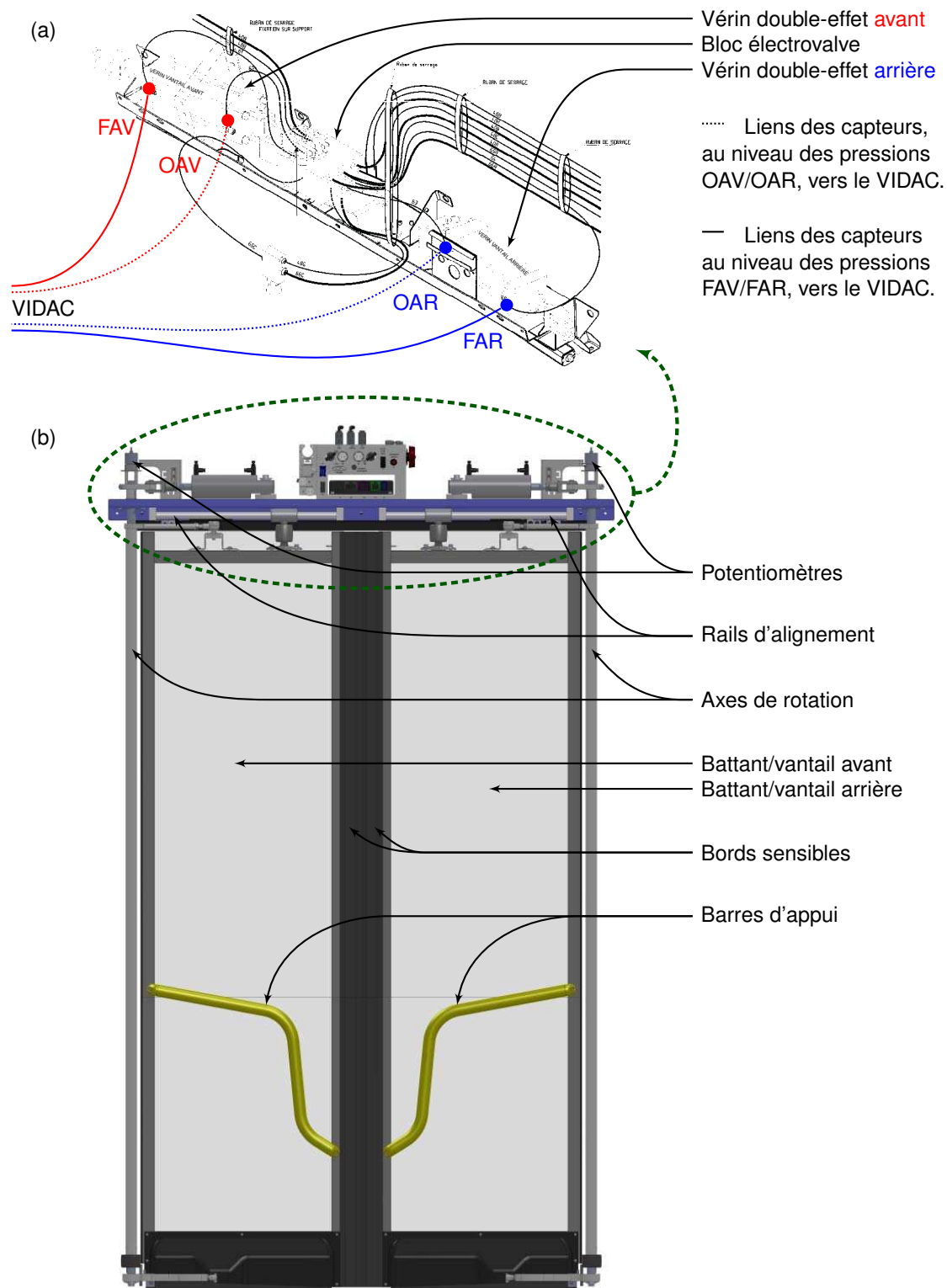


FIGURE 5.3 – Description du système de porte pneumatique et instrumentation.

Les capteurs additionnels de pression, installés sur les vérins double-effet avant et arrière, sont reliés aux ports analogiques du VIDAC (a).

Le mouvement des portes est guidé par le pivot des axes de rotation et le rail de porte sert à aligner chaque battant (b).

Un avantage des portes pneumatiques est la possibilité de régler la vitesse du vérin. En effet, chaque orifice de l'électrovalve, transmettant l'air comprimé aux vérins, est muni d'un *Réducteur de Débit Unidirectionnel* (RDU). Comme le présente la figure 5.4, un RDU est composé d'un limiteur de débit et d'un clapet anti-retour.

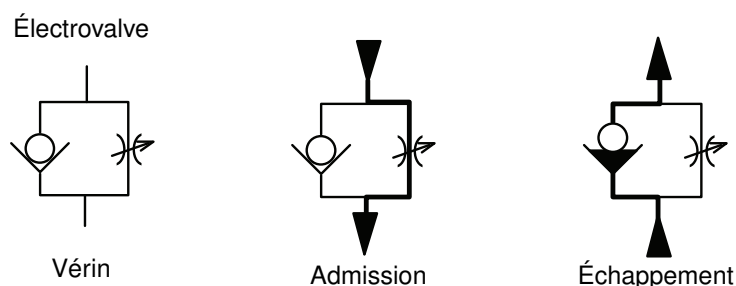


FIGURE 5.4 – Réducteur de Débit Unidirectionnel (RDU).

Sur chaque symbole, le clapet anti-retour est à gauche et le limiteur est à droite. La circulation de l'air est tracée par un trait en gras et le sens est précisé par des flèches.

5.1.3 Démarche des travaux scientifiques

Ce chapitre présente un nouvel outil d'aide à la maintenance dédié à la surveillance des portes décrite en figure 5.3. L'approche adoptée est dite *à base de données* et s'appuie sur l'analyse d'une séquence de courbes. Contrairement au système de freinage, cette stratégie ne nécessite pas de modélisation physique du système.

Ce chapitre vise à décrire une méthode générique adaptée au suivi de fonctionnement d'un système, en particulier, à partir de signaux collectés sur des portes pneumatiques. Lorsque le système est dans un état de fonctionnement normal, nous prendrons l'hypothèse que les signaux mesurés décrivent une certaine zone de l'espace de recherche dans $\mathbb{R}^2$; en mode dégradé, les signaux de ce même système changeront de signature. La section 5.2 propose un modèle de régression flexible afin de représenter la distribution des courbes mesurées. L'apprentissage de ce modèle de mélange est réalisé dans un mode non supervisé et semi-supervisé sur un processus latent discret. A partir de ce modèle de régression paramétrique, nous introduisons une procédure séquentielle de détection dans la section 5.3. On propose de nouveaux critères basés sur le calcul de la vraisemblance complétée pour aider à « localiser » des défauts le long des courbes. Chacune des statistiques par segment, issues de cette approche, est associée à une composante du mélange et l'analyse de ces statistiques aide à la caractérisation d'un nouveau mode de fonctionnement après détection. A partir de défauts connus, il est possible d'identifier l'émergence de nouveaux modes de dégradation. Le détecteur a été évalué sur des cas réels de défaut, dans la section 5.4, suite à une campagne de collecte de données au dépôt en juillet 2011. La procédure de détection/localisation a ensuite été testée sur une séquence de mesures collectées sur des portes en exploitation de janvier à octobre 2012. Cette étude est présentée dans la section 5.5 et s'inscrit dans le cadre d'un cas d'étude du projet européen [EBSF](#).

5.2 Un modèle de régression pour des courbes multivariées

Cette section introduit notre approche de modélisation des courbes, inspirée par le modèle de régression à processus latent proposé par [Chamroukhi \(2010\)](#) déjà présenté dans la sous-section 2.3.2. Nous présentons ici la spécification du modèle de régression pour des courbes (ou trajectoires) **multivariées**, adapté aux signaux de portes, et décrivons l'apprentissage de ce modèle dans un mode non supervisé et semi-supervisé.

5.2.1 Modèle de régression multivarié à processus caché

À l'origine, le modèle de régression à processus caché, intitulé RHLP pour *Regression model with a Hidden Logistic Process* ([Chamroukhi, 2010](#)), était dédié à la modélisation de courbes mono-dimensionnelles. Nous présentons ici une extension de cette approche pour la modélisation de courbes multivariées. À chaque instant, une observation mesurée est un vecteur de taille d . D'autre part, le modèle est décrit de manière à prendre en compte des courbes dont les longueurs peuvent être différentes d'une courbe à l'autre.

Définition du modèle

Soit $\{\mathbf{x}_1, \dots, \mathbf{x}_t\}$, un ensemble de t courbes, où $\mathbf{x}_i = (x_{i1}, \dots, x_{im_i})$ et $x_{ij} \in \mathbb{R}^d$, $d \geq 1$. Chaque courbe $\mathbf{x}_i$ est associée à un vecteur temps $\mathbf{t}_i = (t_1, \dots, t_{m_i})$, où $t_1 < \dots < t_{m_i}$. En outre, on admet qu'à chaque pas de temps t_j , l'observation x_{ij} appartient à un segment parmi K et chaque segment est décrit par un modèle de régression polynomiale d'ordre r . L'observation x_{ij} est supposée avoir été générée par le modèle :

$$x_{ij} = \beta_{z_{ij}} \cdot \mathbf{T}_j + \epsilon_{ij}, \quad \forall (i, j) \in \{1, \dots, t\} \times \{1, \dots, m_i\}, \quad (5.1)$$

où $z_{ij} \in \{1, \dots, K\}$ représente le segment auquel appartient l'observation, $\epsilon_{ij} \sim \mathcal{N}(0, \Sigma_{z_{ij}})$ est un bruit blanc de dimension d et de matrice de covariance $\Sigma_{z_{ij}}$, la matrice $\beta_{z_{ij}}$ est composée des $d \times (r + 1)$ coefficients polynomiaux et $\mathbf{T}_j = (1, t_j, \dots, (t_j)^r)^T$ est le vecteur de covariables. La différence, par rapport à la version mono-dimensionnelle, réside dans la spécification des paramètres β_k et des covariances Σ_k qui sont des matrices dans la version multidimensionnelle. Enfin, il est important de noter que, dans notre application, chaque segment polynomial sera associé à une phase transitoire pendant le fonctionnement des portes (voir la sous-section 5.5.3).

La variable latente Z_{ij} ($\forall i = 1, \dots, t$ et $\forall j = 1, \dots, m_i$) est supposée être distribuée indépendamment selon une loi multinomiale $\mathcal{M}(1, \pi_1(t_j; \mathbf{w}), \dots, \pi_K(t_j; \mathbf{w}))$, où π_k est une fonction logistique :

$$\pi_k(t_j; \mathbf{w}) = P(Z_{ij} = k | t_j) = \frac{\exp(w_{k0} + w_{k1} \cdot t_j)}{\sum_{h=1}^K \exp(w_{h0} + w_{h1} \cdot t_j)}, \quad (5.2)$$

où $\mathbf{w}$ est le vecteur paramètre de la fonction logistique, de dimension $2 \times K$. L'ordre de la logistique a été fixé à 1 afin de garantir la contiguïté entre segments ([Samé et al., 2011](#)). La transformation logistique permet de modéliser dynamiquement des transitions (changements abrupts ou souples) entre segments avec flexibilité (sous-section 2.3.2).

Dans ce cadre, il est possible de montrer que l'observation x_{ij} est distribuée selon un *modèle de mélange* (McLachlan et Peel, 2000) à K composantes :

$$p(x_{ij}|t_j; \theta) = \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \beta_k \cdot \mathbf{T}_j, \Sigma_k), \quad (5.3)$$

où $\theta = (\mathbf{w}, \beta_1, \dots, \beta_K, \Sigma_1, \dots, \Sigma_K)$ est le paramètre du modèle de mélange, π_k est une fonction logistique et représente la probabilité *a priori* que l'observation appartienne à la k^e composante du mélange au temps t_j . Ces quantités pondèrent la combinaison linéaire des K densités normales de chaque modèle de régression simple. Enfin, il est à noter que ces proportions sont strictement positives $\pi_k > 0$ et somment à un : $\sum_{k=1}^K \pi_k = 1$.

Estimation des paramètres

Usuellement, l'estimation des paramètres pour des modèles de mélange s'effectue par la méthode du maximum de vraisemblance (MV). En supposant que les courbes $\mathbf{x}_1, \dots, \mathbf{x}_t$ sont indépendantes conditionnellement au temps $\mathbf{t}$, le paramètre global θ du modèle probabiliste peut être estimé par maximisation de la log-vraisemblance conditionnelle :

$$\begin{aligned} \mathcal{L}(\theta; \mathbf{x}_1, \dots, \mathbf{x}_t, \mathbf{t}) &= \log \prod_{i=1}^t \prod_{j=1}^{m_i} p(x_{ij}|t_j; \theta) \\ &= \sum_{i=1}^t \sum_{j=1}^{m_i} \log \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \beta_k \cdot \mathbf{T}_j, \Sigma_k). \end{aligned} \quad (5.4)$$

La log-vraisemblance est une fonction non linéaire qu'il est impossible de maximiser directement. En effet, le logarithme de la somme ne permet pas d'annuler son gradient. L'optimisation peut être effectuée par l'algorithme Expectation Maximization (EM) qui alterne estimation de la loi *a posteriori* et mise à jour des paramètres, respectivement étapes E et M (Dempster et al., 1977; McLachlan et Krishnan, 2008). Cette stratégie, largement utilisée, se prête bien à l'estimation des paramètres d'un modèle de mélange et, en particulier, ceux d'un modèle génératif à variables latentes comme celui proposé.

La section 2.3.2 décrit la mise en équation de l'algorithme EM dédié au modèle de régression pour des courbes mono-dimensionnelles. On rappelle simplement ici qu'il suffit de maximiser une fonction Q pour faire croître la vraisemblance. Et cette quantité est extraite de la log-vraisemblance complétée $\mathcal{L}_c(\theta; \mathbf{x}, \mathbf{z}, \mathbf{t})$:

$$\begin{aligned} \mathcal{L}_c(\theta; \mathbf{x}, \mathbf{z}, \mathbf{t}) &= \sum_{i=1}^t \sum_{j=1}^{m_i} \log p(x_{ij}, z_{ij}|t_j; \theta) \\ &= \sum_{i=1}^t \sum_{j=1}^{m_i} \sum_{k=1}^K z_{ijk} \cdot \log \left[\pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \beta_k \cdot \mathbf{T}_j, \Sigma_k) \right], \end{aligned} \quad (5.5)$$

avec $z_{ijk} = \begin{cases} 1 & \text{si } z_{ij} = k \\ 0 & \text{sinon} \end{cases}$.

L'algorithme EM alterne les étapes E et M jusqu'à convergence d'un critère sur la vraisemblance (équation 2.25). Conditionnellement au paramètre $\theta^{(q)}$ de l'itération courante q , la fonction Q est l'espérance de la log-vraisemblance complétée conditionnellement aux données observées et au paramètre courant :

$$\begin{aligned} Q(\theta, \theta^{(q)}) &= \mathbb{E}_{\mathbf{Z}} \left[\mathcal{L}_c(\theta; \mathbf{X}, \mathbf{Z}, \mathbf{T}) \mid \mathbf{x}, \mathbf{t}; \theta^{(q)} \right] \\ &= \sum_{i,j,k} \underbrace{\mathbb{E}_{\mathbf{Z}} \left(Z_{ijz} \mid \mathbf{x}, \mathbf{t}; \theta^{(q)} \right)}_{\tau_{ijk}^{(q)}} \cdot \left(\log \pi_k(t_j; \mathbf{w}) + \log \mathcal{N}(x_{ij}; \beta_k \cdot \mathbf{T}_j, \Sigma_k) \right), \end{aligned} \quad (5.6)$$

où $\tau_{ijk}^{(q)}$ représente la loi *a posteriori* qui mesure le degré d'appartenance de l'observation x_{ij} au segment k . L'étape E consiste à calculer les probabilités *a posteriori* à partir des paramètres disponibles à l'itération courante :

$$\begin{aligned} \tau_{ijk}^{(q)} &= P(Z_{ij} = k \mid x_{ij}, t_j; \theta^{(q)}) \\ &= \frac{\pi_k(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_k^{(q)} \cdot \mathbf{T}_j, \Sigma_k^{(q)})}{\sum_{h=1}^K \pi_h(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_h^{(q)} \cdot \mathbf{T}_j, \Sigma_h^{(q)})}. \end{aligned} \quad (5.7)$$

Dans l'étape M, on cherche à mettre à jour chaque paramètre en maximisant la fonction Q . D'après l'équation 5.6, il convient de maximiser deux fonctions Q_1 et Q_2 , respectivement pour la mise à jour des paramètres $\mathbf{w}$ et $(\beta_k, \Sigma_k)_{1 \leq k \leq K}$:

$$\begin{aligned} \theta^{(q+1)} &= \arg \max_{\theta} \left\{ Q(\theta, \theta^{(q)}) \right\} \\ &= \arg \max_{\theta} \left\{ \underbrace{\sum_{i,j,k} \tau_{ijk}^{(q)} \cdot \log \pi_k(t_j; \mathbf{w})}_{Q_1(\theta, \theta^{(q)})} + \underbrace{\sum_{i,j,k} \tau_{ijk}^{(q)} \cdot \log \mathcal{N}(x_{ij}; \beta_k \cdot \mathbf{T}_j, \Sigma_k)}_{Q_2(\theta, \theta^{(q)})} \right\}. \end{aligned} \quad (5.8)$$

Comme dans le cas mono-dimensionnel, le paramètre $\mathbf{w}$ de la fonction logistique est mis à jour via un algorithme Iterative Reweighted Least Squares (IRLS) multi-classe (Green, 1984). Cette stratégie requiert le calcul du vecteur gradient $\nabla_{\mathbf{w}} Q_1$ et celui de la matrice hessienne $\nabla_{\mathbf{w}}^2 Q_1$ - dérivée seconde - de la fonction Q_1 :

$$\frac{\partial}{\partial \mathbf{w}_k} Q_1(\theta, \theta^{(q)}) = \sum_{i,j} \left[\tau_{ijk}^{(q)} - \pi_k(t_j; \mathbf{w}) \right] \cdot \mathbf{v}_j \quad (5.9)$$

$$\frac{\partial^2}{\partial \mathbf{w}_k \partial \mathbf{w}_h^T} Q_1(\theta, \theta^{(q)}) = - \sum_{i,j} \pi_k(t_j; \mathbf{w}) \cdot \left[\delta_{kh} - \pi_h(t_j; \mathbf{w}) \right] \cdot \mathbf{v}_j^T \cdot \mathbf{v}_j, \quad (5.10)$$

où δ_{kh} est le symbole de Kronecker, $\mathbf{v}_j = (1, t_j)^T$ est le vecteur de covariables d'une logistique d'ordre 1, $\frac{\partial}{\partial \mathbf{w}_k} Q_1$ est le k^e élément du gradient et $\frac{\partial^2}{\partial \mathbf{w}_k \partial \mathbf{w}_h^T} Q_1$ est le $(k, h)^e$ élément du hessien (voir la définition A.4).

L'algorithme IRLS est une méthode de type Newton-Raphson adaptée à la régression logistique. Il est possible de montrer que la fonction $Q_1(\theta, \theta^{(q)})$ est convexe même si en pratique, cette fonction peut être complexe à cause de la fonction logistique. L'algorithme IRLS converge itérativement vers un optimum global défini par :

$$\mathbf{w}^{(c+1)} = \mathbf{w}^{(c)} - \left[\nabla_{\mathbf{w}}^2 Q_1(\theta^{(c)}, \theta^{(q)}) \right]^{-1} \cdot \nabla_{\mathbf{w}} Q_1(\theta^{(c)}, \theta^{(q)}), \quad (5.11)$$

où le calcul de $\nabla_{\mathbf{w}} Q_1$ et $\nabla_{\mathbf{w}}^2 Q_1$ est décrit dans les équations 5.9 et 5.10. L'algorithme IRLS itère jusqu'à une stabilisation de la norme (euclidienne) sur le gradient ; nous avons observé en pratique que cette méthode convergeait en moins d'une dizaine d'itérations et que ce nombre d'itérations avait tendance à décroître au cours des itérations de l'EM.

La mise à jour des paramètres $(\beta_k, \Sigma_k)_{1 \leq k \leq K}$ ne nécessite pas la mise en place d'une procédure itérative telle que l'IRLS. D'après l'équation 5.8, il suffit de maximiser la fonction Q_2 par rapport à chaque paramètre, et ce problème peut être résolu analytiquement. Il est succinctement présenté par la suite comment annuler les dérivées partielles de la fonction Q_2 pour les paramètres β_k et Σ_k .

Dans le cas de la mise à jour des coefficients polynomiaux β_k , on reconnaît le problème des *moindres carrés pondérés* par $\tau_{ijk}^{(q)}$ (équation 5.12). L'estimateur qui annule la dérivée partielle $\frac{\partial}{\partial \beta_k} Q_2$ est donné dans l'équation 5.13. La mise à jour de β_k s'écrit :

$$\begin{aligned} \left(\beta_k^{(q+1)} \right)^T &= \arg \max_{\beta_k} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \log \mathcal{N}(x_{ij}; \beta_k \cdot \mathbf{T}_j, \Sigma_k) \\ &= \arg \min_{\beta_k} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \left(x_{ij} - \beta_k \cdot \mathbf{T}_j \right)^T \cdot \Sigma_k^{-1} \cdot \left(x_{ij} - \beta_k \cdot \mathbf{T}_j \right) \end{aligned} \quad (5.12)$$

$$= \left(\sum_{i,j} \tau_{ijk}^{(q)} \cdot \mathbf{T}_j \cdot \mathbf{T}_j^T \right)^{-1} \cdot \left(\sum_{i,j} \tau_{ijk}^{(q)} \cdot \mathbf{T}_j \cdot x_{ij} \right). \quad (5.13)$$

Enfin, la mise à jour de Σ_k correspond au problème d'estimation de la matrice de covariance pour une loi normale multivariée et pondérée par $\tau_{ijk}^{(q)}$ (eq. 5.14). L'estimateur qui annule la dérivée partielle $\frac{\partial}{\partial \Sigma_k} Q_2$ est donné dans l'équation 5.15 :

$$\begin{aligned} \Sigma_k^{(q+1)} &= \arg \max_{\Sigma_k} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \log \mathcal{N}(x_{ij}; \beta_k^{(q+1)} \cdot \mathbf{T}_j, \Sigma_k) \\ &= \arg \min_{\Sigma_k} \sum_{i,j} \tau_{ijk}^{(q)} \cdot \left[\log |\Sigma_k| + (x_{ij} - \beta_k^{(q+1)} \mathbf{T}_j) \Sigma_k^{-1} (x_{ij} - \beta_k^{(q+1)} \mathbf{T}_j)^T \right] \end{aligned} \quad (5.14)$$

$$= \frac{\sum_{i,j} \tau_{ijk}^{(q)} \cdot \left[x_{ij} - \beta_k^{(q+1)} \cdot \mathbf{T}_j \right] \cdot \left[x_{ij} - \beta_k^{(q+1)} \cdot \mathbf{T}_j \right]^T}{\sum_{i,j} \tau_{ijk}^{(q)}}. \quad (5.15)$$

La méthode EM répète les étapes E et M jusqu'à convergence de la log-vraisemblance (monotone et croissante en fonction des itérations de l'algorithme), et retourne une estimation du paramètre global $\hat{\theta}$.

Le pseudo-code 5.1 décrit l'algorithme EM dédié à l'estimation des paramètres du modèle de régression dans une formulation batch. L'algorithme est donné dans un mode non-supervisé et pour des courbes de tailles différentes. La complexité de la méthode proposée dans le pire des cas (Samé et al., 2011), est de $\mathcal{O}(I_{EM} I_{IRLS} t m d^3 K^3 r^3)$, où I_{EM} et I_{IRLS} désignent respectivement le nombre d'itérations des algorithmes EM et IRLS. L'algorithme EM converge vers un optimum local et la qualité de l'estimation du paramètre $\hat{\theta}$ dépend de son initialisation. Si $\theta^{(0)}$ n'est pas fourni, la paramètre θ est initialisé aléatoirement. Concernant l'algorithme IRLS, le paramètre $w^{(c)}$ est initialisé par le paramètre courant $w^{(q)}$ de l'algorithme EM et si $q = 0$, le paramètre $w^{(q)}$ est le vecteur nul.

Algorithme 5.1 : Pseudo-code du modèle de régression à processus latent dans une version multidimensionnelle.

Entrées : Matrice x des t courbes de longueur m_i et de dimension d , nombre K de segments et degré r du polynôme pour chaque courbe

// Initialisation aléatoire de $\theta^{(0)}$ si non fourni

1 $q \leftarrow 0$

2 **Répéter**

 // Étape E : calcul de la loi a posteriori (éq. 5.7)

3 **Pour** $(i, j, k) \in \{1, \dots, t\} \times \{1, \dots, m_i\} \times \{1, \dots, K\}$ **faire**

4

$$\tau_{ijk}^{(q)} \leftarrow \frac{\pi_k(t_j; w^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_k^{(q)} \cdot T_j, \Sigma_k^{(q)})}{\sum_{h=1}^K \pi_h(t_j; w^{(q)}) \cdot \mathcal{N}(x_{ij}; \beta_h^{(q)} \cdot T_j, \Sigma_h^{(q)})} \quad (5.16)$$

 // Étape M : mise à jour des paramètres

 // IRLS : paramètre de la logistique (éq. 5.11)

5

$$w^{(q+1)} \leftarrow \arg \max_w \sum_{i=1}^t \sum_{j=1}^{m_i} \sum_{k=1}^K \tau_{ijk}^{(q)} \cdot \log \pi_k(t_j; w) \quad (5.17)$$

6 **Pour** $k \in \{1, \dots, K\}$ **faire**

7

 // Matrices des coefficients polynomiaux (éq. 5.13)

$$(\beta_k^{(q+1)})^T \leftarrow \left(\sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)} T_j T_j^T \right)^{-1} \cdot \left(\sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)} T_j x_{ij}^T \right) \quad (5.18)$$

8

 // Matrices de variance-covariance (éq. 5.15)

$$\Sigma_k^{(q+1)} \leftarrow \frac{\sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)} \cdot [x_{ij} - \beta_k^{(q+1)} \cdot T_j] \cdot [x_{ij} - \beta_k^{(q+1)} \cdot T_j]^T}{\sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)}} \quad (5.19)$$

9

$q \leftarrow q + 1$

10 **Jusqu'à convergence;**

Sortie : Paramètre estimé $\hat{\theta}$ au sens du maximum de vraisemblance

La mise en œuvre d'un algorithme EM peut être facilitée par la formulation matricielle des équations de mise à jour des paramètres. On reformule donc les équations de mise à jour des paramètres $\mathbf{w}$, $(\beta_k)_{1 \leq k \leq K}$ et $(\Sigma_k)_{1 \leq k \leq K}$. Afin d'alléger la notation, on introduit la variable $m = \sum_{i=1}^t m_i$; on notera par exemple : $\sum_{j=1}^m x_j$ pour $\sum_{i=1}^t \sum_{j=1}^{m_i} x_{ij}$.

• D'après l'équation 5.11, la mise à jour du paramètre $\mathbf{w}$, via l'algorithme IRLS, requiert le calcul du gradient et du hessien de la fonction Q_1 . L'équation 5.9 de mise à jour du gradient, vecteur de taille $2(K-1) \times 1$, peut s'écrire sous la forme matricielle :

$$\nabla_{\mathbf{w}} Q_1 = \mathbb{U} \cdot (\mathbb{T} - \mathbb{\Pi})^T \quad (5.20)$$

$$\text{avec } \begin{cases} \mathbb{U} = \begin{pmatrix} \mathbf{u} & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \mathbf{u} \end{pmatrix}, \text{ une matrice par blocs } (K-1) \times (K-1), \text{ où } \mathbf{u} = \begin{pmatrix} 1 & \cdots & 1 \\ t_1 & \cdots & t_m \end{pmatrix} \\ \mathbb{T} = \begin{pmatrix} \boldsymbol{\tau}_1 & \cdots & \boldsymbol{\tau}_{K-1} \end{pmatrix}, \text{ un vecteur } 1 \times m(K-1), \text{ où } \boldsymbol{\tau}_k = \begin{pmatrix} \tau_{1k}^{(q)} & \cdots & \tau_{mk}^{(q)} \end{pmatrix} \\ \mathbb{\Pi} = \begin{pmatrix} \boldsymbol{\pi}_1 & \cdots & \boldsymbol{\pi}_{K-1} \end{pmatrix}, \text{ un vecteur } 1 \times m(K-1), \text{ où } \boldsymbol{\pi}_k = \begin{pmatrix} \pi_k(t_1; \mathbf{w}) & \cdots & \pi_k(t_m; \mathbf{w}) \end{pmatrix} \end{cases}.$$

Et l'équation 5.10 de mise à jour du hessien, matrice de taille $2(K-1) \times 2(K-1)$, peut s'écrire sous la forme matricielle :

$$\nabla_{\mathbf{w}}^2 Q_1 = \mathbb{U} \cdot \mathbb{V} \cdot \mathbb{U}^T \quad (5.21)$$

$$\text{avec } \mathbb{V} = \begin{pmatrix} V_{11} & \cdots & V_{1,K-1} \\ \vdots & \ddots & \vdots \\ V_{K-1,1} & \cdots & V_{K-1,K-1} \end{pmatrix}, \text{ une matrice de taille } m(K-1) \times m(K-1) \text{ et on a}$$

$$V_{kh} = \begin{pmatrix} \pi_k(t_1; \mathbf{w}^{(c)}) \cdot [\delta_{kh} - \pi_h(t_1; \mathbf{w}^{(c)})] & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \pi_k(t_m; \mathbf{w}^{(c)}) \cdot [\delta_{kh} - \pi_h(t_m; \mathbf{w}^{(c)})] \end{pmatrix}.$$

• D'après l'équation 5.13, la formulation matricielle de l'équation de mise à jour du paramètre $\beta_k^{(q+1)}$ peut s'écrire :

$$\left(\beta_k^{(q+1)} \right)^T = \left(\Gamma_k \cdot \Gamma_k^T \right)^{-1} \cdot \left(\Gamma_k \cdot \mathbb{X}_k^T \right) \quad (5.22)$$

$$\text{avec } \begin{cases} \Gamma_k = \begin{pmatrix} \mathbf{T}_1 \cdot \sqrt{\tau_{1k}^{(q)}} & \cdots & \mathbf{T}_m \cdot \sqrt{\tau_{mk}^{(q)}} \end{pmatrix}, \text{ de taille } (r+1) \times m \\ \mathbb{X}_k = \begin{pmatrix} x_1 \cdot \sqrt{\tau_{1k}^{(q)}} & \cdots & x_m \cdot \sqrt{\tau_{mk}^{(q)}} \end{pmatrix}, \text{ de taille } d \times m \end{cases}.$$

• D'après l'équation 5.15, la formulation matricielle de l'équation de mise à jour du paramètre $\Sigma_k^{(q+1)}$ peut s'écrire :

$$\Sigma_k^{(q+1)} = \frac{1}{\tau_k^{(q)}} \left(\mathbb{X}_k - \beta_k^{(q+1)} \cdot \Gamma_k \right) \cdot \left(\mathbb{X}_k - \beta_k^{(q+1)} \cdot \Gamma_k \right)^T \quad (5.23)$$

$$\text{avec } \tau_k^{(q)} = \sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)}.$$

5.2.2 Cas d'étude simulé dans un cadre non supervisé

A titre d'illustration, nous avons simulé une séquence de courbes homogènes de dimension deux, dont l'échantillonnage est supposé avoir été fixé à 100 Hz. Ces courbes bivariées sont générées à partir d'une unique courbe moyenne représentant un cercle dans $\mathbb{R}^2$, et sont supposées être mesurées sur 10 sec. La courbe moyenne μ est illustrée en rouge dans la figure 5.5 et chaque point de la trajectoire a été simulée en suivant le protocole suivant :

$$\begin{cases} \mu_{j1} &= 50 \cdot (1 - \sin(t_j \cdot \frac{\pi}{5})) \\ \mu_{j2} &= 50 \cdot (1 - \cos(t_j \cdot \frac{\pi}{5})) \end{cases}, \quad \forall 1 \leq j \leq 1000. \quad (5.24)$$

La sinusoïde obtenue est une courbe bivariée de dimension 2×1000 . Un bruit blanc a été ajouté à cette courbe moyenne afin de simuler de nouvelles courbes bruitées. L'ajout de ce bruit est décrit par l'équation 5.1 avec un écart-type fixé à 5. La figure 5.5 illustre la courbe moyenne et trois courbes simulées dans une représentation en fonction du temps (a) et en 2D (b).

Ces signaux ne sont pas a priori décomposables en segments polynomiaux, mais nous allons voir que le modèle RHLP multivarié fournit une solution élégante pour la représentation de ces courbes. Il est à noter que les courbes changent de convexité en dimension 1 (au temps 5 sec) et en dimension 2 (aux temps 2.5 sec et 7.5 sec). Nous allons nous appuyer sur cette connaissance pour segmenter les courbes conjointement sur leurs deux dimensions.

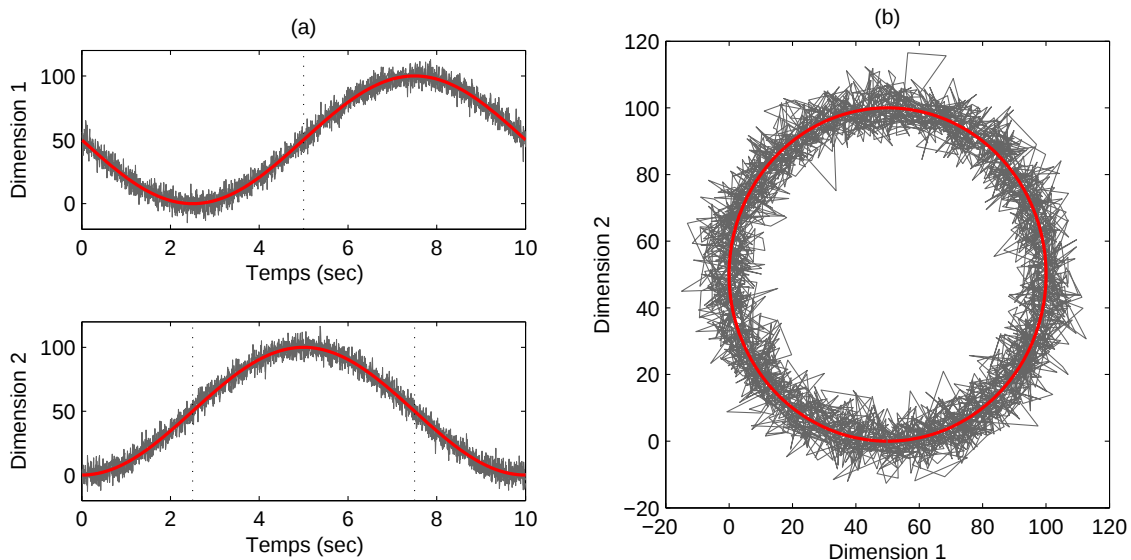


FIGURE 5.5 – Simulation d'un nuage de courbes (grises) et la courbe moyenne (rouge). Représentations 1D de trois courbes simulées et de la courbe moyenne (a), chacune de longueur 1000. Représentation 2D des courbes (b).

Dans cette étude illustrative, le choix a été fait de fixer le nombre de segments à $K = 4$ pour chaque courbe bidimensionnelle. L'ordre des modèles de régression simple a été fixé à $r = 2$ afin de garantir la convexité des segments et d'autoriser un changement de convexité uniquement entre segments. La figure 5.6 représente un modèle de régression appris sur un ensemble de courbes 10 courbes simulées. Chacune des quatre composantes est un modèle de régression de degré 2 appris conjointement sur les deux dimensions des courbes. La courbe moyenne estimée comme mélange des quatre composantes pondéré par la fonction logistique, est illustrée par la figure 5.6(c) en magenta.

Il est évident que les signaux pourraient être également représentés par un unique polynôme de degré plus élevé ($r > 2$). Toutefois nous avons choisi de modéliser les courbes par un mélange de régressions polynomiales avec des transitions très douces. Cette attrait du modèle peut être observé dans la figure 5.6(b) par les pentes douces des fonctions logistiques à l'endroit des transitions. En pratique, les deux hyperparamètres K et r peuvent être élicités par un expert ou, s'ils sont inconnus, doivent faire l'objet d'une recherche préliminaire. Cette tâche de calibration du modèle s'inscrit dans le cadre de la *sélection de modèle* et il existe de nombreux critères de sélection dont une liste non exhaustive est présentée dans la section 2.2.2.

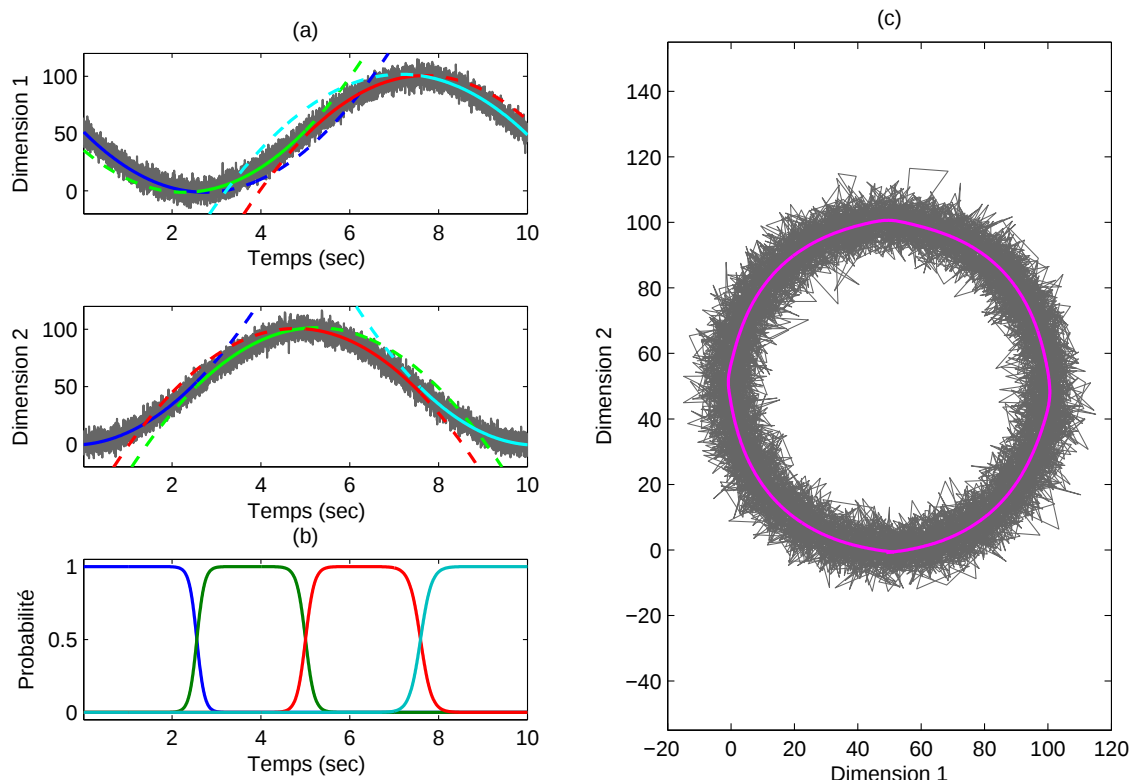


FIGURE 5.6 – Segmentation de dix courbes en quatre morceaux de degré 2. Représentation 1D des courbes (a) et la fonction logistique (b) en fonction du temps. Représentation 2D des courbes et de la courbe moyenne d'un modèle de régression (c).

Afin d'évaluer notre approche, deux critères sont utilisés pour mesurer l'écart entre le processus sous-jacent et le modèle de régression estimé à partir des courbes simulées. Le premier critère est l'*Erreur Quadratique Moyenne* (EQM) qui mesure l'écart entre la courbe moyenne μ et la courbe moyenne estimée $\hat{\mu}$ par le modèle de régression :

$$EQM_1(\mu, \hat{\mu}) = \frac{1}{m \cdot d} \sum_{j=1}^m (\mu_j - \hat{\mu}_j)^T \cdot (\mu_j - \hat{\mu}_j) \quad (5.25)$$

où la courbe moyenne estimée est $\hat{\mu}_j = \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \beta_k \mathbf{T}_j$ (équation 5.3). Le second critère évalue l'EQM entre les temps de changement attendus et ceux estimés par le modèle de régression :

$$EQM_2(\varphi, \hat{\varphi}) = \frac{1}{K-1} \sum_{k=1}^{K-1} (\varphi_k - \hat{\varphi}_k)^2 \quad (5.26)$$

où φ_k représente le temps de changement attendu entre le k^e et le $(k+1)^e$ segment, et $\hat{\varphi}_k$ est le temps de changement entre la k^e et la $(k+1)^e$ composante de la logistique. L'extension semi-supervisée, présentée par la suite, permet de mettre en bijection φ et $\hat{\varphi}$, respectivement les vecteurs de temps de changement attendus et estimés.

Le tableau 5.1 présente quelques indicateurs et notamment la log-vraisemblance de dix estimations du modèle de régression effectuées sur l'ensemble des dix courbes simulées et illustrées dans la figure 5.6. Pour chaque log-vraisemblance obtenue, l'algorithme a été lancé dix fois avec une initialisation aléatoire et le paramètre maximisant la log-vraisemblance a été retenu. On observe que les deux erreurs EQM_1 et EQM_2 sont stables et on remarque notamment une très faible erreur de EQM_2 , 0.015 sec en moyenne. L'algorithme a un comportement raisonnable avec une convergence en une centaine d'itérations pour chaque lancer, en moyenne.

Instance	1	2	3	4	5	6	7	8	9	10
$\mathcal{L}(\theta; \mathbf{x}, \mathbf{t})$	-60409	-60409	-60409	-60410	-60409	-60409	-60409	-60410	-60409	-60409
q	73	75	227	271	132	151	63	278	112	125
EQM_1	0.127	0.132	0.128	0.133	0.130	0.131	0.125	0.126	0.133	0.131
EQM_2	4.57	14.33	4.97	38.83	9.37	19.77	8.93	4.37	25.10	20.73
cputime	69.05	66.91	107.55	107.91	85.80	85.89	80.42	87.95	84.35	90.78

TABLE 5.1 – Quelques indicateurs de performance sur le modèle appris. Pour chaque instance, la log-vraisemblance et le nombre d'itérations correspondent au meilleur candidat parmi dix lancers, après convergence de l'EM. L' EQM_1 mesure l'erreur entre la vraie courbe moyenne et celle estimée. L' EQM_2 mesure l'erreur de segmentation en millisecondes. Le temps CPU (en secondes) mesure le temps nécessaire à l'apprentissage après dix lancers de chaque instance.

La figure 5.7 présente la densité de probabilité d’une estimation du modèle de régression en 2D avec isocontours (a) et en 3D (b). La densité du modèle varie le long de la courbe moyenne et on observe que les zones de transition ainsi que le milieu des segments sont les zones les plus denses. La densité « parfaite » devrait être représentée par une probabilité maximale constante sur la courbe moyenne. La figure permet de visualiser une gradation de probabilité le long de chaque composante et révèle ainsi les limites du modèle de régression proposé.

Enfin, voici le vecteur paramètre $\theta = (\mathbf{w}, \beta_1, \dots, \beta_4, \Sigma_1, \dots, \Sigma_4)$ qui est illustré dans la figure 5.6, et obtenu comme estimateur du maximum de vraisemblance via le modèle régressif avec des matrices de covariance diagonales (Celeux et Govaert, 1995) :

$\mathbf{w}$	β_1	β_2	β_3	β_4
$\begin{pmatrix} -30 & 141 \\ -18 & 111 \\ -8 & 61 \\ 0 & 0 \end{pmatrix}$	$\begin{pmatrix} -7 & 7 \\ -39 & 3 \\ 52 & -1 \end{pmatrix}$	$\begin{pmatrix} 7 & -7 \\ -33 & 73 \\ 36 & -92 \end{pmatrix}$	$\begin{pmatrix} -7 & -7 \\ 112 & 69 \\ -330 & -64 \end{pmatrix}$	$\begin{pmatrix} -7 & 7 \\ 93 & -149 \\ -234 & 762 \end{pmatrix}$
Σ_1	Σ_2	Σ_3	Σ_4	
$\begin{pmatrix} 24.86 & 0 \\ 0 & 23.52 \end{pmatrix}$	$\begin{pmatrix} 24.68 & 0 \\ 0 & 24.43 \end{pmatrix}$	$\begin{pmatrix} 23.47 & 0 \\ 0 & 23.25 \end{pmatrix}$	$\begin{pmatrix} 24.87 & 0 \\ 0 & 24.80 \end{pmatrix}$	

TABLE 5.2 – Paramètres du mélange régressif estimés dans un mode non supervisé sur un ensemble des courbes simulées.

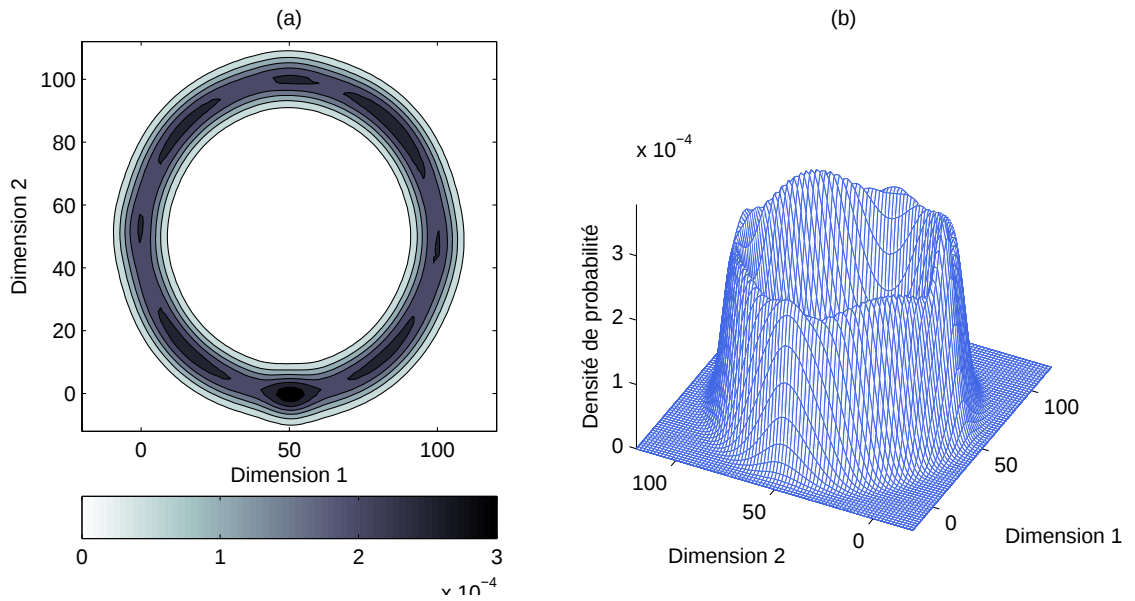


FIGURE 5.7 – Densité de probabilité du modèle de régression (équation 5.3). Représentation de la densité en 2D avec isocontours (a) et en 3D (b).

5.2.3 Apprentissage semi-supervisé des paramètres

Comme présenté dans la sous-section 2.2.2, l'apprentissage *non supervisé* d'un modèle génératif conduit à estimer les étiquettes de toutes les observations. Dans un modèle de mélange, les composantes d'appartenance de chaque observation sont supposées inconnues a priori. L'apprentissage du modèle de régression multivarié à processus latent a été présenté dans un mode non supervisé (algorithme 5.1). Dans l'approche qui a été décrite précédemment, toute l'information disponible sur la segmentation des courbes n'était pas prise en compte. L'optimisation du modèle, portait sur la maximisation de la log-vraisemblance à partir de données non étiquetées (équation 5.4) :

$$\begin{aligned}\mathcal{L}^{\text{NS}}(\boldsymbol{\theta}; \mathbf{x}, \mathbf{t}) &= \sum_{i=1}^t \sum_{j=1}^{m_i} \log p(x_{ij} | \mathbf{t}_i; \boldsymbol{\theta}) \\ &= \sum_{i=1}^t \sum_{j=1}^{m_i} \log \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \boldsymbol{\beta}_k \cdot \mathbf{T}_j, \boldsymbol{\Sigma}_k).\end{aligned}\quad (5.27)$$

L'apprentissage *semi-supervisé* (McLachlan, 1977; Chapelle et al., 2010) se situe à mi-chemin entre l'apprentissage supervisé (toute donnée d'apprentissage est étiquetée) et l'apprentissage non supervisé (aucune donnée n'est étiquetée). Dans notre étude, l'apprentissage semi-supervisé sur le processus caché permet d'intégrer une connaissance partielle de la segmentation sur l'ensemble des courbes de la séquence afin de faciliter l'interprétation des segments. En pratique, on apprend le modèle de régression $\boldsymbol{\theta}$ en supposant connu le sous-ensemble $\mathcal{S} = (S_1 \cup \dots \cup S_K) \subset \{1, \dots, m\}$ formé lui-même de K sous-ensembles disjoints (potentiellement vides), où $m = \max_{1 \leq i \leq t} m_i$ et S_k représente l'ensemble des indices connus pour le k^{e} segment. On obtient alors un critère hybride composé d'une log-vraisemblance complétée sur les observations labellisées et d'une log-vraisemblance classique sur les autres :

$$\begin{aligned}\mathcal{L}^{\text{SS}}(\boldsymbol{\theta}; \mathbf{x}, \mathbf{t}) &= \sum_{i=1}^t \left(\sum_{j \in \mathcal{S}} \log p(x_{ij}, z_{ij} | \mathbf{t}_j; \boldsymbol{\theta}) + \sum_{j \notin \mathcal{S}} \log p(x_{ij} | \mathbf{t}_j; \boldsymbol{\theta}) \right) \\ &= \sum_{i=1}^t \left(\sum_{j \in \mathcal{S}} \log \left(\pi_{z_{ij}}(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \boldsymbol{\beta}_{z_{ij}} \cdot \mathbf{T}_j, \boldsymbol{\Sigma}_{z_{ij}}) \right) \right. \\ &\quad \left. + \sum_{j \notin \mathcal{S}} \log \sum_{k=1}^K \pi_k(t_j; \mathbf{w}) \cdot \mathcal{N}(x_{ij}; \boldsymbol{\beta}_k \cdot \mathbf{T}_j, \boldsymbol{\Sigma}_k) \right)\end{aligned}\quad (5.28)$$

Dans cette extension semi-supervisée, $\mathcal{L}^{\text{SS}}$ est la nouvelle log-vraisemblance du modèle de régression. L'apprentissage d'un tel algorithme converge alors vers un maximum local de la vraisemblance, comme pour l'EM classique en mode non supervisé. En pratique, la prise en compte de l'information des observations labellisées se traduit simplement par une modification de l'étape E. La formulation d'une version semi-supervisée du modèle de régression se réduit donc à calculer différemment les probabilités a posteriori des observations dont les labels de segment sont connus. Seules les probabilités a posteriori des

observations non labellisées sont réestimées à chaque itération, et les autres sont fixées à 0 ou 1 selon la composante d'appartenance. Comme le présente l'algorithme 5.2, la mise en œuvre de l'étape M reste identique au cas non supervisé. Ceci-étant, la mise à jour des paramètres est effectuée avec les vrais probabilités a posteriori pour les observations labellisées.

Algorithme 5.2 : Pseudo-code du modèle de régression à processus latent multivarié dans un mode semi-supervisé.

Entrées : Matrice $\mathbf{x}$ des t courbes de longueur m_i et de dimension d , nombre K de segments et degré r du polynôme pour chaque courbe, ensemble des labels disponibles $\mathcal{S}$

```

// Initialisation aléatoire de  $\boldsymbol{\theta}^{(0)}$  si non fourni
1  $q \leftarrow 0$ 
2 Répéter
    // Étape E : calcul de la loi a posteriori (éq. 5.7)
3   Pour  $(i, k) \in \{1, \dots, t\} \times \{1, \dots, K\}$  faire
4     Pour  $j \in \mathcal{S}$  faire
5        $\tau_{ijk}^{(q)} \leftarrow \mathbb{1}_{z_{ij}=k}$  (5.29)
6     Pour  $j \notin \mathcal{S}$  faire
7        $\tau_{ijk}^{(q)} \leftarrow \frac{\pi_k(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \boldsymbol{\beta}_k^{(q)} \cdot \mathbf{T}_j, \boldsymbol{\Sigma}_k^{(q)})}{\sum_{h=1}^K \pi_h(t_j; \mathbf{w}^{(q)}) \cdot \mathcal{N}(x_{ij}; \boldsymbol{\beta}_h^{(q)} \cdot \mathbf{T}_j, \boldsymbol{\Sigma}_h^{(q)})}$ 
    // Étape M : mise à jour des paramètres
8   // IRLS : paramètre de la logistique (éq. 5.11)
     $\mathbf{w}^{(q+1)} \leftarrow \arg \max_{\mathbf{w}} \sum_{i=1}^t \sum_{j=1}^{m_i} \sum_{k=1}^K \tau_{ijk}^{(q)} \cdot \log \pi_k(t_j; \mathbf{w})$ 
9   Pour  $k \in \{1, \dots, K\}$  faire
10    // Matrices des coefficients polynomiaux (éq. 5.13)
     $(\boldsymbol{\beta}_k^{(q+1)})^T \leftarrow \left( \sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)} \mathbf{T}_j \mathbf{T}_j^T \right)^{-1} \cdot \left( \sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)} \mathbf{T}_j x_{ij}^T \right)$ 
11    // Matrices de variance-covariance (éq. 5.15)
     $\boldsymbol{\Sigma}_k^{(q+1)} \leftarrow \frac{\sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)} \cdot [x_{ij} - \boldsymbol{\beta}_k^{(q)} \cdot \mathbf{T}_j] \cdot [x_{ij} - \boldsymbol{\beta}_k^{(q)} \cdot \mathbf{T}_j]^T}{\sum_{i=1}^t \sum_{j=1}^{m_i} \tau_{ijk}^{(q)}}$ 
12    $q \leftarrow q + 1$ 
13 Jusqu'à convergence;
Sortie : Paramètre estimé  $\hat{\boldsymbol{\theta}}$  au sens du maximum de vraisemblance

```

5.2.4 Cas d'étude simulé dans un cadre semi-supervisé

Reprenons l'exemple du cas d'étude exposé pour la version non supervisée du modèle de régression multivarié. On suppose toujours échantillonner des courbes à 100 Hz sur 10 secondes. La courbe moyenne représente un cercle en 2D et les trois changements de convexité sont interprétés comme des transitions douces entre quatre segments :

$$\forall 1 \leq j \leq 1000, \quad \begin{cases} \mu_{j1} = 50 - 5 \cdot \sin(t_j \cdot \frac{\pi}{5}) \\ \mu_{j2} = 50 - 5 \cdot \cos(t_j \cdot \frac{\pi}{5}) \end{cases} \quad (5.30)$$

Les données ont été fortement bruitées avec un écart-type fixé à 20. Comme illustré par la figure 5.8, il est visuellement difficile de détecter et localiser les changements entre segments à partir des données simulées. Les hyperparamètres restent inchangés par rapport au cas d'étude précédent : nombre de composantes fixé à 4 et ordre des modèles de régression simple fixé à 2 afin d'éviter tout changement de convexité par composante.

La figure 5.9 décrit les résultats en terme d'estimation de la courbe moyenne et de la segmentation des courbes de manière non-supervisée et de manière semi-supervisée (supervision de 20% sur le processus caché). L'intérêt de la semi-supervision est illustré dans la figure 5.9(c). Plus la largeur des bandes jaunes est grande, plus la quantité de labels connus pour chaque segment est importante. On conjecture que le fait de fixer la composante d'appartenance de certaines observations a pour effet d'assurer des transitions entre composantes à l'extérieur des zones labellisées. Ainsi, la forme des transitions (brusque/douce) n'est pas garantie mais leur localisation sera assurée dans des zones non labellisées, appelées par la suite **zones d'incertitude**.

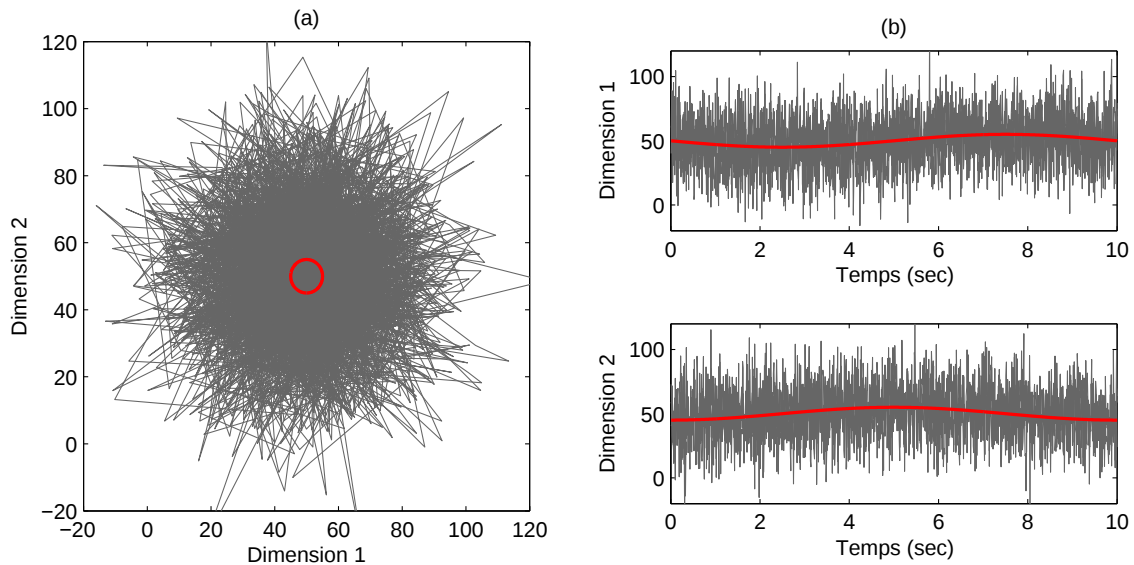


FIGURE 5.8 – Simulation d'un nuage de courbes (grises) et la courbe moyenne (rouge).

Représentation 2D de trois courbes simulées et de la courbe moyenne (a).

Représentation des courbes en fonction du temps (b).

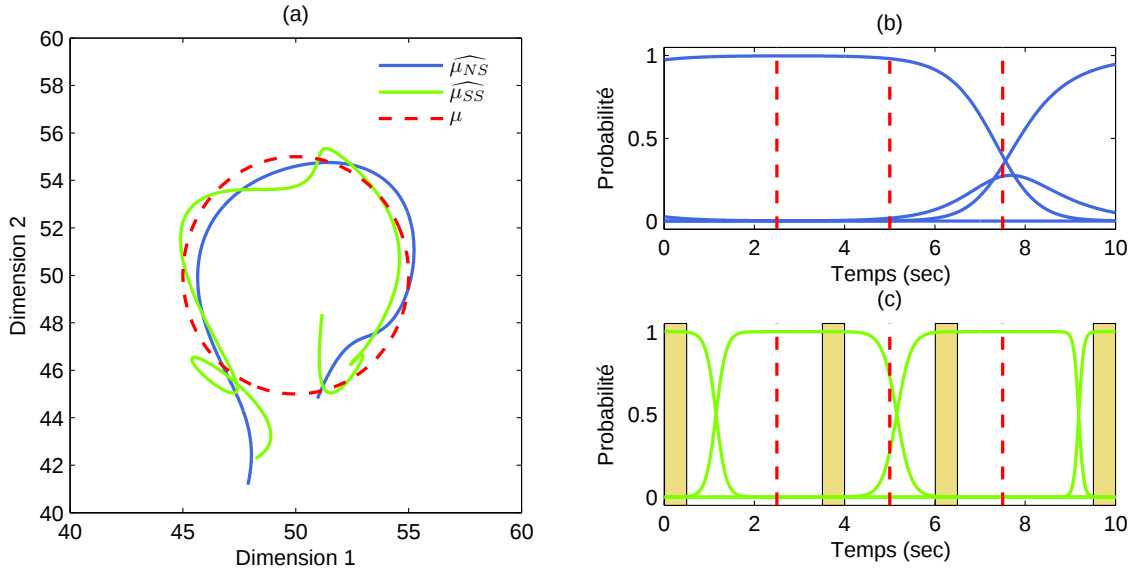


FIGURE 5.9 – Courbes moyennes estimées et segmentation temporelle de deux modèles appris en mode non-supervisé (bleu) ou semi-supervisé (vert).
 Vraie courbe moyenne et estimée en mode non / semi-supervisé à 20% (a).
 Segmentation attendue et fonctions logistiques en fonction du temps (b, c);
 les bandes jaunes représentent les 20% d'étiquettes connus (c).

Le tableau 5.3 présente quelques indicateurs concernant l'apprentissage semi-supervisé de notre modèle de régression sur une séquence de dix courbes. Une colonne correspond à un ratio de labels connus dans chaque segment ; le mode non supervisé est représenté par un ratio de 0%. L'apport d'une semi-supervision sur le processus caché se traduit par une meilleure identification de la segmentation (EQM_2) et une accélération de la convergence de l'algorithme (q et $cputime$). En effet, un apprentissage sans supervision de courbes bruitées ne conduit pas toujours à identifier le vrai nombre de segments sous-jacents ; seuls trois segments ont été identifiés dans la figure 5.9(b). En revanche, l'estimation semi-supervisée de la courbe moyenne n'est pas améliorée (EQM_1) ce qui peut être expliqué par l'importance du bruit pour chaque signal simulé (figure 5.8).

$ \mathcal{S} /m$	0% (non-sup.)	20%	40%	60%	80%	100%
$\mathcal{L}^{SS}(\theta)$	-88396 (3)	-88398 (1)	-88398 (1)	-88398 (0)	-88398 (0)	-88399 (0)
q	35 (22)	40 (13)	28 (6)	23 (8)	19 (3)	2 (0)
EQM_1	0.49 (0.13)	0.54 (0.08)	0.50 (0.02)	0.52 (0.03)	0.50 (0.04)	0.47 (0.00)
EQM_2	2.66 (4.38)	0.71 (0.64)	0.24 (0.31)	0.22 (0.14)	0.07 (0.03)	0.00 (0.00)
$cputime$	30.97 (7.16)	31.49 (2.55)	23.61 (2.35)	19.19 (1.01)	16.56 (1.69)	3.26 (1.46)

TABLE 5.3 – Quelques indicateurs sur le modèle appris en mode semi-supervisé. Pour chaque modèle, la log-vraisemblance et le nombre d'itérations correspondent au meilleur candidat parmi dix lancers, après convergence de l'EM. L' EQM_1 mesure l'erreur entre la vraie courbe moyenne et celle estimée. L' EQM_2 mesure l'erreur de segmentation (sec). Le temps CPU (sec) mesure le temps nécessaire à l'apprentissage de chaque modèle après dix lancers.

5.3 Stratégie de détection séquentielle dans un flux de courbes

Les méthodes statistiques de détection de changement sur une séquence d'observations sont généralement formulées sous la forme de tests séquentiels d'hypothèses. Des états de l'art détaillés de règles de détection classiques et de leurs applications peuvent être notamment trouvés dans les ouvrages de [Basseville et Nikiforov \(1993\)](#) et [Lai \(1995, 2001\)](#). Le chapitre 3 présente le cadre de ce type d'approches et décrit quelques méthodes séquentielles. Ces techniques, comme les cartes de contrôle, prennent le plus souvent des données multivariées en entrée. En d'autres termes, des vecteurs de taille fixe peuvent être traités à chaque instant mais pas des courbes multivariées.

Dans cette section, on introduit une nouvelle stratégie de détection séquentielle d'anomalies dans une séquence de courbes basée sur le modèle de régression multivarié présenté dans la partie précédente (section 5.2). On cherche ici à détecter tout changement dans la distribution de probabilité des courbes, mélange de plusieurs modèles de régression polynomiale, ce qui sera interprété comme un changement d'état du système surveillé. La section 2.3 présente quelques méthodes d'analyse de données fonctionnelles ; on cite le modèle MixRHLP ([Chamroukhi et al., 2013](#)), mélange de modèles de RHLP univariés, pour le regroupement en sous-classes de courbes dans un cadre non séquentiel. Cette section commence par un rapide rappel de quelques règles pour la détection dans une séquence de courbes, puis on introduit une nouvelle stratégie de localisation des défauts. Enfin, on présente une méthode d'estimation des seuils de détection.

5.3.1 Rappel de quelques règles de détection en ligne

On rappelle que la procédure classique du CUSUM (CUMulative SUM), introduite par [Page \(1954\)](#) et présentée dans la sous-section 3.2.1, fait l'hypothèse d'une réception séquentielle d'observations indépendantes. On étend ce type de procédure à des données fonctionnelles supposées indépendantes : $\mathbf{x}_1, \dots, \mathbf{x}_t, \dots$. Dans le cadre d'une détection en ligne, la règle du CUSUM consiste à prendre une décision, dès la réception d'une nouvelle observation, entre les deux hypothèses suivantes :

$$\left\{ \begin{array}{ll} (\mathcal{H}_0) & \mathbf{x}_i \stackrel{\text{iid}}{\sim} p(\mathbf{x}_i; \theta_0) \quad \forall i = 1, \dots, t \\ (\mathcal{H}_1) & \begin{array}{ll} \mathbf{x}_i \stackrel{\text{iid}}{\sim} p(\mathbf{x}_i; \theta_0) & \forall i = 1, \dots, p-1 \\ \mathbf{x}_i \stackrel{\text{iid}}{\sim} p(\mathbf{x}_i; \theta_1) & \forall i = p, \dots, t \end{array} \end{array} \right. \quad (5.31)$$

où p est le *temps de changement*, et $p(\cdot; \theta_0)$ et $p(\cdot; \theta_1)$ sont respectivement les modèles, supposés stationnaires, avant et après changement. Dans la règle du CUSUM, les paramètres θ_0 et θ_1 sont supposés connus. La statistique de détection g_t peut alors s'écrire :

$$g_t = \max_{1 < p \leq t} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i; \theta_1)}{p(\mathbf{x}_i; \theta_0)}. \quad (5.32)$$

Le temps de changement est usuellement détecté par la règle d'arrêt suivante :

$$t_a = \inf \{t \geq 1 : g_t \geq h\} \quad (5.33)$$

où t_a est le *temps d'alarme* et h est un seuil fixé (son estimation est détaillée dans la sous-section 5.3.4). L'optimalité de cette règle a été obtenue dans un cadre asymptotique (Lorden, 1971) au sens du critère 3.1 d'optimalité de Lorden, et dans un cadre non asymptotique (Moustakides, 1986; Ritov, 1990) au sens du critère 3.2 d'optimalité de Pollak et Siegmund (1975). Il est à noter que le critère 3.3 d'optimalité de Lai (1998) est une extension du critère de Lorden et considère une séquence constituée de données dépendantes ou indépendantes.

Les hypothèses du CUSUM ne correspondent pas toujours aux situations rencontrées dans des applications réelles. En particulier, la densité $p(\cdot; \theta_1)$ est rarement connue car il est difficile de connaître à l'avance la forme exacte de la distribution après changement. Le test du rapport de vraisemblance généralisé (GLR) (Basseville et Nikiforov, 1993) peut être envisagé lorsque le paramètre θ_1 est inconnu mais appartient à un certain espace Θ . Ce paramètre peut alors être estimé par le principe du maximum de vraisemblance (MV). La *statistique du GLR* est définie par :

$$g_t = \max_{1 < p \leq t} \sup_{\theta_1 \in \Theta} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i; \theta_1)}{p(\mathbf{x}_i; \theta_0)}. \quad (5.34)$$

Une procédure de type GLR est notamment proposée par Siegmund et Venkatraman (1995) au sens du critère 3.1 de Lorden pour un changement de moyenne dans une séquence de variables normales. Lai (1998) obtient, sous certaines hypothèses, quelques résultats d'optimalité pour la règle du GLR sur des données dépendantes ou indépendantes. Contrairement à la statistique du CUSUM (équation 5.32), celle du GLR n'a pas de formulation récursive (Basseville et Nikiforov, 1993, page 38). En revanche, aussi bien le test du GLR que celui du CUSUM cherchent un point de changement entre les temps 1 et t . Nous présentons par la suite quelques reformulations de ces tests qui nous permettent de nous affranchir de quelques difficultés de passage à l'échelle.

La théorie de la détection séquentielle vise à proposer des procédures qui nécessitent peu de temps de calcul et de mémoire, tout en préservant le mieux possible les propriétés d'optimalité des détecteurs (Lai, 1995). Une procédure de type GLR est souvent coûteuse en temps de calcul à cause de l'estimation de θ_1 et peut s'avérer, en pratique, difficilement implémentable. Une première alternative a été proposée par Willsky et Jones (1976) afin de réduire la complexité computationnelle de l'étape 5.34 en maintenant un temps de calcul constant à chaque instant t . L'idée est de borner la recherche d'un point de changement sur une fenêtre de taille fixée W , et de ne plus démarrer cette recherche depuis la première observation. La *statistique du GLR à fenêtre limitée* est définie par :

$$g_t = \max_{t-W < p \leq t} \sup_{\theta_1 \in \Theta} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i; \theta_1)}{p(\mathbf{x}_i; \theta_0)}. \quad (5.35)$$

D'autres solutions ont été proposées afin de réduire la complexité computationnelle d'une procédure de type GLR. L'*approche asymptotique locale*, décrite par [Basseville et Nikiforov \(1993, pages 126-129\)](#), permet de remplacer la tâche de maximisation de la vraisemblance par des approches locales ; l'idée est ici de surveiller des résidus à partir du gradient de la log-vraisemblance (fonction score). D'autre part, [Nikiforov \(2000, 2003\)](#) propose une statistique avec une formulation récursive et l'optimalité de la procédure est établie à partir d'un critère lié à un délai moyen maximum de détection/localisation.

5.3.2 Approche globale : détection d'anomalies

Dans notre étude, nous considérons que les paramètres θ_0 et θ_1 (avant et après changement) sont inconnus. Ces deux paramètres sont estimés séparément par le principe du maximum de vraisemblance (équation 5.4) via l'algorithme EM synthétisé par le pseudo-code 5.1. Ce choix de modélisation induit un domaine de définition $\Theta \subset \mathbb{R}^v$ des paramètres, où v est la dimension du paramètre θ . Dans la version du test GLR à fenêtre limitée de taille W , le calcul de chaque statistique implique donc $2 \times W$ maximisations. A chaque instant t , on estime le *temps de changement courant* par :

$$p^* = \arg \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i | t; \hat{\theta}_1)}{p(\mathbf{x}_i | t; \hat{\theta}_0)} \quad (5.36)$$

où $\hat{\theta}_0$ et $\hat{\theta}_1$ sont les paramètres maximisant respectivement les vraisemblances $p(\mathbf{x}_1, \dots, \mathbf{x}_{p-1})$ et $p(\mathbf{x}_p, \dots, \mathbf{x}_t)$. Afin de faciliter l'implémentation en ligne de cette procédure et de réduire sa complexité computationnelle, nous introduisons quelques modifications supplémentaires.

La première idée est de borner l'intervalle de recherche globale du détecteur par une fenêtre $[t - M; t]$ de taille $M + 1$, où M ($M > W$) sera appelée *mémoire du détecteur*. Sans cette précaution, le paramètre $\hat{\theta}_0$ d'avant changement serait estimé sur un échantillon dont la taille tendrait vers l'infini quand $t \rightarrow \infty$. Pour chacun des W rapports de log-vraisemblance cumulés ($t - W < p \leq t$), on estime les paramètres avant et après changement sur des échantillons respectivement de taille $M + p - t$ et $t - p + 1$:

$$\begin{cases} \hat{\theta}_0 = \sup_{\theta \in \Theta} \mathcal{L}(\theta; \mathbf{x}_{t-M}, \dots, \mathbf{x}_{p-1}, t) \\ \hat{\theta}_1 = \sup_{\theta \in \Theta} \mathcal{L}(\theta; \mathbf{x}_p, \dots, \mathbf{x}_t, t) \end{cases} \quad (5.37)$$

où la log-vraisemblance $\mathcal{L}$ du modèle de régression est donnée par l'équation 5.4.

La seconde idée est d'éviter de tester plusieurs fois le même point de changement en calculant une seule statistique pour toutes les W courbes¹. Ainsi, chaque courbe de la séquence n'est évaluée qu'une seule fois en tant que « courbe de changement ». Cette stratégie permet d'accélérer la procédure de détection d'un facteur $\mathcal{O}(W)$ mais la prise de décision (hypothèses 5.31) n'est menée qu'après la réception de W nouvelles courbes.

1. Cette stratégie pouvant s'avérer sous-optimale, elle sera abandonnée pour le cas de données réelles.

La troisième idée est d'accélérer l'estimation des paramètres par l'algorithme EM (équation 5.37) en réduisant leur domaine de définition Θ . Les paramètres avant et après changement du modèle de régression multivarié sont estimés de manière semi-supervisée. Comme illustré dans le cas d'étude simulé, l'apport de la semi-supervision se traduit notamment par une accélération de la convergence de l'algorithme. La sous-section 5.5.3 détaille l'identification de l'ensemble $\mathcal{S}$ des labels pour notre application. D'autre part, l'initialisation non aléatoire, qui sera précisée par la suite, permet d'accélérer le comportement de l'algorithme EM.

L'algorithme 5.3 présente la stratégie proposée pour la détection d'anomalies sur une séquence de courbes avec la prise en compte des trois idées décrites précédemment. L'estimation du seuil de détection est détaillée dans la sous-section 5.3.4.

Algorithme 5.3 : Pseudo-code de la stratégie proposée pour la détection séquentielle sur des données fonctionnelles.

Entrées : Séquence de courbes $(\mathbf{x}_1, \mathbf{x}_2, \dots)$, mémoire M et nombre de courbes sous $\mathcal{H}_0$, taille W des fenêtres locales, nombre K de segments et degré r du polynôme pour chaque courbe, et seuil de détection h

```

1  $t_a \leftarrow 1$ 
  // Initialisation de l'hypothèse nulle  $\mathcal{H}_0$ 
2  $t \leftarrow t_a + M - 1$ 

3 Tant que il y a  $W$  nouvelles courbes faire
4   // Calcul de la statistique de détection
   
$$g_t \leftarrow \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i | t; \hat{\theta}_1)}{p(\mathbf{x}_i | t; \hat{\theta}_0)}$$


5   // Estimation du temps de changement (équation 5.36)
   
$$p^* \leftarrow \arg \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i | t; \hat{\theta}_1)}{p(\mathbf{x}_i | t; \hat{\theta}_0)}$$


   // Estimations des paramètres avant/après changement par
   // le principe du MV (équation 5.37)
6   Avec  $\begin{cases} \hat{\theta}_0 \leftarrow \text{EM}((\mathbf{x}_{t-M}, \dots, \mathbf{x}_{p-1}), r, K) \\ \hat{\theta}_1 \leftarrow \text{EM}((\mathbf{x}_p, \dots, \mathbf{x}_t), r, K) \end{cases}$ 

7   Si  $g_t < h$  alors
8      $t \leftarrow t + W$  // Recherche sur des fenêtres locales contigües
9   Sinon
10     $t_a \leftarrow p^*$ 
    // Retourne le temps d'alarme
    // Puis, démarre une nouvelle initialisation

```

Sortie : Temps de changement estimés t_a

5.3.3 Approche locale : localisation des défauts sur la signature

Notre approche de *localisation* est motivée par le fait que les modes de dégradations des portes affectent certains segments de la signature d'ouverture/fermeture plus que d'autres. Le terme « localisation » sera utilisé pour désigner l'opérateur de recherche de zones sur la signature temporelle qui traduisent l'existence de ces dégradations.

Si on choisit de modéliser la densité des courbes par le modèle de mélange donné par l'équation 5.3, la **statistique de détection** peut alors s'écrire comme :

$$\begin{aligned}
 g_t &= \sum_{i=p^*}^t \sum_{j=1}^{m_i} \log \frac{p(x_{ij}|t_j; \theta_1)}{p(x_{ij}|t_j; \theta_0)} \\
 &= \sum_{i=p^*}^t \sum_{j=1}^{m_i} \left(\log \sum_{k=1}^K p(x_{ij}, Z_{ij} = k|t_j; \theta_1) \right. \\
 &\quad \left. - \log \sum_{h=1}^K p(x_{ij}, Z_{ij} = h|t_j; \theta_0) \right) \quad (5.38)
 \end{aligned}$$

où p^* est le temps de changement estimé courant (équation 5.36). La statistique de l'équation 5.38 utilise la vraisemblance (équation 5.4) des modèles de régression avant et après changement. Cette formulation ne nous permet pas de calculer un indicateur par segment à cause des logarithmes de sommes sur k . Les vraisemblances complétées (équation 5.5) supposent connues les composantes d'appartenance de chaque observation. On obtient une nouvelle statistique décomposable par segment :

$$\begin{aligned}
 g_t^{(c)} &= \sum_{i=p^*}^t \sum_{j=1}^{m_i} \sum_{k=1}^K z_{ijk}(\theta_1) \cdot \log p(x_{ij}, Z_{ij} = k|t_j; \theta_1) \\
 &\quad - z_{ijk}(\theta_0) \cdot \log p(x_{ij}, Z_{ij} = k|t_j; \theta_0) \quad (5.39)
 \end{aligned}$$

On suppose ici que les composantes des mélanges $p(\cdot; \theta_0)$ et $p(\cdot; \theta_1)$ sont ordonnées de la même manière. On choisit z_{ijk} par la règle du *maximum a posteriori* (MAP) tel que :

$$z_{ijk}(\theta) = \begin{cases} 1 & \text{si } k \text{ a la plus grande probabilité a posteriori (équation 5.7)} \\ 0 & \text{sinon} \end{cases} \quad (5.40)$$

On remarque que : $g_t^{(c)} = \sum_{k=1}^K g_{t,k}^{(c)}$, où $g_{t,k}^{(c)}$ est définie par :

$$\begin{aligned}
 g_{t,k}^{(c)} &= \sum_{i=p^*}^t \sum_{j=1}^{m_i} z_{ijk}(\theta_1) \cdot \log \left[\pi_k(t_j; \mathbf{w}_1) \cdot \mathcal{N}(x_{ij}; (\boldsymbol{\beta}_1)_k \cdot \mathbf{T}_j, (\boldsymbol{\Sigma}_1)_k) \right] \\
 &\quad - z_{ijk}(\theta_0) \cdot \log \left[\pi_k(t_j; \mathbf{w}_0) \cdot \mathcal{N}(x_{ij}; (\boldsymbol{\beta}_0)_k \cdot \mathbf{T}_j, (\boldsymbol{\Sigma}_0)_k) \right] \quad (5.41)
 \end{aligned}$$

Cette quantité peut être considérée comme une statistique par segment, que nous désignons également sous le nom de **statistique de localisation**. Comme pour la tâche de détection (équation 5.3.4), les K statistiques résultantes sont comparées à des seuils dont les estimations seront précisées dans la section 5.3.4. Après détection, un test sur les statistiques de localisation peut nous aider à identifier le mode de dégradation.

L'algorithme 5.4 présente la stratégie proposée pour la détection et localisation de courbes anormales dans un flux de courbes multivariées. Incrémenter la procédure de 1 à chaque itération ralentit l'algorithme (ligne 8) mais permet de prendre une décision dès la réception d'une nouvelle courbe. L'estimation des seuils est détaillée dans la prochaine sous-section.

Algorithme 5.4 : Pseudo-code de la stratégie proposée pour la détection séquentielle et la localisation de défauts sur des données fonctionnelles.

Entrées : Séquence de courbes $(\mathbf{x}_1, \mathbf{x}_2, \dots)$, mémoire M et nombre de courbes sous $\mathcal{H}_0$, taille W des fenêtres locales, nombre K de segments et degré r du polynôme pour chaque courbe, seuil de détection h et seuils de localisation pour chaque segment $(h_k)_{1 \leq k \leq K}$

```

1  $t_a \leftarrow 1$ 
  // Initialisation de l'hypothèse nulle  $\mathcal{H}_0$ 
2  $t \leftarrow t_a + M - 1$ 
3 Tant que il y a une nouvelle courbe faire
4   // Calcul de la statistique de détection
      
$$g_t \leftarrow \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i | t; \hat{\theta}_1)}{p(\mathbf{x}_i | t; \hat{\theta}_0)}$$

5   // Estimation du temps de changement (équation 5.36)
      
$$p^* \leftarrow \arg \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(\mathbf{x}_i | t; \hat{\theta}_1)}{p(\mathbf{x}_i | t; \hat{\theta}_0)}$$

      // Estimations des paramètres avant/après changement par
      // le principe du MV (équation 5.37)
6   Avec  $\begin{cases} \hat{\theta}_0 \leftarrow \text{EM}((\mathbf{x}_{t-M}, \dots, \mathbf{x}_{p-1}), r, K) \\ \hat{\theta}_1 \leftarrow \text{EM}((\mathbf{x}_p, \dots, \mathbf{x}_t), r, K) \end{cases}$ .
7   Si  $g_t < h$  alors
8     |  $t \leftarrow t + 1$ 
9   Sinon
10    |  $t_a \leftarrow p^*$ 
      // Retourne le temps d'alarme
11    | Pour  $k \in \{1, \dots, K\}$  faire
12      | // Statistique de localisation (équation 5.41)
          
$$g_{t,k}^{(c)} \leftarrow \sum_{i=t_a}^t \sum_{j=1}^{m_i} z_{ijk}(\hat{\theta}_1) \cdot \log p(\mathbf{x}_{ij}, Z_{ij} = k | t_j; \hat{\theta}_1) \\ - z_{ijk}(\hat{\theta}_0) \cdot \log p(\mathbf{x}_{ij}, Z_{ij} = k | t_j; \hat{\theta}_0)$$

13    |  $\mathcal{K} \leftarrow \{k : g_{t,k}^{(c)} \geq h_k\}$ 
      // Puis, démarre une nouvelle initialisation

```

Sortie : Temps de changement estimés t_a et segments anormaux $\mathcal{K}$

5.3.4 Estimation des seuils de détection et de localisation

La formulation d'un problème de détection sous la forme de tests séquentiels d'hypothèses nécessite de choisir un seuil de test (voir équation 5.3.4). Comme décrit dans le chapitre 3, on cherche à définir une famille de détecteurs afin de satisfaire deux objectifs antagonistes : minimiser le nombre de fausses alarmes et minimiser le retard à la détection. Estimer un seuil d'une valeur trop faible augmenterait le nombre de fausses alarmes, et une valeur de seuil trop haute allongerait le retard à la détection voire provoquerait l'absence de toute détection. On se place dans le cadre des méthodes séquentielles étudiées par [Lorden \(1971\)](#) ou encore [Moustakides \(1986\)](#) afin de vérifier la contrainte d'ARL sur la durée moyenne entre deux fausses alarmes (équation 3.6). Pour rappel, on définit le temps d'alarme de notre détecteur de sorte que :

$$t_a = \inf \{t \geq 1 : g_t \geq h\}$$

où g_t est la statistique de détection. On choisit $h = h_\gamma$ afin que la contrainte d'ARL (Average Run Length) soit vérifiée ([Lai, 1995](#)) :

$$\mathbb{E}(t_a; \theta_0) \geq \gamma(1 + o(1)) \quad \text{quand } \gamma \rightarrow \infty \quad (5.42)$$

où $\mathbb{E}(t_a; \cdot)$ est la fonction ARL et γ est la durée moyenne entre deux fausses alarmes. En nous plaçant dans le cadre asymptotique de [Lorden \(1971\)](#), on pose alors $\mathbb{E}(t_a | \theta_0) = \gamma$. En d'autres termes, le seuil de détection h_γ doit être choisi afin que la durée moyenne entre deux fausses alarmes soit γ . Il est parfois possible de calculer ou d'approcher la fonction ARL d'un détecteur. [Basseville et Nikiforov \(1993\)](#), sous-section 5.2.2) présentent une étude de la fonction ARL pour le test du CUSUM avec données indépendantes dans laquelle est décrit le calcul des approximations et des bornes sur l'ARL. Afin d'approcher cette fonction, les auteurs s'appuient sur les approximations de Wald ou de Siegmund. En outre, il est possible d'approcher la fonction ARL par une méthode empirique à base de simulations ; nous nous sommes inspirés de cette approche afin d'estimer le seuil de notre procédure de type GLR. [Crosier \(1988, Appendix A\)](#) décrit une méthode itérative pour estimer le seuil de détection d'une version multivariée du CUSUM dont la détection concerne un changement de moyenne dans une distribution gaussienne multivariée. L'auteur propose de fixer un seuil puis de calculer les statistiques de test à partir de nombreuses simulations. Si la durée moyenne entre deux fausses alarmes est en deçà de la durée γ , alors il faut augmenter la valeur du seuil ; sinon, il faut diminuer la valeur du seuil. Cette méthode itère sur chaque nouvelle valeur du seuil et celle-ci est incrémentée ou décrétementée d'une petite valeur de 10^{-1} à chaque itération jusqu'à converger vers une valeur de seuil $\hat{h}_\gamma$. Cette stratégie est similaire à celle que propose [Verdier \(2007, page 61\)](#) afin d'estimer la fonction ARL de façon empirique. On fixe une valeur de seuil et on applique un test sur m trajectoires simulées sous l'hypothèse $\mathcal{H}_0$. La durée moyenne entre deux fausses alarmes est alors la moyenne des m temps d'alarmes. On répète cette procédure en faisant varier le seuil jusqu'à atteindre la durée moyenne souhaitée γ .

L'étude de la fonction ARL s'avère souvent complexe et peut difficilement être intégrée à une procédure de type GLR dans le cadre d'une détection en ligne. Il a été montré dans la sous-section 3.2.3 qu'il existe un lien entre la durée γ et le taux α de fausses alarmes : $\gamma = 1/\alpha$ (équation 3.33). En supposant que le processus est stationnaire, il est possible de fixer un seuil $h = h_\alpha$ où α est considéré comme constant pendant toute la procédure avec comme justification l'optimalité asymptotique d'une procédure de type GLR (Lai, 1998). On estime alors h_α comme le quantile d'ordre $(1 - \alpha)$ de la distribution de la statistique de test g_t . Comme pour la fonction ARL, l'étude de la statistique de test peut être assez complexe. Il est alors possible de choisir une loi a priori sur la distribution des statistiques de test afin d'estimer ses quantiles. On cite notamment les travaux de Gustafsson et Palmqvist (2000) qui estiment un seuil en s'appuyant sur une distribution paramétrique de la statistique ; cet a priori est motivé par une étude basée sur des données réelles et simulées, et permet d'approcher des seuils y compris lorsque les taux de fausses alarmes sont très faibles.

Nous proposons une stratégie d'estimation des seuils de détection et de localisation en utilisant conjointement l'approche de ré-échantillonnage à base de simulations (comme pour l'étude de la fonction ARL) et l'estimation d'un quantile de la distribution de la statistique de test. Un modèle de régression θ_0 est appris sur un échantillon de courbes sous $\mathcal{H}_0$ supposé en bon fonctionnement (de taille M). Dans un second temps, de nombreuses courbes sont simulées à partir du modèle génératif paramétré par θ_0 et défini par l'équation 5.3. Enfin, des statistiques de test sont calculées à partir des courbes générées et à partir d'un taux α de fausses alarmes fourni par l'utilisateur, le seuil est approché comme le quantile d'ordre $(1 - \alpha)$ de leur distribution. Le seuil de détection h_α est estimé à partir des statistiques de détection définies par l'équation 5.38. Et les K seuils de localisation $h_{\alpha,k}$ sont estimés à partir des statistiques de localisation définies par l'équation 5.41. On remarque que $h_\alpha \geq h_{\alpha,k}, \forall k \in \{1, \dots, K\}$. Il est à noter que notre objectif n'est pas ici d'obtenir des approximations précises des seuils de test. On propose simplement une méthode facilement implémentable afin d'estimer un seuil qui garantit en moyenne un taux de fausses alarmes choisi par l'utilisateur.

L'algorithme 5.5 décrit la stratégie de détection et localisation avec une estimation des seuils à chaque initialisation du détecteur. Le choix de fixer les seuils sur toute la durée de la procédure peut sembler « rigide » malgré un gain de temps important. Toutefois, un aspect adaptatif a été intégré à la procédure. Les paramètres d'initialisation des modèles de régression sont actualisés à chaque pas de temps (ligne 10) dans le but de relaxer légèrement l'hypothèse de stationnarité sur les modes définis par les paramètres θ_0 et θ_1 tout en autorisant une certaine variabilité du processus observé en fonction du temps. A chaque itération, la fenêtre de recherche du détecteur est de M , où M est la mémoire du détecteur (nombre de courbes connues sous $\mathcal{H}_0$). La complexité dans le pire des cas est donc : $\mathcal{O}(W \times I_{EM} I_{IRLS} M m d^3 K^3 r^3)$, où W est la taille de la fenêtre locale, I_{EM} et I_{IRLS} représentent respectivement le nombre d'itérations des algorithmes EM et IRLS.

Algorithme 5.5 : Pseudo-code de la stratégie proposée pour la détection séquentielle et la localisation de défauts sur des données fonctionnelles, avec estimation des seuils de détection/localisation.

Entrées : Séquence de courbes $(x_1, x_2, \dots)$, mémoire M ou nombre de courbes sous $\mathcal{H}_0$, taille W des fenêtres locales, nombre K de segments et degré r du polynôme pour chaque courbe, taux α de fausses alarmes choisi par l'utilisateur

```

1  $t_a \leftarrow 1$ 
   // Initialisation de l'hypothèse nulle  $\mathcal{H}_0$ 
2  $t \leftarrow t_a + M - 1$ 
3  $\hat{\theta}_0^{\text{init}} \leftarrow \text{EM}((x_{t_a}, \dots, x_t), r, K)$ 
4  $(h_\alpha, (h_{\alpha,k})_{1 \leq k \leq K}) \leftarrow \text{estimation des seuils à en fonction de } \alpha$ 
5 Tant que il y a une nouvelle courbe faire
6   // Calcul de la statistique de détection
      
$$g_t \leftarrow \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(x_i | t; \hat{\theta}_1)}{p(x_i | t; \hat{\theta}_0)}$$

7   // Estimation du temps de changement (équation 5.36)
      
$$p^* \leftarrow \arg \max_{t-W < p \leq t} \sum_{i=p}^t \log \frac{p(x_i | t; \hat{\theta}_1)}{p(x_i | t; \hat{\theta}_0)}$$

   // Estimations des paramètres avant/après changement par
   // le principe du MV (équation 5.37)
8   Avec  $\begin{cases} \hat{\theta}_0 \leftarrow \text{EM}((x_{t-M}, \dots, x_{p-1}), r, K, \hat{\theta}_0^{\text{init}}) \\ \hat{\theta}_1 \leftarrow \text{EM}((x_p, \dots, x_t), r, K, \hat{\theta}_1^{\text{init}}) \end{cases}$  .

9   Si  $g_t < h_\alpha$  alors
10      $(\hat{\theta}_0^{\text{init}}, \hat{\theta}_1^{\text{init}}) \leftarrow (\hat{\theta}_0, \hat{\theta}_1)$ 
11      $t \leftarrow t + 1$ 
12   Sinon
13      $t_a \leftarrow p^*$ 
      // Retourne le temps d'alarme
14     Pour  $k \in \{1, \dots, K\}$  faire
15       // Statistique de localisation (équation 5.41)
          
$$g_{t,k}^{(c)} \leftarrow \sum_{i=t_a}^t \sum_{j=1}^{m_i} z_{ijk}(\hat{\theta}_1) \cdot \log p(x_{ij}, Z_{ij} = k | t_j; \hat{\theta}_1) \\ - z_{ijk}(\hat{\theta}_0) \cdot \log p(x_{ij}, Z_{ij} = k | t_j; \hat{\theta}_0)$$

16      $\mathcal{K} \leftarrow \{k : g_{t,k}^{(c)} \geq h_{\alpha,k}\}$ 
      // Puis, démarre une nouvelle initialisation

```

Sortie : Temps de changement estimés t_a et segments anormaux $\mathcal{K}$

5.4 Évaluation du détecteur sur données simulées réalistes

5.4.1 Travaux expérimentaux

Afin d'évaluer notre approche de détection sur le système de porte d'autobus décrite dans la sous-section 5.1.2, une campagne de collecte de données a été menée au dépôt en juillet 2011 sur un autobus articulé de 18 mètres. Cette campagne préliminaire a été réalisée avant que l'architecture télématique installée (décrite dans la sous-section 1.3.2) ne soit pleinement opérationnelle. Cette section présente les résultats que nous avons pu obtenir sur la base de cette expérimentation relativement légère. Pour l'évaluation du détecteur, le choix a été fait de générer plusieurs séquences de courbes à l'aide d'un modèle MRHLP appris sur les données réelles collectées. Par la suite, la section 5.5 approfondit notre approche de surveillance des portes et présente des travaux menés sur les données réelles d'un autobus en exploitation dans le cadre du projet EBSF.

Un ensemble de 11 cycles en bon fonctionnement et un cycle avec défaut (blocage) a été collecté en juillet 2011 avec une fréquence de 100 Hz. La figure 5.10 présente un exemple d'ouverture et de fermeture en fonction du temps à partir de six mesures : les quatre mesures de pression (en bar) sont décrites dans la sous-section 5.1.2 et les deux positions des battants avant et arrière sont normalisées (en %). On précise d'une part que les quatre pressions de porte sont mesurées par des capteurs spécifiques à cette étude qui ne correspondent pas à ceux installés dans le cadre du projet EBSF. Ceci-dit, les mesures de pressions sont compatibles avec la description donnée dans la figure 5.3(a). D'autre part, les positions ont été collectées directement à partir des tensions mesurées aux bornes des potentiomètres préexistants (voir figure 5.3(b)) sans utilisation du CAN.

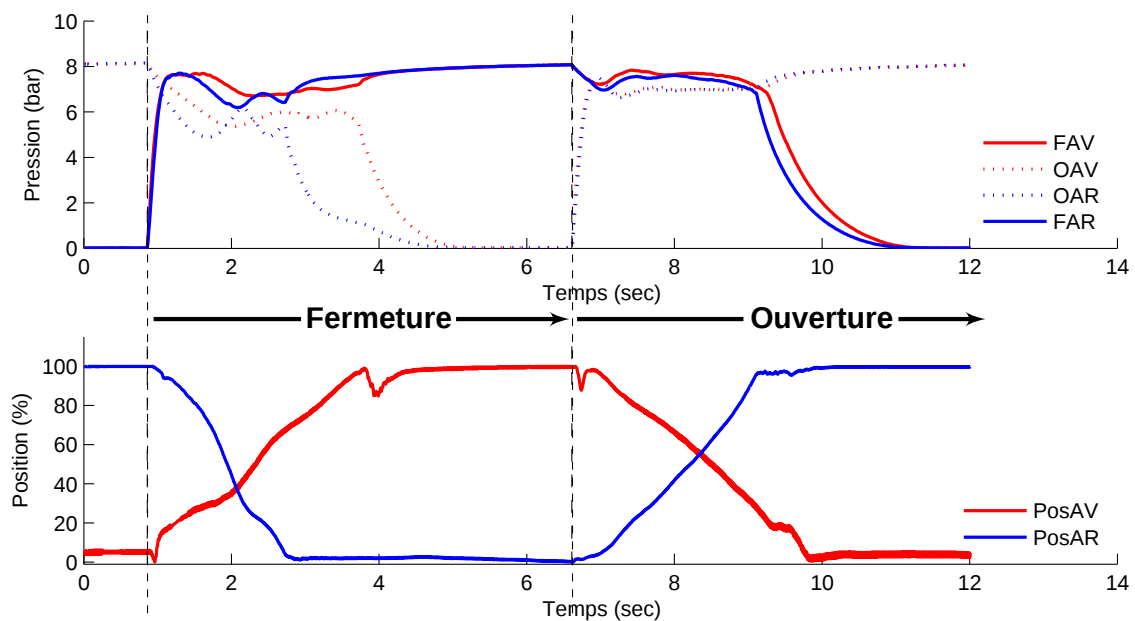


FIGURE 5.10 – Exemple d'un cycle d'ouverture / fermeture des battants de porte avant et arrière en pression (bar) et position (%), collecté le 12 juillet 2011.

On remarque sur la figure 5.10 que les deux battants s'ouvrent simultanément. En revanche lors d'une fermeture, on observe que la fermeture du battant arrière précède celle du battant avant. La durée d'ouverture/fermeture pour chaque battant est régie par des réglages mécaniques effectués par l'exploitation et peut varier d'un bus à l'autre ; cette durée dépasse rarement 5 secondes. D'autre part, la position du battant avant n'a pu être mesurée de manière satisfaisante. On observe des « creux » inattendus sur le signal de position avant lorsque la porte est fermée.

Remarque 5.1 (Choix des pressions à surveiller pendant le mouvement d'une porte)

Il est fréquent que les pressions d'ouverture OAV et OAR (sous-section 5.1.2) fluctuent lorsque la porte est ouverte. On observe le même phénomène pour les pressions de fermeture FAV et FAR lorsque la porte est fermée. Cela s'explique par le fait que de l'air comprimé est régulièrement réinjecté dans les vérins afin de maintenir les battants de porte en position ouverte/fermée. Il a donc été privilégié de traiter les courbes d'ouverture à partir des pressions de fermeture, et les courbes de fermeture à partir des pressions d'ouverture, qui s'annulent à la fin du mouvement de la porte.

5.4.2 Réglage de la structure du modèle de régression

Dans cette étude, on s'intéresse à la surveillance du battant arrière d'une porte pendant des mouvements de fermeture. Un modèle de régression est appris sur un ensemble de onze courbes bivariées (pression OAR et position AR) collectées sur des portes en bon fonctionnement. Pour cela, il est nécessaire de choisir un « bon » modèle m parmi une collection de modèles. Cette étape importante et non triviale, appelée *sélection de modèle* (sous-section 2.2.2), consiste essentiellement à résoudre le dilemme biais-variance. Pour notre modèle de régression multivarié, il y a deux hyperparamètres à régler : le nombre K de segments, et l'ordre r des polynômes de chaque composante. Deux mesures classiques à maximiser sont les critères BIC et ICL :

$$\text{BIC}(m) = \mathcal{L}(\hat{\theta}_m; \mathbf{x}, \mathbf{t}) - \frac{\nu_m}{2} \log(\bar{m}) \quad (5.43)$$

$$\text{ICL}(m) = \mathcal{L}_c(\hat{\theta}_m; \mathbf{x}, \hat{\mathbf{z}}, \mathbf{t}) - \frac{\nu_m}{2} \log(\bar{m}) \quad (5.44)$$

où $\mathcal{L}(\hat{\theta}; \mathbf{x}, \mathbf{t})$ est la log-vraisemblance du modèle de régression (équation 5.4), $\mathcal{L}_c(\hat{\theta}; \mathbf{x}, \hat{\mathbf{z}}, \mathbf{t})$ est la log-vraisemblance complétée (équation 5.5) après convergence et $\bar{m} = \sum_{i=1}^t m_i$. ν_m est la dimension du paramètre $\hat{\theta}_m$; dans le cas de notre modèle de régression avec une matrice de variance diagonale (Celeux et Govaert, 1995), on a :

$$\nu_m = \underbrace{2(K-1)}_w + \underbrace{d(r+1)K}_\beta + \underbrace{dK}_\Sigma \quad (5.45)$$

Il est à noter que d'autres critères pourraient être employés, par exemple les mesures AIC et AIC3. Ces deux critères sont des formulations pénalisées de la log-vraisemblance, tout comme BIC, et leurs valeurs sont du même ordre que le critère BIC. La figure 5.11 représente les critères BIC, ICL et le temps CPU après 100 lancements (avec initialisation

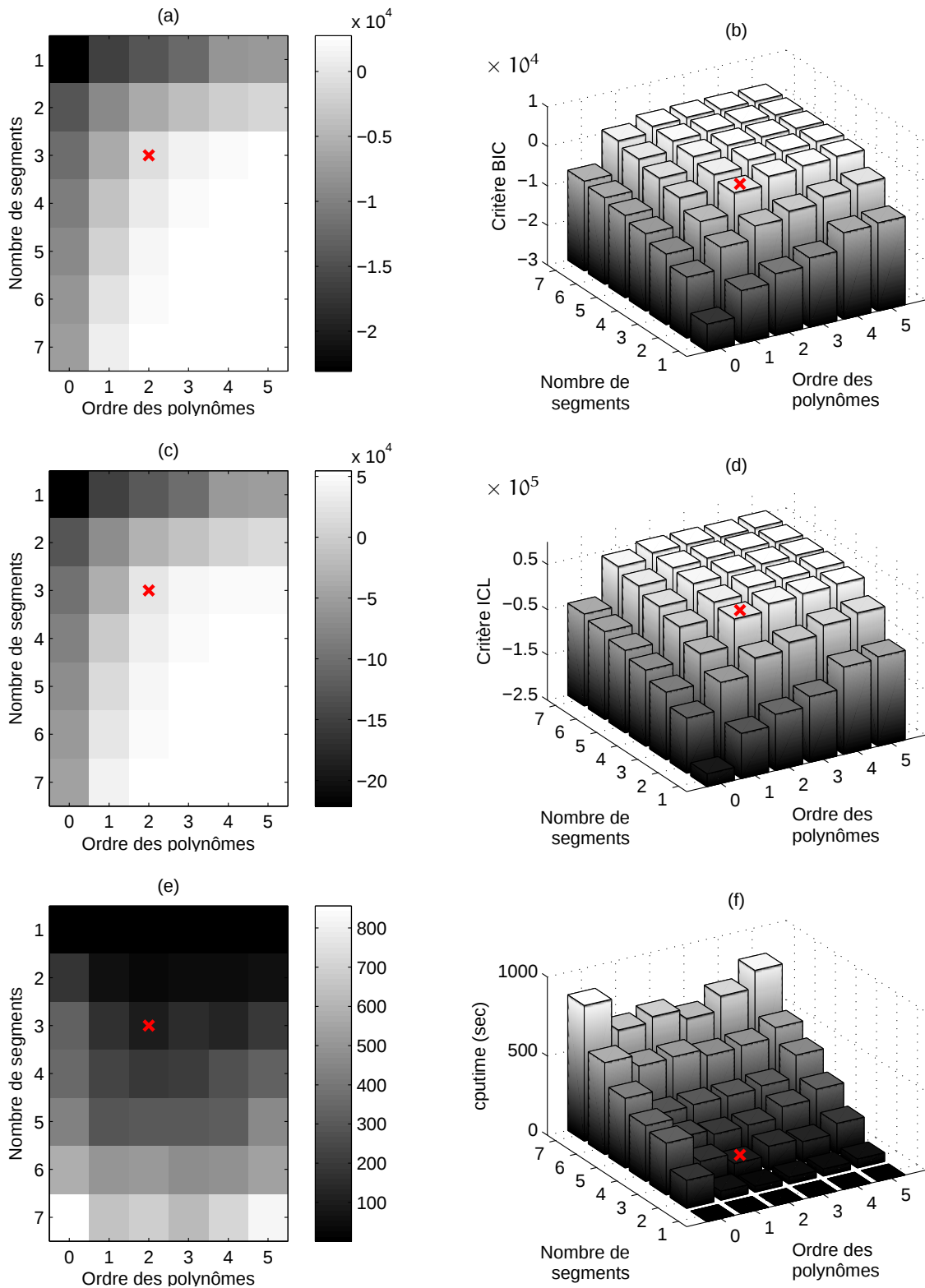


FIGURE 5.11 – Sélection du modèle de régression multivarié à partir des critères BIC (a,b) et ICL (c,d), et temps CPU (e,f).

Le modèle choisi (x) est composé de 3 segments dont chaque polynôme est d'ordre 2.

aléatoire) à partir de nos onze exemples de fermeture de porte en bon fonctionnement. Plusieurs configurations conjointes du nombre des segments $1 \leq K \leq 7$ et de l'ordre des polynômes $0 \leq r \leq 5$ ont été mises en compétition dans un cadre non supervisé. Les critères n'ont pas été maximisés car malgré les termes de pénalisations, ces critères continuent à augmenter en fonction de K et r . Le modèle à trois segments ($K = 3$) d'ordre $r = 2$ a été retenu pour notre étude. En effet, on observe une augmentation des critères puis une stabilisation à partir du modèle choisi. Ce modèle de régression est représenté par la figure 5.12. Il est à noter que le temps CPU augmente avec le nombre de segments car la dimension du paramètre du modèle est factorisable par K , le nombre de segments (équation 5.45).

La stratégie de détection illustrée par l'algorithme 5.3 a été testée sur des données simulées réalistes à partir des données collectées en juillet 2011. Le nombre de segments du modèle de régression et le degré de chaque polynôme ont été fixés à 3 et 2 suite à une étude de sélection de modèles (voir paragraphe précédent). Nous avons observé en pratique que la qualité de la détection n'est pas affectée par la taille W des fenêtres locales. La taille W a été fixée expérimentalement à 5. La prochaine sous-section présente quelques résultats sur données simulées afin de caractériser la performance du détecteur proposé.

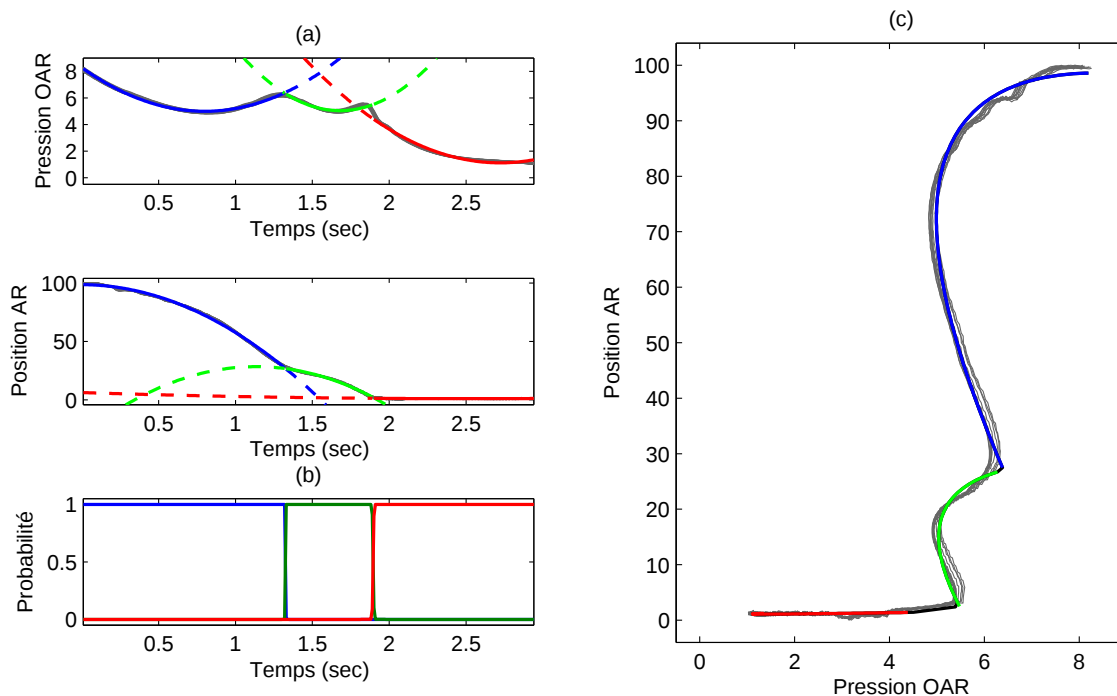


FIGURE 5.12 – Segmentation de 11 courbes (pression, position) de fermeture de portes en bon fonctionnement, collectées le 12 juillet 2011. Représentation 1D du mélange de régressions (a) et probabilités logistiques (b). Courbes bivariées et modèle RHLP multivarié dans $\mathbb{R}^2$ (c).

5.4.3 Résultats expérimentaux

Une séquence de courbes contenant un changement a été générée à partir de deux instances d'un modèle RHLP multivarié à 3 segments d'ordre 2 avec des matrices de covariance diagonales (Celeux et Govaert, 1995) : 10 000 courbes avant changement, suivies de 500 courbes après changement. Le tableau 5.4 détaille les vecteurs paramètre θ_0 et θ_1 , respectivement avant et après changement.

	w	β_1	β_2	β_3
θ_0	$\begin{pmatrix} -1922 & 2803 \\ -449 & 852 \\ 0 & 0 \end{pmatrix}$	$\begin{pmatrix} 5 & -8 \\ 8 & -40 \\ -1 & 99 \end{pmatrix}$	$\begin{pmatrix} 10 & -34 \\ 33 & -45 \\ 103 & -30 \end{pmatrix}$	$\begin{pmatrix} 5 & -26 \\ 37 & 1 \\ -4 & 6 \end{pmatrix}$
θ_1	$\begin{pmatrix} -1901 & 2716 \\ -859 & 1388 \\ 0 & 0 \end{pmatrix}$	$\begin{pmatrix} 5 & -8 \\ 8 & -41 \\ 2 & 98 \end{pmatrix}$	$\begin{pmatrix} 19 & -62 \\ 53 & -100 \\ 280 & -163 \end{pmatrix}$	$\begin{pmatrix} 1 & -8 \\ 14 & 2 \\ -9 & 39 \end{pmatrix}$
	Σ_1	Σ_2	Σ_3	
θ_0	$\begin{pmatrix} 0.01 & 0 \\ 0 & 1.00 \end{pmatrix}$	$\begin{pmatrix} 0.01 & 0 \\ 0 & 0.38 \end{pmatrix}$	$\begin{pmatrix} 0.02 & 0 \\ 0 & 0.11 \end{pmatrix}$	
θ_1	$\begin{pmatrix} 0.01 & 0 \\ 0 & 0.94 \end{pmatrix}$	$\begin{pmatrix} 0.01 & 0 \\ 0 & 0.20 \end{pmatrix}$	$\begin{pmatrix} 0.01 & 0 \\ 0 & 0.06 \end{pmatrix}$	

TABLE 5.4 – Paramètres des modèles avant et après changement utilisés pour simuler des courbes réalistes.

Différentes situations ont été testées en faisant varier un niveau de bruit donné sur l'ensemble de la séquence des courbes simulées ; chaque courbe a été générée selon le modèle de mélange, défini dans l'équation 5.3, de la manière suivante :

$$x_{ij} = \sum_{k=1}^K \pi_k(t_j; w^*) \cdot [\beta_k^* \cdot T_j + \lambda \cdot \epsilon_{ij}], \quad \forall (i, j) \in \{1, \dots, t\} \times \{1, \dots, m_i\}, \quad (5.46)$$

où $* \in \{0, 1\}$, λ est un *niveau de bruit* de valeur 5, 25 ou 50, et $\epsilon_{ij} \sim \mathcal{N}(0, \Sigma_k^*)$ est un bruit blanc de dimension d . La stratégie séquentielle, décrite dans le pseudo-code 5.3, a été évaluée sur 7 contraintes de type ARL avec un taux de fausses alarmes théorique : $\alpha = \{1/2, 1/10, 1/20, 1/100, 1/200, 1/1000, 1/2000\}$.

Afin d'évaluer notre stratégie de détection, on utilise deux critères liés aux erreurs de première et de seconde espèce dans un cadre séquentiel. Les deux métriques de performances sont définies comme :

- Le *taux de fausses alarmes* (False Alarm Rate, ou FAR)

$$\text{FAR}(t_a) = 1/\mathbb{E}(t_a | t_a < t_c) \quad (5.47)$$

- Le *délai moyen de détection* (Average of Detection Delay, ou ADD) :

$$\text{ADD}(t_a) = \mathbb{E}(t_a - t_c | t_a \geq t_c) \quad (5.48)$$

où t_a est le temps d'alarme et t_c est le vrai temps de changement. Il est à noter que les deux critères sont calculés empiriquement (en moyennant les espérances). La figure 5.13 illustre quelques caractéristiques du détecteur en utilisant les deux métriques définies précédemment. On retrouve notamment ces modes d'évaluation dans (Tartakovsky et al., 2005).

Les trois sous-figures (a,c,e) de la figure 5.13 présentent les graphiques $-\log(\hat{\alpha})$ vs. $-\log(\alpha)$ pour les trois niveaux de bruits considérés, où $\hat{\alpha}$ est le taux de fausses alarmes expérimental calculé selon l'équation 5.47. Les pointillés désignent une estimation idéale des seuils sous la contrainte du taux de fausses alarmes, avec les mêmes valeurs estimées et théoriques des taux de fausses alarmes : $\hat{\alpha} = \alpha$. Il est à noter que la plage de valeurs de $-\log(\alpha)$ allant de 0.7 à 7.6 correspond à un temps espéré entre deux fausses alarmes allant de 2 à 2 000 fenêtres de courbes. Pour chaque niveau de bruit, le taux empirique de fausses alarmes $\hat{\alpha}$ est très corrélé au taux théorique α et on observe une certaine stabilité entre ces deux quantités. Le critère du $\hat{\alpha}$ est calculé en prenant l'inverse de la durée avant la première fausse alarme (sous l'hypothèse $\mathcal{H}_0$), comme décrit dans l'équation 5.47. On suppose alors que ce temps moyen est le même que le temps moyen entre deux fausses alarmes. Dans cette configuration, le critère FAR est mieux estimé lorsque la contrainte sur le taux théorique de fausses alarmes concerne une grande valeur de α (petites valeurs de $-\log(\alpha)$). Les seuils ont été correctement estimés par la procédure décrite dans la sous-section 5.3.4 car les courbes estimées et théoriques (celles en pointillées) sont proches.

Les trois sous-figures (b,d,f) de la figure 5.13 présentent les graphiques ADD vs. $-\log(\hat{\alpha})$. Ce graphique peut être vu comme une version « séquentielle » de la courbe ROC. Il y a un continuum de valeurs ADD qui correspondent aux nombres de courbes suivant $p(\cdot; \theta_1)$ avant détection. On rappelle que la règle de détection optimale recherchée doit garantir un faible délai de détection tout en maintenant un taux faible de fausses alarmes (Gustafsson et Palmqvist, 2000). Le détecteur idéal atteint le coin inférieur droit des sous-figures (b,d,f). On observe que le critère ADD est toujours nul lorsque les courbes sont faiblement bruitées. On l'explique par la différence notable des distributions $p(\cdot; \theta_0)$ et $p(\cdot; \theta_1)$. Le délai peut être important lorsque le bruit est important ; dans ce cas, les distributions se chevauchent et le point de changement peut-être difficilement détecté.

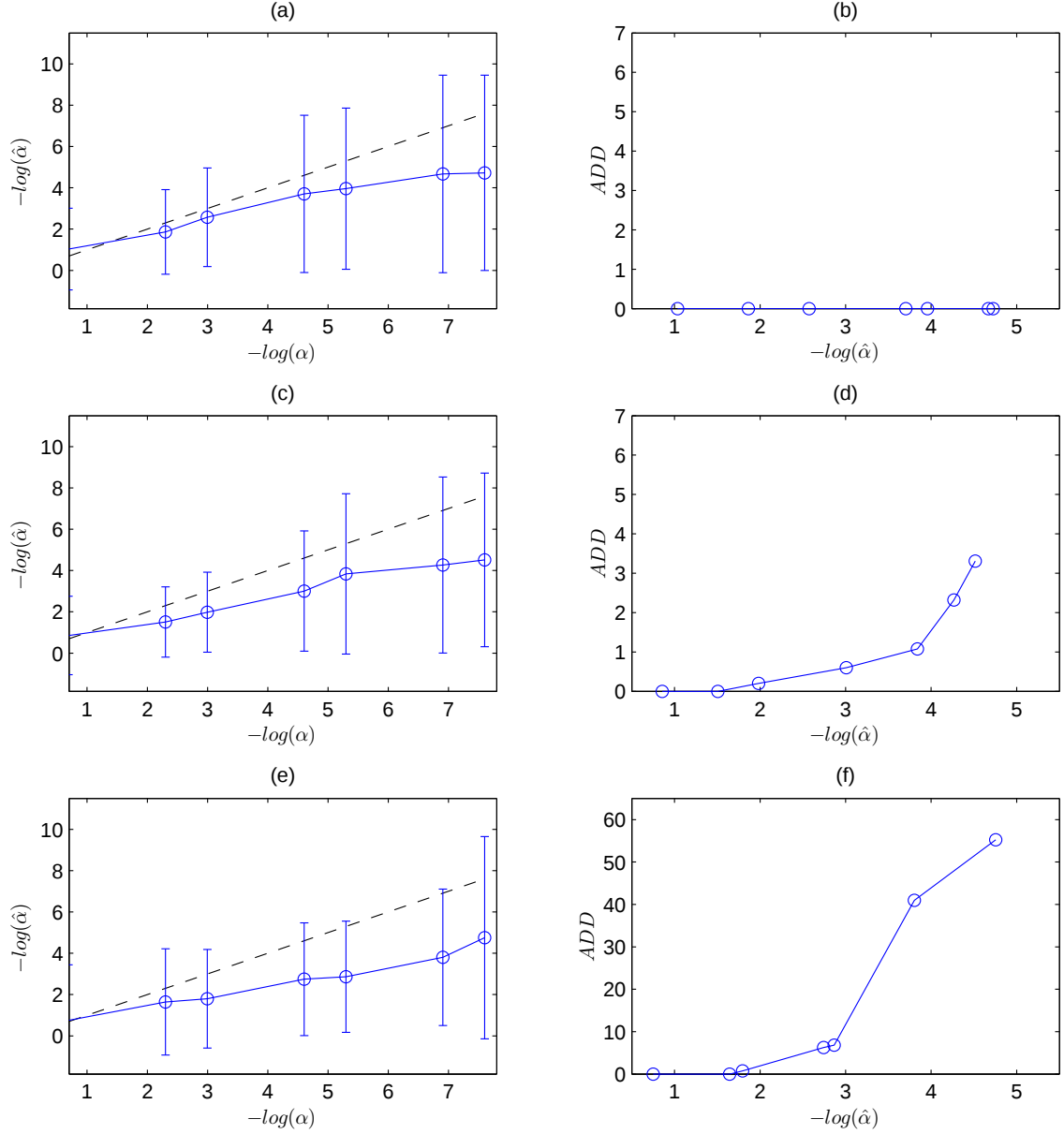


FIGURE 5.13 – Évaluation graphique du détecteur proposé en fonction des taux de fausses alarmes et des délais moyens à la détection. Les figures (a,c,e) représentent les taux de fausses alarmes expérimentaux $\hat{\alpha}$ en fonction des taux théoriques α . Et les figures (b,d,f) représentent les délais moyens de détection ADD en fonction des taux de fausses alarmes expérimentaux $\hat{\alpha}$.

5.5 Expérimentations sur données réelles

5.5.1 Acquisition des données de fonctionnement

Dans le cadre du UseCase de Brunoy du projet [EBSF](#), une flotte d'autobus Citelis (IRISBUS-IVECO), standards et articulés, a été équipée d'une architecture télématique innovante (connectique, capteurs additionnels, calculateurs embarqués, serveur backoffice,...) décrite dans la sous-section 1.3.2. Une seule porte est instrumentée sur chaque véhicule. Il s'agit de la porte arrière dans le cas d'un Citelis standard, comme le présente la figure 5.14, et de la porte milieu dans le cas d'un bus articulé.

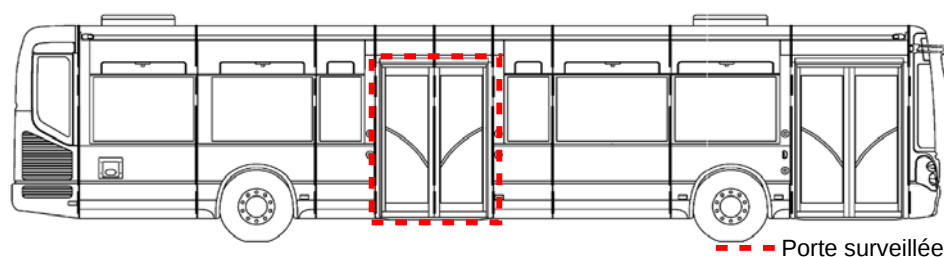


FIGURE 5.14 – Autobus standard Citelis et porte surveillée.

La sous-section 5.1.2 décrit l'instrumentation, embarquée dans les véhicules, qui est dédiée à notre étude et les mesures qui sont collectées. Un serveur backoffice a été mis en place au dépôt de bus afin de décharger et centraliser les données de porte. Après une tournée et le retour d'un véhicule au dépôt, les données de fonctionnement des portes sont déchargées par WiFi sur le serveur backoffice. Ces données sont enregistrées sous la forme de fichiers propriétaires puis ces fichiers binaires sont mis à disposition par DIGIGROUP, partenaire du projet et détenteur des calculateurs embarqués VIDAC.

5.5.2 Description des données réelles collectées

Des données sont collectées par des capteurs embarqués pendant le mouvement des portes pneumatiques. Sur la porte instrumentée de chaque véhicule (voir figure 5.3), on mesure essentiellement les quatre pressions des vérins, les deux positions instantanées des battants et l'état d'ouverture/fermeture de la porte, en fonction du temps. Le tableau 5.5 présente l'ensemble des variables mises à disposition pour la surveillance du système porte. Avec une fréquence d'échantillonnage de 100 Hz, le calculateur embarqué VIDAC enregistre ces mesures pendant 20 secondes à partir du déclenchement de la commande d'ouverture de la porte surveillée, via le bouton situé sur le tableau de bord. En pratique, on observe un nombre variable d'ouvertures et fermetures selon l'échantillon dans l'intervalle des 20 secondes, appelé *trigger*.

Variable	Unité	Utilisée pour la surveillance du battant mouvement			
		avant	arrière	ouverture	fermeture
Position AV	%	X		X	X
Position AR	%		X	X	X
Pression FAV__	10 ³ Pa	X		X	
Pression OAV...	10 ³ Pa	X			X
Pression OAR...	10 ³ Pa		X		X
Pression FAR__	10 ³ Pa		X	X	
Statut de la porte	{0, 1}	X	X	X	X
Temps	sec	X	X	X	X

TABLE 5.5 – Description des données collectées pour la surveillance des portes.
 Les quatre pressions sont mesurées en analogique par le calculateur VIDAC.
 Les deux positions de chaque battant et le statut (vert) sont transmis
 au VIDAC par le CAN constructeur.

Les triggers de données sont enregistrés les uns à la suite des autres dans des fichiers qui sont déchargés régulièrement au dépôt de bus (voir la sous-section précédente). La figure 5.15(a) présente un exemple de flux de données collectées en mai 2012, avec un facteur de 200 pour le statut (vert). L'extraction des courbes d'ouverture/fermeture de ces flux n'est pas une tâche triviale sans indicateur de début et de fin pour chaque mouvement. Il est à noter que l'acquisition de ces signaux étant nouvelle, les équipes opérationnelles n'ont pu apporter leur expertise pour le traitement de ces données. En outre, la dynamique d'un mouvement peut différer selon les conditions d'exploitation : sa durée peut varier pour un même bus et un mouvement d'ouverture/fermeture peut être interrompu à tout moment par un obstacle dans la course de la porte.

Chaque trigger est déterminé à partir du temps puis on extrait les mouvements à partir du signal binaire Statut de la porte qui correspond à l'information affichée par la led du bouton d'ouverture/fermeture sur le tableau de bord. Le début d'une ouverture est facilement identifiable par ce signal. Ce n'est pas le cas de la fin d'une ouverture ni d'une fermeture. En pratique, on identifie ces fenêtres par recherche locale d'un point de rupture sur les pressions. Après cette étape, on constitue une séquence d'ouvertures et fermetures dont les durées varient généralement entre 2.5 et 5 secondes. La figure 5.15(b) illustre un cycle d'ouverture/fermeture extrait du flux de données mesuré en mai 2012 sur un bus standard. On remarque notamment que les signaux de position transmis par le CAN constructeur ressemblent assez peu à ceux mesurés en juillet 2011 sur les bornes des potentiomètres (figure 5.10). Les mesures de position sont discrétisées puis échangées très grossièrement entre plusieurs ECU avant d'être transmises au VIDAC sous la forme de signaux codés sur très peu de bits significatifs (6 bits) ; la prochaine sous-section décrit comment nous affranchir des difficultés liées au traitement de la donnée de position.

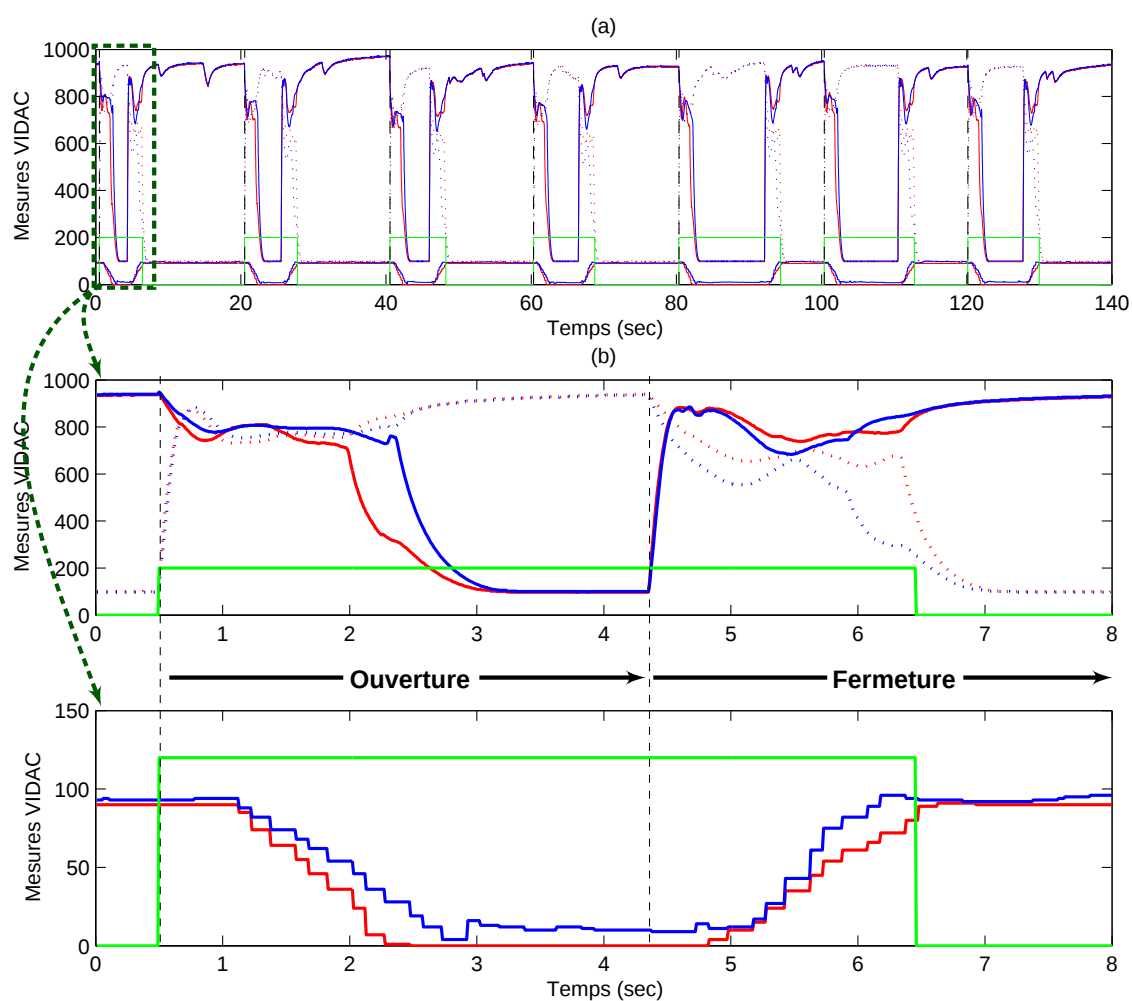


FIGURE 5.15 – Échantillon d'un flux de données (a) et zoom sur un cycle d'ouverture / fermeture (b) collecté sur la porte instrumentée d'un bus standard en mai 2012.
Les légendes sont décrites dans le tableau 5.5.

Démarche adoptée pour la surveillance des portes d'un autobus

On présente ici les résultats obtenus à partir de données réelles sur un seul autobus pendant huit mois et demi : du 30/01/2012 au 12/10/2012. Ce véhicule a été choisi pour nos travaux car il en a été extrait l'échantillon le plus important (près de 5 000 triggers) d'une part, et d'autre part, le système porte de ce bus a subi plusieurs défaillances jugées pertinentes pour notre problème. Les résultats de diagnostic sont de nature plutôt qualitative en terme de détection, de localisation et d'interprétation des défaillances retournées par notre stratégie de surveillance. Dans les prochaines sous-sections, nous commencerons par présenter l'intérêt d'ajouter un a priori sur la structure de nos signaux (ajout d'une semi-supervision au modèle de régression). Puis, nous présenterons les résultats obtenus par la stratégie proposée pour la détection séquentielle et la localisation des défauts (algorithme 5.5) appliquée à une séquence d'ouvertures.

5.5.3 A priori physique et sélection du modèle

Pour rappel, l'ajout d'une semi-supervision pour l'apprentissage au modèle de régression améliore la qualité de l'estimation et accélère la convergence de l'algorithme (sous-section 5.2.3). On cite notamment les travaux de Côme (2009, chapitre 4) qui illustrent l'utilité d'un étiquetage pour un problème de séparation de sources. Cette sous-section a pour objet de motiver notre choix d'un modèle de régression spécifique adapté à nos données fonctionnelles.

Commandé au niveau du tableau de bord, le mouvement de chaque battant est marqué par trois instants de changement illustrés par des cercles rouges dans la figure 5.17(a), et donc décomposables en quatre segments. Le premier et troisième points correspondent respectivement au début et à la fin du mouvement d'un battant. Dans le cas d'une ouverture, le second point marque l'instant où le piston appuie sur le résidu d'air de la chambre 2 (chambre 1 pour une fermeture) afin de vider entièrement cette partie du vérin. Nous avons exploité cet a priori physique lié à la cinétique de la porte en choisissant un modèle de régression avec $K = 4$.

Après multiplexage, les mesures de position sont collectées sous la forme de signaux très peu précis comme illustrées par la figure 5.16(a, bas). Plusieurs méthodes de débruitage ont été testées sur ces signaux (splines, filtre moyen,...). Mais aucun prétraitement n'a pu être jugé satisfaisant car les signaux sources sont inconnus et donc, aucune mesure d'erreur n'a pu être employée afin de favoriser une méthode ou une autre. Nous avons donc préféré conserver les signaux bruts de position de manière à préserver le plus d'information possible. La forme « en escalier » de ces signaux peut mettre en échec le modèle de régression car chaque marche est potentiellement interprétée comme une transition entre composantes. Afin d'assurer une stabilité du modèle et d'améliorer la qualité de nos estimations, l'extension semi-supervisée du modèle de régression, décrite dans la sous-section 5.2.3, a été mise en œuvre. L'étiquetage partiel de la localisation des segments a permis d'isoler les quatre phases et d'éviter l'inversion des composantes (problème classique en clustering), tout en maintenant une certaine variabilité des transitions entre segments. La figure 5.17(a) illustre les trois zones d'incertitude (grisées) pour vingt courbes d'ouverture supposées en bon fonctionnement. En pratique, ces trois zones sont formées sur les distributions empiriques des trois instants de changement identifiés par recherche locale sur chacune des courbes disponibles.

Suite à la caractérisation du nombre de segments et l'étiquetage partiel du modèle de régression, nous avons sélectionné l'ordre du modèle à partir des critères BIC et ICL (équations 5.43 et 5.44). Une nouvelle étude de sélection de modèle est nécessaire car les positions diffèrent de celles mesurées en juillet 2011. Un ordre 2 a été choisi car les critères n'augmentent plus beaucoup à partir de cet ordre (figure 5.16(b)). A titre d'illustration, un modèle appris en mode semi-supervisé est représenté en figure 5.17, sur un ensemble de vingt courbes d'ouverture en bon fonctionnement.

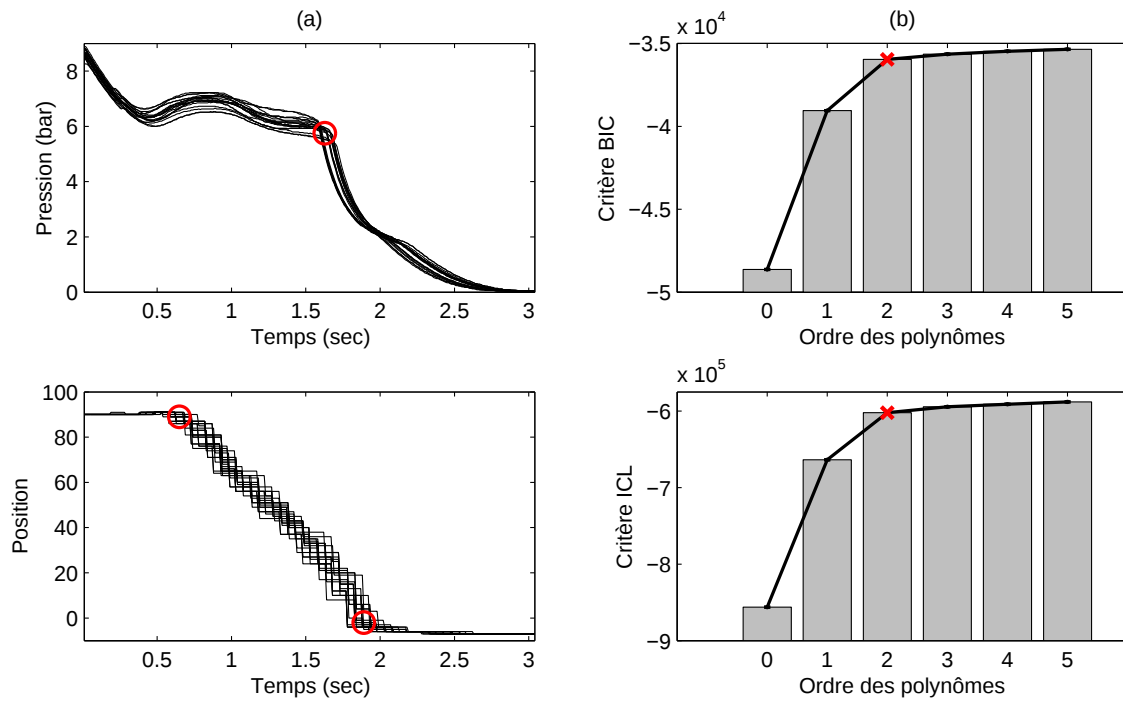


FIGURE 5.16 – Représentation de vingt courbes d’ouverture de portes avec trois points de changements (a) marqués par des cercles (o). Sélection d’un ordre deux (x) pour le modèle de régression (b).

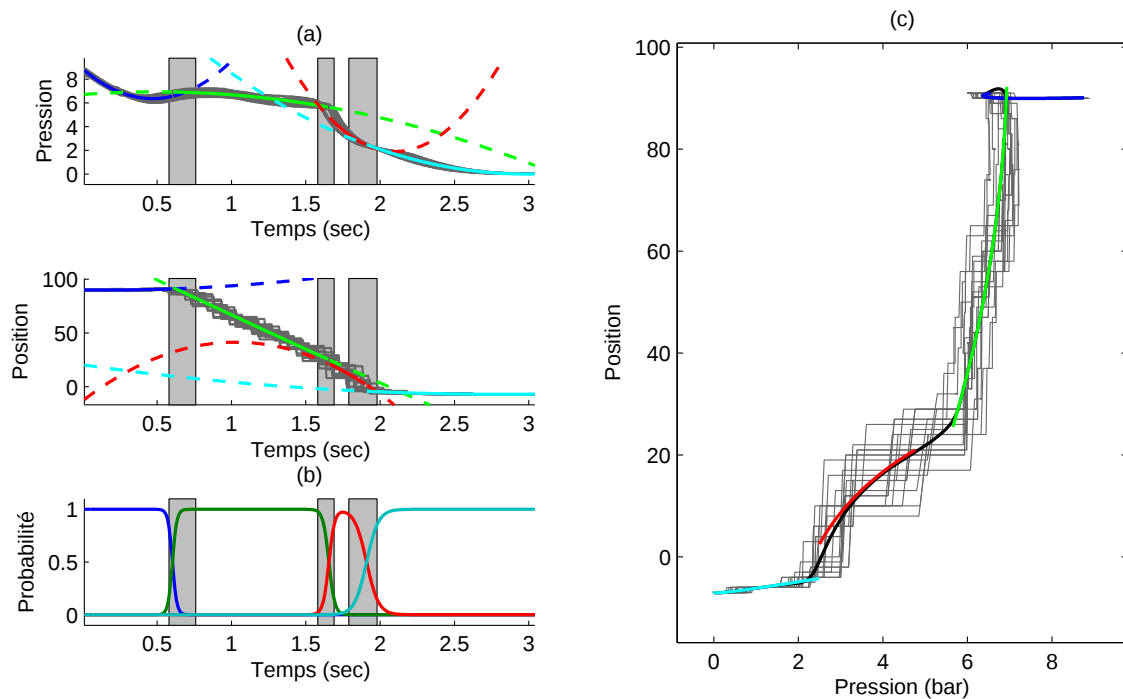


FIGURE 5.17 – Segmentation de 20 courbes (pression, position) d’ouverture de portes en bon fonctionnement, collectées le 30 janvier 2012. Représentation 1D du modèle de régression (a) et fonction logistique (b). Courbes bivariées et courbe moyenne du modèle MRHLP dans $\mathbb{R}^2$ (c).

5.5.4 Interprétation des défauts réels

Le diagnostic d'un système complexe nécessite de détecter une défaillance puis de l'identifier de manière à opérer une action appropriée (voir la sous-section 2.1.2). La stratégie proposée (pseudo-code 5.5) a pour but de détecter une anomalie et de localiser le défaut sur la signature temporelle. Dans le cadre de nos travaux, il a été possible de dégrader physiquement les systèmes porte de manière temporaire. Dans cette sous-section, nous présentons succinctement une campagne de dégradations provoquées au dépôt en juillet 2012.

La caractérisation formelle des défauts a été difficilement réalisable uniquement sur la base d'un suivi des interventions notifiées dans WINATEL, l'outil de GMAO de l'exploitation (voir sous-section 1.3.4). Comme pour le système de freinage, une liste de plusieurs modes de défaillances a pu être établie sur la base d'échanges avec l'équipe de maintenance. L'apport de connaissance experte, motivé par l'expérience métier, a été déterminant pour conduire ce travail. Contrairement au système de freinage, il a été possible de reproduire trois types de défauts (provisaires et réparables) occasionnellement observables sur des systèmes de porte pneumatique. Une campagne d'essais a été réalisée pendant l'été 2012 afin de collecter un échantillon de quelques courbes représentatives de chacun des modes de défaillance suivants :

- *Blocage de porte* : des obstacles ont été positionnés dans la course de chaque battant (fréquent en exploitation). Comparés aux signaux normaux, les deux derniers segments (après blocage) diffèrent et semblent ne former qu'une seule composante (voir figure 5.18). Dans nos essais, la fin des signaux est variable car la position de l'obstacle a été difficilement maintenue immobile sous l'effet de la force des battants ;
- *Déréglage de potentiomètre* : le calibrage des potentiomètres peut être altéré avec le mouvement répété des battants. Les signaux sont toujours décomposables en quatre phases mais l'étendue des positions est modifiée (valeurs translatées) ;
- *Fuite d'air* : le réservoir d'air qui alimente les vérins est purgé. Les signaux sont toujours décomposables en quatre phases mais les pressions de départ diminuent à mesure que le réservoir se vide.

Nos observations sont similaires dans le cas d'un mouvement d'ouverture/fermeture et/ou du battant avant/arrière. Les trois situations de dégradation décrites précédemment sont illustrées par les figures 5.18 et 5.19, pendant l'ouverture d'un battant avant.

La prochaine sous-section présente les résultats obtenus en terme de détection d'anomalies et d'identification de défauts sur des signaux réels mesurés sur un système porte pneumatique. La connaissance préalable des signatures particulières de défaut par rapport à celles en bon fonctionnement nous a permis d'identifier les défauts après détection.

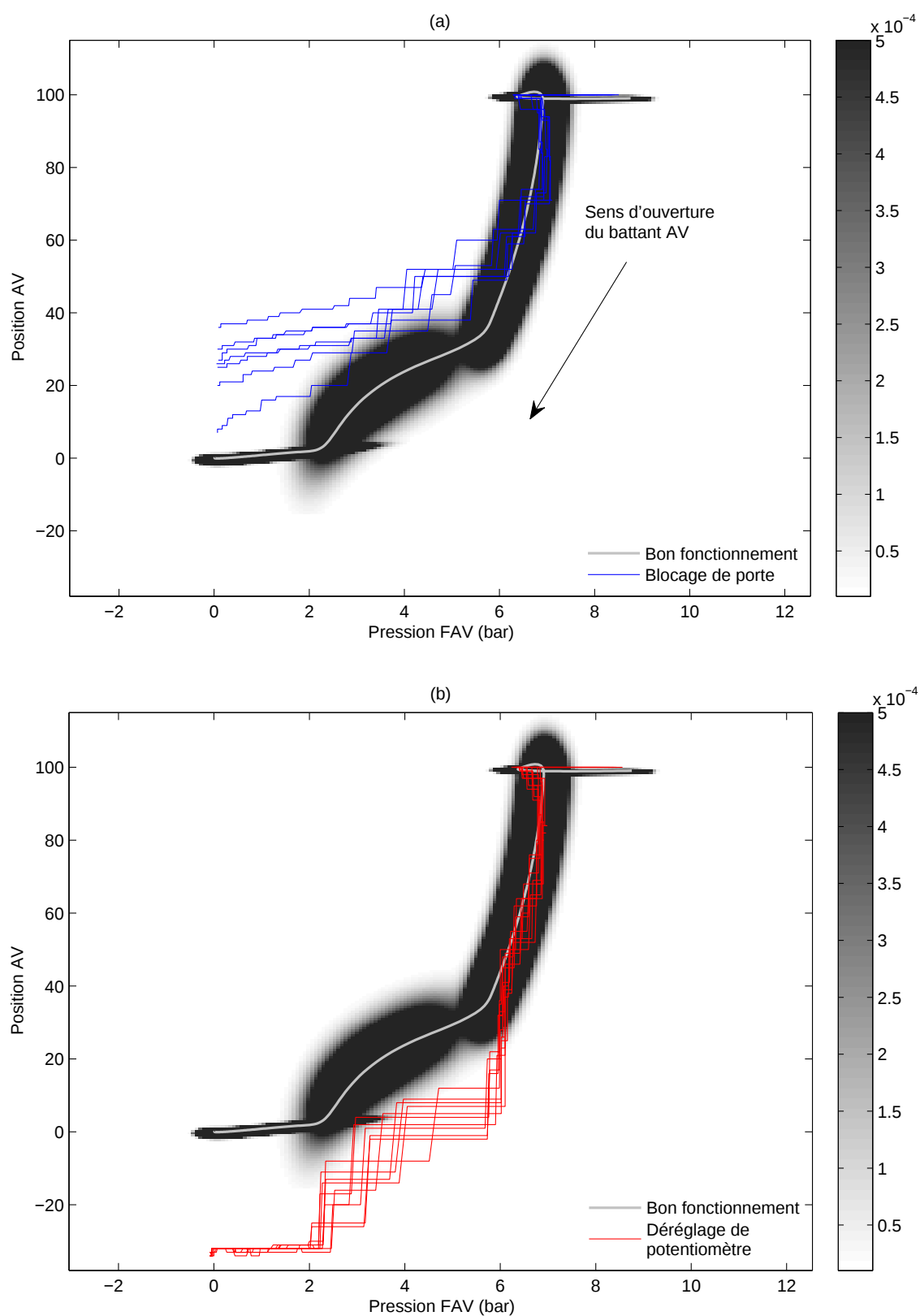


FIGURE 5.18 – Densité de signatures en bon fonctionnement (cf. figure 5.17) et la courbe moyenne du modèle MRHLP. Quelques dégradations provoquées : blocage dans la course du battant avant (a) et déréglage du potentiomètre avant (b).

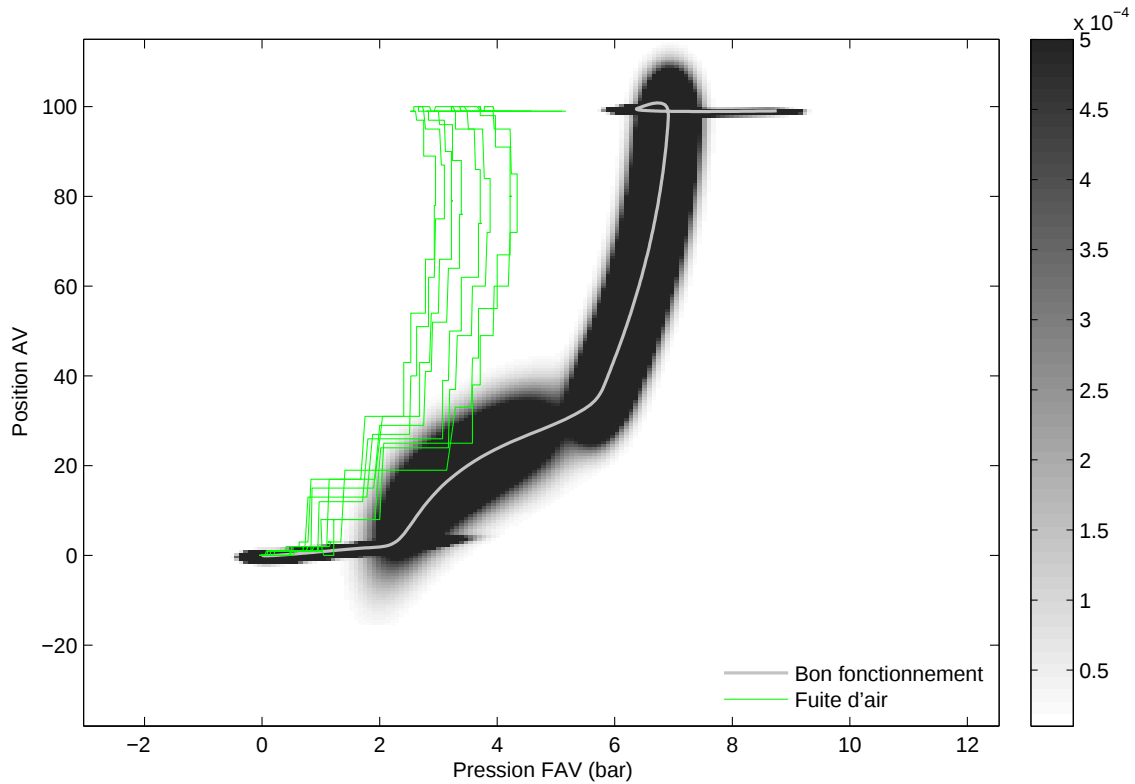


FIGURE 5.19 – Densité de signatures en bon fonctionnement (cf. figure 5.17). Quelques signatures de fuite d’air des servitudes alimentant le vérin avant.

5.5.5 Résultats en terme de détection, localisation et interprétation

Dans cette sous-section, la procédure de détection et de localisation de défauts (algorithme 5.4) est testée sur un jeu de données réelles, mesurées sur un système de portes pneumatiques entre janvier et octobre 2012. La stratégie est évaluée qualitativement sur la séquence de courbes décrite dans la sous-section 5.5.2. L’objectif est de surveiller le fonctionnement de la porte instrumentée à partir de la procédure décrite par l’algorithme 5.5 ; le détecteur est réinitialisé lorsqu’un changement apparaît dans la distribution des courbes. Cette densité de probabilité est représentée par un mélange de sous-modèles polynomiaux spécifiques à nos données fonctionnelles décrit dans la section 5.2 ; la modélisation intègre plusieurs a priori motivés par l’application dont les choix sont présentés dans la sous-section 5.5.3.

La figure 5.20 illustre le test proposé pour la détection d’anomalies sur une séquence de données fonctionnelles, appliqué à une séquence de 3 108 ouvertures mesurées entre janvier et octobre 2012. La séquence est ordonnée selon l’ordre d’acquisition des courbes d’ouverture et le nombre de cycles pouvant varier selon les jours d’exploitation, le pas du vecteur de date n’est pas constant. Le détecteur a été paramétré avec une mémoire M

de taille 20, ce qui correspond au nombre moyen de cycles sur une journée. La taille W de la fenêtre locale a été expérimentalement fixée à 5. La procédure séquentielle permet de prendre une décision sur l'état du système dès la réception d'une nouvelle courbe. Rappelons que dans l'étude expérimentale présentée en section 5.4, la séquence de courbes simulées ne comportait qu'un seul changement. Cette étude suppose également que le temps moyen d'alarme (avant un premier changement) correspond au temps moyen entre deux fausses alarmes. Après chaque nouvelle détection, la méthode est réinitialisée afin d'apprendre un nouveau mode de bon fonctionnement sous l'hypothèse $\mathcal{H}_0$ et un nouveau seuil d'alarme est estimé comme quantile à partir d'une loi lognormale à deux paramètres apprise sur la mémoire des courbes sous $\mathcal{H}_0$. Ces réinitialisations successives sont observables dans les figures 5.20 et 5.21 par les « espaces vides » laissés après chaque détection. La procédure de détection, formulée dans un cadre général, permet ainsi de détecter séquentiellement plusieurs changements. Dans cette étude, les seuils sont estimés de manière à assurer un taux de fausses alarmes théoriques de 10^{-4} .

La figure 5.20 met en évidence 35 détections signalées entre janvier et octobre 2012 (marquées par un symbole $\circ$), toutes suivies d'une réinitialisation du détecteur. Ces événements correspondent à des blocages de portes ou des dérèglages de potentiomètres (sous-section 5.5.4), et on note la présence de 4 fausses alarmes (non explicables a priori). Trois dérèglages ont pu être observés : changement de potentiomètre par l'équipe de maintenance le 5 mars, un dérèglement faible et un autre important que nous avons provoqués les 23 et 27 juillet. Les statistiques de détection et de localisation sont représentés par des couleurs bleues/vertes autour de ces changements, respectivement dans les figures 5.20 et 5.21. Ce choix de représentation permet de distinguer facilement les différentes translations de position après dérèglement, dans la figure 5.21. Ces trois dérèglages ont bien été détectés à chaque fois par la statistique de détection et au moins une statistique de localisation par segment. On remarque que les statistiques de détection au 5 mars et au 23 juillet sont de moindre amplitude, ce qui traduit un dérèglement faible. En revanche, la valeur très élevée de la statistique du 27 juillet révèle un dérèglement important, et ce constat est également visible sur les statistiques de localisation. D'autre part, on observe que certains blocages importants ont été signalés par le détecteur. Ce résultat est intéressant dans le sens où la procédure proposée permet à la fois de détecter des changements structurels comme un dérèglement, et des changements sporadiques (outliers) comme les blocages de porte, dans une séquence de courbes.

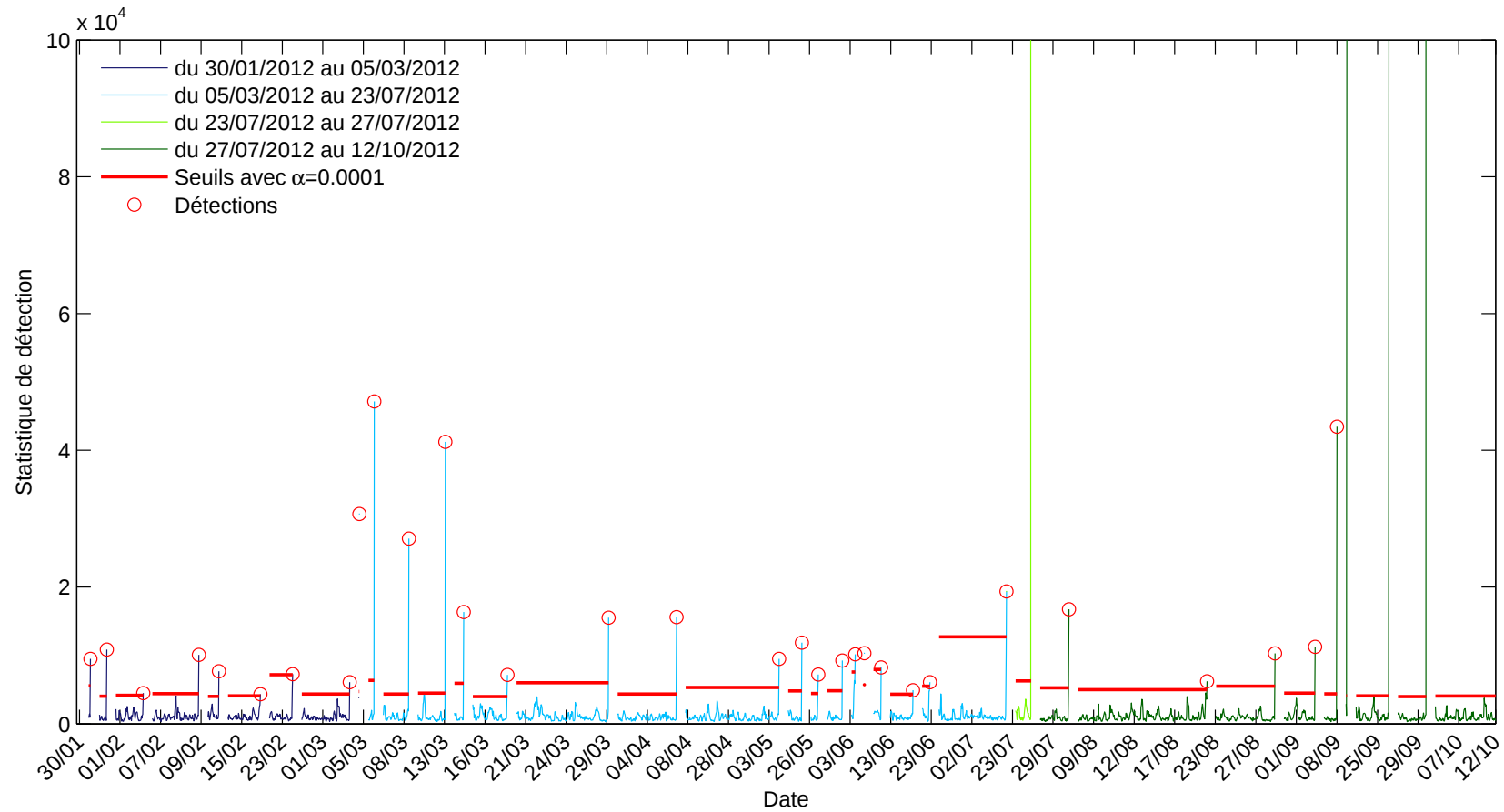


FIGURE 5.20 – Statistique du test séquentiel proposé pour la détection de défauts et appliqué à des données fonctionnelles mesurées sur des portes pneumatiques. Entre janvier et octobre 2012, 35 anomalies ont été détectées avec un taux de fausses alarmes théorique fixé à 10^{-4} .

La figure 5.21 présente les statistiques de localisation par segment (équation 5.41) en fonction du temps. Pour l'ensemble des courbes de la fenêtre courante, on rappelle que cette statistique est rapportée à chaque segment par la règle du MAP (voir équation 5.40). Dans nos expérimentations, il a été observé que les statistiques de localisation sont globalement positives et montrent une certaine sensibilité aux transitions entre composantes du modèle de mélange. Lorsqu'un défaut est détecté par la procédure de détection, ces mesures nous offrent une information complémentaire de fonctionnement du système sur chaque segment. En d'autres termes, le test de détection est un outil macroscopique utile pour le suivi de fonctionnement global du système et l'étape de localisation permet un classement de la nouvelle forme détectée.

Il est à noter qu'une estimation biaisée (à partir de données atypiques fortement bruitées) des seuils de détection a été la cause de deux fausses alarmes (sur quatre) signalées par le détecteur. A chaque réinitialisation du détecteur, on réestime des seuils de détection et de localisation (sous-section 5.3.4) ainsi qu'une segmentation partielle du modèle de régression (sous-section 5.5.3) à partir d'un ensemble de courbes, appelé *mémoire*, supposées vérifier l'hypothèse $\mathcal{H}_0$. En pratique, des blocages peuvent exister dans cet ensemble de courbes supposées en bon fonctionnement, et cela affecte l'estimation des seuils ainsi que celle de la segmentation partielle. Les blocages de porte impactant plus particulièrement les segments 2 et 3, on observe dans la figure 5.21 que les seuils de localisation de ces segments peuvent varier de manière relativement importante, comparés aux seuils des segments 1 et 4.

La figure 5.22 permet de visualiser la séquence des 3108 courbes d'ouvertures de porte. Les courbes de couleurs bleues et vertes représentent les courbes collectées entre les trois dérèglages mentionnés précédemment. Ces dérèglages se traduisent par des translations des valeurs absolues de position. On remarque notamment le dérèglement important du 27 juillet illustré par le passage des courbes vertes claires aux courbes vertes foncées. Les courbes en rouge représentent les blocages détectés par la procédure proposée. Cette visualisation de la séquence de courbes permet d'illustrer facilement les régions denses de l'espace des courbes et celles moins denses, où appartiennent les blocages en particulier.

Enfin, la procédure s'exécute en un temps raisonnable. On a observé une durée moyenne de 5.77 secondes en temps CPU à chaque itération de la procédure séquentielle entre deux réinitialisations. Les durées de réinitialisation du détecteur (avec estimation de θ_0 et des seuils) sont de 175 secondes en moyenne. Le coût de calcul quasi constant et relativement faible nous permet d'envisager une implémentation de cette stratégie dans des calculateurs embarqués et intégrée à des applications temps réel.

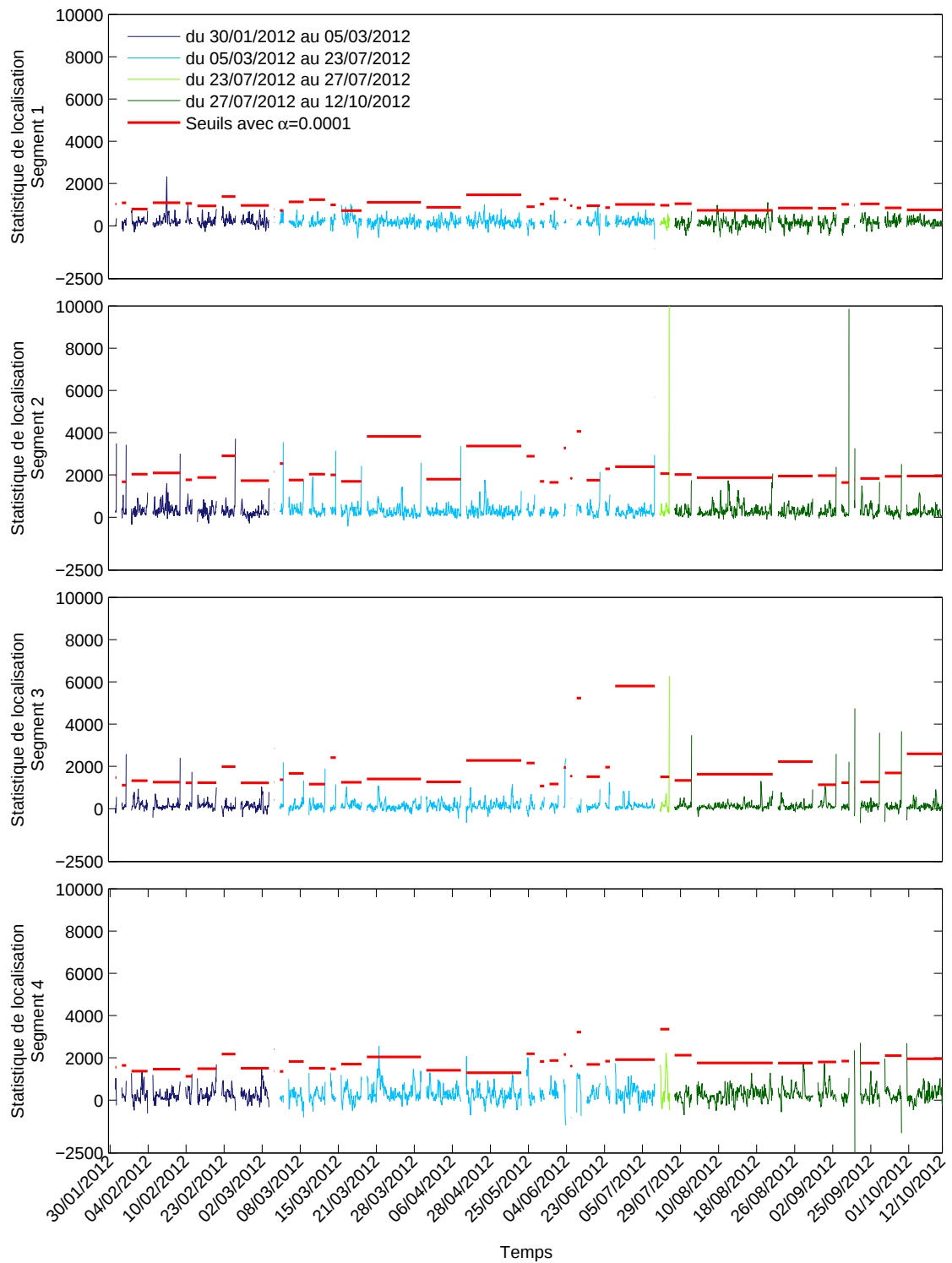


FIGURE 5.21 – Statistiques de localisation des défauts par segment calculées à partir d’une séquence de courbes mesurées sur des portes pneumatiques entre janvier et octobre 2012.

Les localisations par segment nous ont permis d’aider à l’identification de défauts détectés.

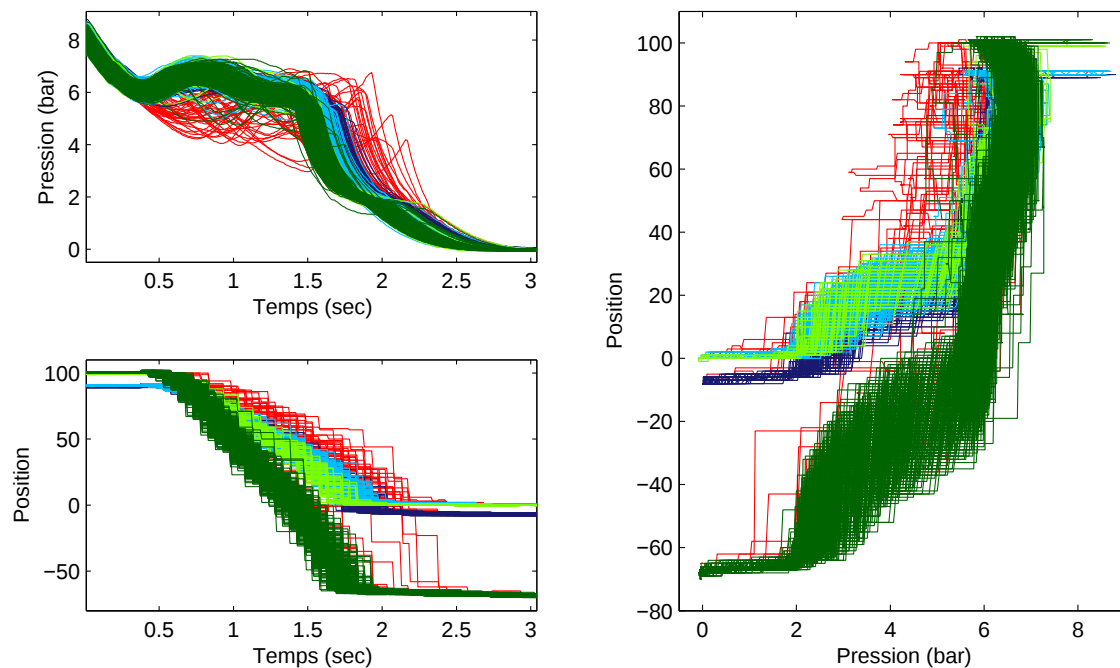


FIGURE 5.22 – Représentation des courbes mesurées entre janvier et octobre 2012.
 Les courbes entre dérèglages avérés sont de couleurs bleues et vertes (voir figure 5.21).
 Les blocages de portes détectés ont été représentés en rouge.

5.6 Conclusion

Ce chapitre présente une nouvelle procédure séquentielle de détection et localisation dans un flux de courbes multivariées (profils multidimensionnels). Cet outil générique est motivé par le suivi de fonctionnement de portes pneumatiques dans un autobus. La méthode de détection de changement est formulée sous la forme de tests séquentiels d'hypothèses qui s'appuient sur l'estimation de la distribution courante des courbes collectées. Cette densité de probabilité est représentée par un mélange de sous-modèles de régression qui constitue l'extension multivariée du modèle de régression à processus latent ([Chamroukhi, 2010](#)).

Le test de détection est une procédure de type GLR qui nécessite de mettre à jour les paramètres des distributions à chaque étape. L'estimation s'effectue par le principe du maximum de vraisemblance (MV) dont une extension a été présentée en mode semi-supervisé. Cette approche présente plusieurs avantages et notamment celui de pouvoir améliorer la segmentation des courbes et ainsi, de faciliter l'interprétation des segments identifiés à des phases successives sur le système physique. En outre, des tests supplémentaires opérés segment par segment permettent l'identification des nouveaux modes de fonctionnement après détection.

Il est à noter que la stratégie proposée intègre une double notion de temps sans avoir pris d'hypothèses de type markoviennes. Chaque courbe est une suite d'observations dont les transitions dynamiques entre régimes sont régies par une loi logistique. Et d'autre part, le flux des courbes est traité séquentiellement.

Le détecteur a été validé sur une séquence de courbes générées à partir de signaux mesurés en juillet 2011. L'évaluation de cette stratégie séquentielle a été menée à partir des courbes $-\log(\text{FAR})$ vs. $-\log(\alpha)$ et ADD vs. $-\log(\text{FAR})$, et montre notamment la capacité de la méthode à estimer des seuils de détection pertinents. La procédure de détection et localisation par segment a ensuite été validée sur une séquence de signaux réels collectée de janvier à octobre 2012 dans le cadre du projet [EBSF](#). De nombreux blocages de portes ont pu être signalés et trois dérèglages de potentiomètres ont bien été détectés. L'outil se réinitialise automatiquement après chaque détection et montre ainsi sa capacité d'adaptabilité après plusieurs changements. La procédure proposée est décrite dans un cadre générique et permet d'imaginer un large champs d'applications.

Conclusions et perspectives

Synthèse des travaux

Le problème étudié dans le cadre de cette thèse porte essentiellement sur l'étape de détection de défaut dans un processus de diagnostic industriel. Ces travaux sont motivés par la surveillance de deux systèmes critiques d'autobus impactant la disponibilité des véhicules et leur coût de maintenance : le système de freinage et les portes. Afin de surveiller le fonctionnement de ces deux systèmes, on choisit une approche à base de reconnaissance des formes, détaillée dans le chapitre 2, qui s'appuie sur l'analyse de données collectées à partir d'une nouvelle architecture télématique en partie embarquée dans les autobus et développée dans le cadre du projet européen [EBSF](#). Cette approche s'est révélée particulièrement pertinente dans un contexte où l'expertise des données observées est faible et lorsque les modèles physiques des systèmes sont indisponibles.

L'objectif est de détecter un changement structurel dans un flux de données traité séquentiellement, puis éventuellement d'identifier la cause de ce changement. Dans ces travaux de thèse, le problème de la détection séquentielle, détaillée dans le chapitre 3, est considéré aussi bien pour le suivi d'un système de freins que celui de portes. En prévision d'une implémentation en embarqué, on choisit de composer des méthodes en ligne à base de tests d'hypothèses dont le coût calculatoire reste relativement peu élevé. A chaque instant d'analyse, le système surveillé peut être signalé hors de contrôle dès lors qu'un changement significatif apparaît dans la distribution des observations. Les méthodes proposées pour les études des freins et des portes, respectivement dans les chapitres 4 et 5 de cette thèse, sont évaluées en termes de détection à partir de données réelles. Des situations de dégradation n'étant pas toujours disponibles, des résultats ont été obtenus à partir de défaillances simulées. Toutefois, l'étude sur les portes a pu être illustrée par quelques dégradations réelles suite à un travail préliminaire de caractérisation des défauts. La procédure que nous avons proposée a ainsi permis de détecter puis localiser un dérèglement des potentiomètres de porte.

En outre, les méthodes proposées ne nécessitent pas de disposer, au préalable, de données labellisées par des états de fonctionnement avec certitude et précision, comme dans un problème de classification supervisée. Dans ce contexte, nous avons intégré aux méthodes de détection des connaissances disponibles sur les systèmes industriels (modèle physique du système de freinage, connaissance partielle de la segmentation des signatures de porte,...) de manière à améliorer les performances en termes de détection. Pour le système de freinage, on restreint l'espace des caractéristiques en étudiant les variables de sortie (liées au freinage) d'un modèle dynamique du véhicule qui a été validé expérimentalement dans le cadre nos travaux. Cette étude est présentée dans le chapitre 4 de cette thèse. Pour l'étude des portes, le détecteur s'appuie sur un nouveau modèle de régression multidimensionnel adapté à des données fonctionnelles et proposé pour la représentation de la densité jointe des courbes mesurées et leurs états latents qui définissent une segmentation sur les courbes. Ce modèle génératif est appris dans un mode semi-supervisé où certains labels de la segmentation sont fixés à partir d'une connaissance partielle de la cinétique des portes, justifiée par les experts des systèmes.

Par ailleurs, la méthode de détection appliquée aux freins s'appuie sur un modèle mathématique dont les variables de sortie sont physiquement interprétables. Ce choix de modélisation permet de sélectionner avec précision les variables d'intérêt puis d'alimenter efficacement le détecteur. Contrairement à cette approche basée sur la connaissance d'un modèle a priori, la stratégie proposée pour le suivi des portes s'appuie directement sur l'analyse des courbes mesurées à chaque cycle de fonctionnement du système. Cette seconde approche de détection et localisation dans un flux de courbes multivariées se révèle particulièrement intéressante car cette procédure répond à un problème assez générique et commun à de nombreuses applications réelles dont certaines sont brièvement évoquées par la suite.

Perspectives

Les outils proposés dans cette thèse résultent de travaux sur l'analyse de données nouvellement collectées sur des freins et des portes d'autobus en exploitation. Il serait intéressant d'approfondir les propriétés théoriques des méthodes proposées. D'autre part, nos travaux étant motivés par une application réelle, de nombreuses extensions au problème peuvent être envisagées, selon le système complexe à étudier. Quelques extensions des méthodes proposées et des problèmes étudiés sont décrites par la suite.

Concernant l'étude sur le système de freinage, le modèle physique du véhicule décrit dynamiquement son mouvement longitudinal. Cette modélisation simplifiée est justifiée par un compromis entre simplicité et précision d'une part et d'autre part, les données disponibles pour cette étude correspondent aux entrées du modèle longitudinal. A court terme, il serait utile de disposer d'information sur les actions relatives aux différents systèmes de freinage employés dans un autobus (cf. tableau 4.3), en particulier le ralentisseur

à commande électrique. En outre, un modèle qui prend en compte le mouvement latéral du véhicule pourrait être formulé sur la base des travaux de [Nouvelière \(2002, section 2.8\)](#) mais cette modélisation nécessite l’acquisition de données supplémentaires (cap, déplacement latéral). Une telle approche apporterait une précision plus fine à la modélisation des phases de freinage mais augmenterait automatiquement le degré de liberté du modèle. En outre, le modèle considéré dans cette thèse est de type « bicyclette » dans le sens où chaque essieu est ramené à une seule roue. Dans ce cas, les équations des moments sont appliquées sur l’axe principal de chaque roue mais on pourrait également étudier un modèle qui prend en compte la rotation au niveau de chacune des roues. Dans une autre approche, le système de freinage pourrait être surveillé sans s’appuyer sur un modèle analytique complexe. En d’autres termes, le fonctionnement des freins pourrait être directement analysé à partir de la dizaine de données CAN mesurées. Des méthodes de sélection de variables ou de projection pourraient être envisagées de manière à réduire la dimension des observations (sous-section 2.2.1). Cette approche nécessite de prendre en compte des données de natures différentes (continues, discrètes, constantes) d’une part, et d’autre part de valider la pertinence du sous-espace d’indicateurs choisi. Également pour cette thèse, il est préférable de constituer une base historisée la plus exhaustive possible et constituée d’exemples étiquetés en bon et mauvais fonctionnement afin d’apprendre à détecter des événements pertinents.

Par ailleurs, l’approche préconisée pour la surveillance des portes se base sur la paramétrisation des courbes par un modèle génératif dédié. L’apprentissage de ce modèle a été formulé de manière à respecter les conditions classiques d’un algorithme EM dont la convergence vers un maximum local est garantie. Plusieurs extensions pourraient être envisagées pour ce type de modèle. On cite notamment trois extensions :

- EM en-ligne : montée de gradient stochastique ([Titterton, 1984](#); [Sato, 2000](#); [Samé et al., 2007](#)), mise à jour récursive des paramètres ([Cappé et Moulines, 2009](#); [Cappé, 2011](#)) ;
- Initialisation non aléatoire : stratégie basée sur une connaissance disponible des données afin d’aider à répondre au problème posé ;
- Meilleure représentation de l’incertitude et de l’imprécision des labels : supervision partielle ([Ambroise et al., 2001](#)) où chaque étiquette peut être labellisée par un ensemble de classes, ou bien une supervision douce ([Côme, 2009](#), chapitre 3) où chaque étiquette peut être représentée par une fonction de croyance.

D’autre part, il est à noter que le seuil de détection, dans l’étude sur les données réelles, est estimé comme le quantile d’une loi a priori (lognormale) apprise sur les statistiques de test d’un ensemble de courbes sous $\mathcal{H}_0$, et non à partir d’une procédure de rééchantillonnage. Ce choix est motivé par le fait que les statistiques calculées à partir des courbes réelles et celles calculées à partir des courbes simulées par le modèle génératif, ne suivent pas expérimentalement la même loi de probabilité. A notre connaissance, cette différence

s'explique par la forte variabilité des courbes bidimensionnelles mesurées et notamment celle des courbes de position dont les valeurs étaient trop fortement discrétisées. Le modèle génératif, présenté dans la section 5.2, admet une variance importante autour d'une courbe moyenne mais il serait intéressant d'ajouter un effet aléatoire (*random effect*) sur les coefficients polynomiaux de régression $(\beta_k)_{1 \leq k \leq K}$, à la manière de James et Sugar (2003) sur les coefficients d'une base de splines. L'ajout d'un bruit additif sur les coefficients de régression autorise une variance sur la courbe moyenne. Suivant la paramétrisation des matrices de covariances (Celeux et Govaert, 1995), cette extension augmente plus ou moins la complexité du modèle génératif obtenu. En retour, cette nouvelle formulation permettrait de générer des courbes « plus variées » et ainsi, conduirait à une meilleure estimation du seuil de détection par une stratégie de rééchantillonnage.

A court terme, on pourrait utiliser le modèle génératif appliqué au suivi des portes pour la modélisation ou la segmentation d'autres données fonctionnelles, ou données longitudinales. L'extension multidimensionnelle du modèle de régression à processus latent, introduit dans la section 5.2, n'a pas fait l'objet d'une étude comparative dans le cadre de cette thèse en terme de segmentation ou de modélisation des courbes. A la manière de Chamroukhi (2010) pour la version univariée, il serait intéressant de comparer notre modèle à d'autres méthodes de régression multivariée. D'autre part, il est à noter que notre stratégie de détection sur un flux de données fonctionnelles, décrite dans la section 5.3, intègre une double notion de temps : chaque courbe est la discrétisation d'une fonction indexée par le temps et la séquence de courbes est ordonnée en fonction de la réception des courbes. Notre approche propose de segmenter les courbes observées puis de détecter un nouveau groupe de courbes à partir d'un test séquentiel. L'approche de *clustering par blocs*, également appelée *bi-clustering* ou *co-clustering* (Govaert et Nadif, 2003), serait intéressante à étudier car cette approche permet de grouper des variables sur deux dimensions simultanément. Dans ce sens, Nadif et Govaert (2005) proposent notamment un modèle de mélange par bloc, estimé via un algorithme de type EM par le principe du maximum de vraisemblance. L'idée de ce type de méthodes est de partitionner les individus et leurs variables en modélisant des états latents sur ces deux dimensions. Dans notre application, ce problème devrait prendre en compte la contrainte sur le temps dans le sens où l'objectif serait de chercher les limites des segments et celles des modes de fonctionnement, sans permutation des lignes ou des colonnes.

Toutes les méthodes proposées dans cette thèse pour la surveillance des systèmes de freinage et portes visent à détecter un changement structurel dans le flux de données observées. Ces méthodes sont construites à partir de tests séquentiels d'hypothèses et les observations sont supposées indépendantes. Or il est vraisemblable que cette hypothèse d'indépendance est quelque peu restrictive. Il serait donc intéressant de modifier les tests séquentiels de détection à la manière de Lai (1998) afin de prendre en compte des données dépendantes entre elles et ainsi, exploiter les relations temporelles entre les variables observées. La sous-section 3.3.2 liste quelques extensions de ce type de procédures à base

de tests d'hypothèses. Le formalisme des modèles markoviens pourrait être également considéré pour des données dépendantes comme le souligne la sous-section 3.4.2. Enfin, les propriétés d'optimalité des détecteurs proposés n'ont pas été explorées dans le cadre de cette thèse. Les garanties d'optimalité en détection sont généralement obtenues à partir de processus gaussiens ou lorsque les distributions paramétriques avant et après un changement sont connues. Dans nos applications, la difficulté d'établir les propriétés d'optimalité des méthodes s'explique par la non-connaissance des distributions avant/après changement et dans le cas des portes, par la nature fonctionnelle des données.

Enfin, ces travaux de thèse ouvrent plusieurs perspectives d'application industrielle. A titre d'illustration et sans volonté de dresser une liste exhaustive de cas d'application, on décrit brièvement trois problématiques liées à nos travaux concernant des données fonctionnelles :

- Dans le cas d'une gestion de la production électrique, des signatures quotidiennes de consommation électrique sont collectées par les exploitants. Le suivi de ces signaux permet d'adapter la production du système électrique en fonction de la consommation journalière.
- Dans le cas du suivi de fonctionnement d'un réseau d'eau potable, les mesures collectées sur le réseau (débit, concentration en chlore,...) peuvent être représentées par des données longitudinales (spatialement réparties sur le réseau). Une rupture dans le flux de ces signatures peut être le signe d'une intrusion ou d'une contamination dans le réseau d'eau.
- Afin de détecter des situations d'accidents en motos, les systèmes de déclenchement d'un gilet « airbag » à gonflage pyrotechnique nécessitent de se doter d'outils de surveillance à partir de capteurs embarqués. Le suivi des signaux temporels mesurés par une centrale inertielle (3 accéléromètres et 3 gyromètres) permet de détecter des changements brusques provoqués par des chutes ou des collisions.

Annexe **A**

Rappel de quelques outils mathématiques

Sommaire

A.1 Quelques définitions	178
A.2 Quelques opérateurs et inégalités	180

A.1 Quelques définitions

Définition A.1 (Espérance et variance)

L'espérance mathématique d'une variable aléatoire discrète X est la moyenne arithmétique des différentes valeurs de X pondérées par leurs probabilités :

$$\mathbb{E}(X) = \sum_{i=1} x_i \cdot P(X = x_i). \quad (\text{A.1})$$

La variance de X mesure la dispersion des différentes valeurs de X autour de son espérance, notée μ . On calcule cette quantité (positive ou nulle) comme suit

$$\text{Var}(X) = \mathbb{E}[(x - \mu)^2] = \mathbb{E}(x^2) - \mu^2 \quad (\text{A.2})$$

$$= \sum_{i=1} (x_i - \mu)^2 \cdot P(X = x_i). \quad (\text{A.3})$$

Enfin, il est à noter que l'écart-type de X est simplement la racine carrée de la variance :

$$\sigma_X = \sqrt{\text{Var}(X)}. \quad (\text{A.4})$$

Définition A.2 (Fonction continue)

Soit $f : \mathbb{R}^d \rightarrow \mathbb{R}$, une fonction et $d \in \mathbb{N}$. On dit que f est une fonction continue sur $I \subset \mathbb{R}^d$, si

$$\lim_{x \rightarrow a} f(x) = f(a) \quad \forall a \in I. \quad (\text{A.5})$$

Définition A.3 (Fonction convexe)

Soit $f : \mathbb{R}^d \rightarrow \mathbb{R}$, une fonction où $d \in \mathbb{N}$. On dit que f est une fonction convexe, si

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y), \quad (\text{A.6})$$

où $x, y \in \mathbb{R}^d$ et $\lambda \in [0, 1]$.

On dit que f est une fonction strictement convexe, si

$$f(\lambda x + (1 - \lambda)y) < \lambda f(x) + (1 - \lambda)f(y), \quad (\text{A.7})$$

où $x, y \in \mathbb{R}^d$, $x \neq y$ et $\lambda \in [0, 1]$.

Définition A.4 (Différentiabilité, Vecteur gradient et Matrice hessienne)

Soit $f : \mathbb{R}^d \rightarrow \mathbb{R}$, une fonction. On dit que f est une fonction différentiable en $x \in \mathbb{R}^d$ s'il

existe un vecteur $\nabla f(x) \in \mathbb{R}^d$ et une fonction $\epsilon : \mathbb{R}^d \rightarrow \mathbb{R}$ tels que

$$f(x+h) = f(x) + \langle \nabla f(x), h \rangle + \|h\| \epsilon(h) \quad \forall h \in \mathbb{R}^d, \quad (\text{A.8})$$

où $\langle \cdot, \cdot \rangle$ désigne un produit scalaire, et $\epsilon(h) \rightarrow 0$ quand $\|h\| \rightarrow 0$. S'il existe, le vecteur $\nabla f(x)$ est appelé *vecteur gradient* de la fonction f en x et défini par

$$\nabla f(x) = \begin{pmatrix} \frac{\partial}{\partial x_1} f(x) \\ \vdots \\ \frac{\partial}{\partial x_d} f(x) \end{pmatrix}, \quad (\text{A.9})$$

où $\frac{\partial}{\partial x_i} f(x)$ représentent les *dérivées partielles* de f .

D'autre part, la fonction f est deux fois différentiable en $x \in \mathbb{R}^d$ s'il existe une matrice $\nabla^2 f(x) \in \mathbb{R}^{d \times d}$. Cette matrice est appelée *matrice hessienne* et $\nabla^2 f(x)$ est définie comme la matrice des dérivées partielles secondes de f en x

$$\nabla^2 f(x) = \begin{pmatrix} \frac{\partial^2}{\partial x_1^2} f(x) & \cdots & \frac{\partial^2}{\partial x_1 \partial x_d} f(x) \\ \vdots & \ddots & \vdots \\ \frac{\partial^2}{\partial x_d \partial x_1} f(x) & \cdots & \frac{\partial^2}{\partial x_d^2} f(x) \end{pmatrix}. \quad (\text{A.10})$$

Cette matrice est symétrique sous quelques conditions, d'après le théorème de Schwarz.

Définition A.5 (Optimum global et local)

Dans un problème d'optimisation, on cherche à maximiser (ou minimiser) une fonction objectif $f : \mathbb{R}^d \rightarrow \mathbb{R}$. On distingue deux types de solution dans le cas d'une maximisation :

- Un point $x \in \mathbb{R}^d$ est un *maximum global* de f si

$$f(x) \geq f(y) \quad \forall y \in \mathbb{R}^d. \quad (\text{A.11})$$

- Un point $x \in \mathbb{R}^d$ est un *maximum local* de f si

$$\exists \epsilon > 0 \quad f(x) \geq f(y) \quad \forall y \in \mathcal{B}(x, \epsilon), \quad (\text{A.12})$$

où $\mathcal{B}(x, \epsilon)$ est une boule centrée en x et contient l'ensemble des points dont la distance à x est inférieure à ϵ .

Il est à noter que si f est une fonction convexe, tout maximum local x^* correspond au maximum global et $0 \in \nabla f(x^*)$.

A.2 Quelques opérateurs et inégalités

L'inégalité de Jensen ([Cover et Thomas, 2006](#), section 2.6) est utile pour établir l'algorithme EM, défini dans la sous-section 2.2.3. On présente cette inégalité comme suit, dans la définition A.6.

Définition A.6 (Inégalité de Jensen)

Soit X , une variable aléatoire. Si f est une fonction convexe (définition A.3), alors l'inégalité de Jensen est définie par

$$f(\mathbb{E}(X)) \leq \mathbb{E}(f(X)). \quad (\text{A.13})$$

De même, si f est une fonction concave ($\nabla^2 f \leq 0$), alors l'inégalité de Jensen est définie par

$$f(\mathbb{E}(X)) \geq \mathbb{E}(f(X)). \quad (\text{A.14})$$

Par ailleurs, le critère 3.1 de [Lorden \(1971\)](#) est un des premiers critères d'optimalité proposé pour une méthode de détection séquentielle. Ce résultat s'appuie sur l'opérateur de la borne supérieure essentielle. On définit cet opérateur comme suit, dans la définition A.7.

Définition A.7 (Borne supérieure essentielle (supremum essentiel))

Soit $(\Omega, \mathbb{F}, \mu)$, un espace mesuré où $f \in \mathbb{F}$ est une fonction (mesure borélienne) telle que $f : \Omega \rightarrow \bar{\mathbb{R}}$. La borne supérieure essentielle de f est le plus petit nombre $\alpha \in \bar{\mathbb{R}}$ parmi lesquels f dépasse α seulement sur un ensemble de mesure nulle. Il est à noter que la borne supérieure de f dépasse son supremum essentiel : $\text{ess sup } (f) \leq \sup(f)$.

Plus formellement, on définit $\text{ess sup } (f)$ comme suit. Soit $\alpha \in \mathbb{R}$, on définit

$$M_\alpha = \{x \in \Omega : f(x) > \alpha\}, \quad (\text{A.15})$$

comme le sous-ensemble de Ω où $f(x)$ est supérieur à α . On définit également

$$A_0 = \{\alpha \in \mathbb{R} : \mu(M_\alpha) = 0\}, \quad (\text{A.16})$$

comme le sous-ensemble de $\mathbb{R}$ pour lequel M_α est de mesure nulle. Le supremum essentiel de f est défini par

$$\text{ess sup } (f) = \inf A_0. \quad (\text{A.17})$$

Le supremum essentiel étant défini sur $\bar{\mathbb{R}} = \mathbb{R} \cup \{-\infty, \infty\}$, on remarque que $\text{ess sup } (f) = \infty$ si $A_0 = \emptyset$ et $\text{ess sup } (f) = -\infty$ si $A_0 = \mathbb{R}$.

Bibliographie

- AFNOR. Maintenance - Indicateurs de performance clés pour la maintenance. Technical Report NF EN 15341, Juin 2007.
- AFNOR. Maintenance - Terminologie de la maintenance. Technical Report NF EN 13306, Oct. 2010.
- H. Akaike. A new look at the statistical model identification. *Automatic Control, IEEE Transactions on*, 19(6) :716–723, 1974.
- P. Aknin. *De la mesure à sa représentation et sa discrimination. Application au diagnostic des infrastructures ferroviaires*. HDR thesis, ENS Cachan, Déc. 2008.
- H. Albazzaz et X. Z. Wang. Statistical process control charts for batch operations based on independent component analysis. *Industrial & engineering chemistry research*, 43(21) : 6731–6741, 2004.
- C. Ambroise, T. Denoeux, G. Govaert, et P. Smets. Learning from an imprecise teacher : probabilistic and evidential approaches. Dans *Proceedings of Applied Stochastic Models and Data Analysis (ASMDA)*, 2001.
- N. Amenc, B. Maffei, et H. Till. Les causes structurelles du troisième choc pétrolier. Technical report, EDHEC Risk and Asset Management Research Centre, nov 2008.
- A. Amiri et S. Allahyari. Change point estimation methods for control chart postsignal diagnostics : a literature review. *Quality and Reliability Engineering International*, 28(7) : 673–685, 2012.
- F. Aparisi. Hotelling’s t^2 control chart with adaptive sample sizes. *International Journal of Production Research*, 34(10) :2853–2862, 1996.
- S. Arlot, A. Celisse, et Z. Harchaoui. Kernel change-point detection. *arXiv preprint arXiv :1202.3878*, 2012.
- J. A. D. Aston et C. Kirch. Detecting and estimating changes in dependent functional data. *Journal of Multivariate Analysis*, 109 :204–220, Août 2012.

- M. Basseville. Detecting changes in signals and systems – a survey. *Automatica*, 24(3) : 309–326, 1988.
- M. Basseville et M.-O. Cordier. Surveillance et diagnostic de systèmes dynamiques : approches complémentaires du traitement de signal et de l’intelligence artificielle. Technical report, INRIA, 1996.
- M. Basseville et I. V. Nikiforov. *Detection of abrupt changes : theory and application*, volume 104. Prentice Hall Englewood Cliffs, Avr. 1993.
- C. Baysse, D. Bihannic, A. Gégout-Petit, M. Prenat, et J. Saracco. Hidden Markov Model for the detection of a degraded operating mode of optronic equipment. *arXiv preprint arXiv :1212.2358*, Déc. 2012.
- R. Bellman. *Adaptive control processes : a guided tour*. Rand Corporation Research studies. Princeton University Press, 1961.
- A. Ben Salem. *Modèles Probabilistes de Séquences Temporelles et Fusion de Décisions. Application à la Classification de Défauts de Rails et à leur Maintenance*. PhD thesis, Université de Nancy 1, Mar. 2008.
- Y. Bengio et P. Frasconi. Input-Output HMMs for Sequence Processing. *Neural Networks, IEEE Transactions on*, 7(5) :1231–1249, 1996.
- S. Bersimis, S. Psarakis, et J. Panaretos. Multivariate Statistical Process Control Charts : An Overview. *Quality and Reliability Engineering International*, 23(5) :517–543, Août 2007.
- C. Biernacki. Initializing EM using the properties of its trajectories in Gaussian mixtures. *Statistics and Computing*, 14(3) :267–279, Août 2004.
- C. Biernacki, G. Celeux, et G. Govaert. Assessing a mixture model for clustering with the integrated completed likelihood. *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, 22(7) :719–725, 2000.
- C. M. Bishop. *Pattern recognition and machine learning*. Springer New York, Première édition, 2006.
- M. Borga, T. Landelius, et H. Knutsson. A unified approach to PCA, PLS, MLR and CCA. Technical report, Computer Vision Laboratory, Department of Electrical Engineering, Linköping University, S-581 83 Linköping, Sweden, Nov. 1992.
- Bosch. Road vehicles – Controller Area Network (CAN) – Part 1 : Data link layer and physical signalling. Technical Report ISO 11898, 2003.
- B. E. Boser, I. M. Guyon, et V. N. Vapnik. A training algorithm for optimal margin classifiers. Dans *Proceedings of the 5th annual workshop on Computational Learning Theory, COLT ’92*, pages 144–152, New York, NY, USA, 1992. ACM.

- H. A. Boubacar. *Classification Dynamique de données non-stationnaires :Apprentissage et Suivi de Classes évolutives*. PhD thesis, Université des Sciences et Technologie de Lille - Lille I, 2006.
- G. Bouchard. *Les modèles génératifs en classification supervisée et applications à la catégorisation d'images et à la fiabilité industrielle*. PhD thesis, Université Joseph Fourier - Grenoble 1, 2005.
- S. Boyd et L. Vandenberghe. *Convex Optimization*. Cambridge University Press, Première édition, Mar. 2004.
- H. Bozdogan. Model Selection and Akaike Information Criteria. *Psychometrika*, 52 :345–370, 1987.
- A. P. Bradley. The use of the area under the ROC curve in the evaluation of machine learning algorithms. *Pattern Recognition*, 30(7) :1145–1159, Juil. 1997.
- M. Breunig, H.-P. Kriegel, R. T. Ng, et J. A. Sander. LOF : Identifying Density-Based Local Outliers. Dans *Proceedings of the International Conference on Management of Data, SIGMOD '00*, pages 93–104. ACM, 2000.
- F. Bu et H. S. Tan. Pneumatic Brake Control for Precision Stopping of Heavy-Duty Vehicles. *Control Systems Technology, IEEE Transactions on*, 15(1) :53–64, Jan. 2007.
- O. Cappé. Online EM Algorithm for Hidden Markov Models. *Journal of Computational and Graphical Statistics*, 20(3) :728–749, Jan. 2011.
- O. Cappé et E. Moulines. On-line expectation-maximization algorithm for latent data models. *Journal of the Royal Statistical Society : Series B (Statistical Methodology)*, 71(3) : 593–613, Sept. 2009.
- C. Celeux et G. Govaert. A classification EM algorithm for clustering and two stochastic versions. *Computational Statistics and Data Analysis*, 14(3) :315–332, 1992.
- G. Celeux et G. Govaert. Gaussian parsimonious clustering models. *Pattern Recognition*, 28(5) :781–793, Mai 1995.
- Certu. Bus a Haut Niveau de Service (BHNS) : du choix du système à sa mise en oeuvre. Technical report, Certu, Cete, Gart, INRETS, UTP, Nov. 2009.
- F. Chamroukhi. *Hidden process regression for curve modeling, classification and tracking*. PhD thesis, Université de Technologie de Compiègne, Déc. 2010.
- F. Chamroukhi, A. Samé, G. Govaert, et P. Aknin. A hidden process regression model for functional data description. application to curve discrimination. *Neurocomputing*, 73(7) :1210–1221, 2010.

- F. Chamroukhi, H. Glotin, et A. Samé. Model-based functional mixture discriminant analysis with hidden process regression for curve classification. *Neurocomputing*, Mar. 2013.
- V. Chandola, A. Banerjee, et V. Kumar. Anomaly detection : A survey. *ACM Comput. Surv.*, 41(3) :1–58, Juil. 2009.
- O. Chapelle, B. Schölkopf, et A. Zien. *Semi-Supervised Learning*. Adaptive Computation and Machine Learning. MIT Press, Première édition, 2010.
- S. Charbonnier. *Surveillance de systèmes continus à l'aide de méthodes n'utilisant pas de modèles formels : Application aux systèmes médicaux*. HDR thesis, Institut National Polytechnique de Grenoble (INPG), Nov. 2006.
- S. Chatterjee et P. Qiu. Distribution-Free Cumulative Sum Control Charts Using Bootstrap-Based Control Limits. *The Annals of Applied Statistics*, 3(1) :349–369, 2009.
- V. Cherkassky et F. M. Mulier. *Learning from data : concepts, theory, and methods*. Wiley-IEEE Press, 2007.
- P. Cocheteux. *Contribution à la maintenance proactive par la formalisation du processus de pronostic des performances de systèmes industriels*. PhD thesis, Université Henri Poincaré - Nancy I, Nov. 2010.
- E. Côme. *Apprentissage de modèles génératifs pour le diagnostic de systèmes complexes avec labellisation douce et contraintes spatiales*. PhD thesis, Université de Technologie de Compiègne, Jan. 2009.
- T. Cover et J. Thomas. *Elements of Information Theory*. Wiley Series in Telecommunications and Signal Processing. John Wiley & Sons, Seconde édition, 2006.
- R. B. Crosier. Multivariate generalizations of cumulative sum quality-control schemes. *Technometrics*, 30 :291–303, Août 1988.
- Daimler, MAN, Scania, Volvo, IrisBus, et VDL. Bus FMS-Standard Interface description, Seconde édition. Technical report, Working Group BCEI, Juil. 2009.
- H. Dassanayake. *Fault diagnosis for a new generation of intelligent train door systems*. PhD thesis, University of Birmingham, 2002.
- H. Dassanayake, C. Roberts, C. J. Goodman, et A. M. Tobias. Use of parameter estimation for the detection and diagnosis of faults on electric train door systems. *Journal of Risk and Reliability*, 223 :271–278, 2009.
- R. Davis et J. Franke. Mini-Workshop : Time Series with Sudden Structural Changes. *Oberwolfach Reports*, pages 557–586, 2008.

- A. P. Dempster, N. M. Laird, et D. B. Rubin. Maximum Likelihood from Incomplete Data via the EM Algorithm. *Journal of the Royal Statistical Society. Series B (Methodological)*, 39 (1) :1–38, 1977.
- T. G. Dietterich. Machine Learning for Sequential Data : A Review. Dans *Proceedings of the Joint IAPR International Workshop on Structural, Syntactic, and Statistical Pattern Recognition*, pages 15–30, London, UK, 2002. Springer-Verlag.
- T. M. T. Do et T. Artières. Champs de markov conditionnels pour le traitement de séquences. *Conférence Francophone sur l'Extraction et la Gestion des Connaissances : EGC'2006*, 2006.
- R. Donat. *Modélisation de la Fiabilité et de la Maintenance par Modèles Graphiques Probabilistes - Application à la Prévention des Ruptures de Rails*. PhD thesis, INSA de Rouen, Nov. 2009.
- V. P. Dragalin, A. G. Tartakovsky, et V. V. Veeravalli. Multihypothesis Sequential Probability Ratio Tests-Part I : Asymptotic Optimality. *Information Theory, IEEE Transactions on*, 45(7) :2448–2461, 1999.
- B. Dubuisson. *Automatique et statistiques pour le diagnostic*. Traité IC2. Série productique. Hermes Science Publications, 2001a.
- B. Dubuisson. *Diagnostic, intelligence artificielle et reconnaissance des formes*. Traité IC2. Série productique. Hermes Science Publications, 2001b.
- R. Duda, P. Hart, et D. Stork. *Pattern classification*. Pattern Classification and Scene Analysis : Pattern Classification. Wiley, Seconde édition, 2001.
- EBSF Consortium. Projet European Bus System of the Future. <http://www.ebsf.eu/>, Oct. 2012.
- B. Efron et R. Tibshirani. *An Introduction to the Bootstrap*. Monographs on Statistics and Applied Probability. Taylor & Francis, Première édition, 1994.
- B. Efron, T. Hastie, I. Johnstone, et R. Tibshirani. Least Angle Regression. *The Annals of Statistics*, 32(2) :407–499, 2004.
- J. Fan et S.-K. Lin. Test of significance when data are curves. *Journal of the American Statistical Association*, 93(443) :1007–1021, 1998.
- P.-P. Faure. *Une approche à base de modèles fondés sur les intervalles pour la génération automatique d'arbres de diagnostic optimaux. Application au domaine de l'automobile*. PhD thesis, Université Paris 13, Juin 2001.
- Federal Motor Carrier Safety Administration (FMCSA). Report to Congress on the Large Truck Crash Causation Study. Technical report, U.S. Department of Transportation, Mar. 2006.

- F. Ferraty et Y. Romain. *The Oxford Handbook of Functional Data Analysis* (Oxford Handbooks). Oxford University Press, USA, Première édition, Jan. 2011.
- L. Fillatre. *Contributions en Détection et Classification Statistique Paramétrique*. HDR thesis, Université de Technologie de Compiègne, Déc. 2011.
- D. Fisher. Brake System Component Dynamic Performance Measurement and Analysis. Technical report, SAE Technical Paper, 1970.
- D. Freund et D. Skorupski. Commercial vehicle safety technologies : applications for brake performance monitoring. Technical report, US Departement of Transportation, Avr. 2009.
- S. Frühwirth-Schnatter. *Finite Mixture and Markov Switching Models*. Springer Series in Statistics. Springer, Première édition, 2006.
- S. J. Gaffney. *Probabilistic Curve-Aligned Clustering and Prediction with Regression Mixture Models*. PhD thesis, University of California, Irvine, 2004.
- J. C. Gerdes. *Decoupled Design of Robust Controllers for Nonlinear Systems : As Motivated by and Applied to Coordinated Throttle and Brake Control for Automated Highways*. PhD thesis, University of California, Berkeley, 1996.
- G. Govaert. *Analyse des données*. Traité IC2. Série traitement du signal et de l'image. Hermes Science Publications, Première édition, 2003.
- G. Govaert et M. Nadif. Clustering with block mixture models. *Pattern Recognition*, 36(2) :463–473, 2003.
- P. J. Green. Iteratively Reweighted Least Squares for Maximum Likelihood Estimation, and some Robust and Resistant Alternatives. *Journal of the Royal Statistical Society. Series B (Methodological)*, 46(2), 1984.
- F. Gustafsson et J. Palmqvist. Change Detection Design For Low False Alarm Rates. Technical report, Department of Electrical Engineering, Linköping University, 2000.
- I. Guyon et A. Elisseeff. An introduction to variable and feature selection. *J. Mach. Learn. Res.*, 3 :1157–1182, Mar. 2003.
- A. Haghani et Y. Shafahi. Bus maintenance systems and maintenance scheduling : model formulations and solutions. *Transportation Research Part A : Policy and Practice*, 36(5) : 453–482, Juin 2002.
- D. J. Hand. L'incohérence de l'aire sous la courbe ROC, que faire à ce propos ? *Revue MODULAD*, 74(42), 2010.
- T. Hastie, R. Tibshirani, et J. Friedman. *The Elements of Statistical Learning : Data Mining, Inference, and Prediction*. Springer Series in Statistics. Springer, Seconde édition, 2009.

- J. K. Hedrick et M. Uchanski. Brake System Modeling and Control. Technical report, University of California, Berkeley, Jan. 2001.
- J. K. Hedrick, J. C. Gerdes, D. B. Maciucă, et D. Swaroop. Brake System Modeling, Control and Integrated Brake/Throttle Switching : Phase I. Technical report, University of California, Berkeley, Mai 1997.
- K. Hedrick, A. Howell, et B. Song. Fault Tolerant Longitudinal Control of Transit Buses : Modeling and Control. Dans *2001 PATH Conference*. Dept. of Mechanical Engineering University of California, Berkeley, 2001.
- V. Hodge et J. Austin. A survey of outlier detection methodologies. *Artificial Intelligence Review*, 22(2) :85–126, 2004.
- A. E. Hoerl et R. W. Kennard. Ridge regression : Biased estimation for nonorthogonal problems. *Technometrics*, 12(1) :55–67, 1970.
- S. Huang et W. Ren. Vehicle longitudinal control using throttles and brakes. *Robotics and Autonomous Systems*, 26(4) :241–253, Mar. 1999.
- R. Isermann. Process fault detection based on modeling and estimation methods - A survey. *Automatica*, 20(4) :387–404, Juil. 1984.
- R. Isermann. Perspectives of automatic control. *Control Engineering Practice*, 19(12) :1399–1407, Déc. 2011.
- J. E. Jackson. *A User's Guide to Principal Components*. Wiley Series in Probability and Mathematical Statistics. John Wiley & Sons, Inc., 2004.
- J. Jacques. *Contribution à l'apprentissage statistique à base de modèles génératifs pour données complexes*. HDR thesis, Université des Sciences et Technologie de Lille - Lille I, Nov. 2012.
- G. M. James et C. A. Sugar. Clustering for Sparsely Sampled Functional Data. *Journal of the American Statistical Association*, 98(462) :397–408, 2003.
- T. Jebara. *Machine Learning : Discriminative and Generative*. Kluwer International series in Engineering and Computer Science. Kluwer Academic Publishers, Première édition, 2004.
- M. K. Jeong, J.-C. Lu, et N. Wang. Wavelet-based spc procedure for complicated functional data. *International Journal of Production Research*, 44(4) :729–744, 2006.
- L. Johnson, P. Fancher, T. Gillespie, U. of Michigan. Highway Safety Research Institute, et M. V. M. Association. *An Empirical Model for the Prediction of the Torque Output of Commercial Vehicle Air Brakes*. Highway Safety Research Institute, 1978.

- P. Karthikeyan, D. Sonawane, et S. Subramanian. Model-based control of an electro-pneumatic brake system for commercial vehicles. *International Journal of Automotive Technology*, 11(4) :507–515, Août 2010.
- P. Karthikeyan, C. Siva Chaitanya, N. Jagga Raju, et S. C. Subramanian. Modelling an electropneumatic brake system for commercial vehicles. *Electrical Systems in Transportation, IET*, 1(1) :41–48, Mar. 2011.
- T. Kempowsky. *Surveillance de procédés à base de méthodes de classification : conception d'un outil d'aide pour la détection et le diagnostic des défaillances*. PhD thesis, LAAS-CNRS, Déc. 2004.
- Y. Khan, P. Kulkarni, et K. Youcef-Toumi. Modeling, Experimentation and Simulation of a Brake Apply System. *Journal of Dynamic Systems, Measurement, and Control*, 116(1) : 1–12, Mar. 1994.
- R. Killick, C. F. Nam, J. Aston, et E. I.A. Changepoint.info : The changepoint repository. <http://changepoint.info/>, 2012.
- H.-J. Kim, B. Yu, et E. J. Feuer. Selecting the number of change-points in segmented line regression. *Statistica Sinica*, 19(2) :597, 2009.
- W. J. Krzanowski et D. J. Hand. *ROC Curves for Continuous Data*. Chapman & Hall/CRC Monographs on Statistics & Applied Probability, Première édition, Mai 2009.
- J. Lafferty, A. McCallum, et F. Pereira. Conditional Random Fields : Probabilistic Models for Segmenting and Labeling Sequence Data. Dans *Proceedings of the 18th International Conference on Machine Learning, ICML '01*. Morgan Kaufmann, San Francisco, CA, 2001.
- T. L. Lai. Sequential Changepoint Detection in Quality Control and Dynamical Systems. *Journal of the Royal Statistical Society. Series B (Methodological)*, 57(4) :613–658, 1995.
- T. L. Lai. Information bounds and quick detection of parameter changes in stochastic systems. *Information Theory, IEEE Transactions on*, 44(7) :2917–2929, Nov. 1998.
- T. L. Lai. Sequential analysis : some classical problems and new challenges. *Statist. Sinica*, 11(2) :303–408, 2001.
- T. L. Lai et H. Xing. Sequential Change-Point Detection When the Pre- and Post-Change Parameters are Unknown. *Sequential Analysis*, 29(2) :162–175, Avr. 2010.
- M. Larsson. Road Slope Estimation. Master's thesis, Linköping University, Automatic Control, Jan. 2010.
- J. Lee, B. Kang, et S.-H. Kang. Integrating independent component analysis and local outlier factor for plant-wide process monitoring. *Journal of Process Control*, 21(7) :1011–1021, Août 2011.

- N. Lehasab. *A generic fault detection and isolation approach for single-throw mechanical equipment*. PhD thesis, University of Birmingham, 1999.
- N. Lehasab, H. P. B. Dassanayake, C. Roberts, S. Fararooy, et C. J. Goodman. Industrial fault diagnosis : pneumatic train door case study. *Journal of Rail and Rapid Transit*, 216 : 175–183, 2002.
- R. Limpert. *Brake Design and Safety*. SAE International, Troisième édition, Oct. 2011.
- G. Lorden. Procedures for Reacting to a Change in Distribution. *The Annals of Mathematical Statistics*, 42(6), 1971.
- C. A. Lowry et D. C. Montgomery. A review of multivariate control charts. *IIE transactions*, 27(6) :800–810, 1995.
- C. A. Lowry, W. H. Woodall, C. W. Champ, et S. E. Rigdon. A multivariate exponentially weighted moving average control chart. *Technometrics*, 34 :46–53, Fév. 1992.
- M. Luong, V. Perduca, et G. Nuel. Hidden Markov Model Applications in Change-Point Analysis, Déc. 2012.
- J. MacGregor et T. Kourti. Statistical process control of multivariate processes. *Control Engineering Practice*, 3(3) :403–414, Mar. 1995.
- D. B. Maciuca. *Nonlinear Robust and Adaptive Control with Application to Brake Control for Automated Highway Systems*. PhD thesis, University of California, Berkeley, 1997.
- D. B. Maciuca, C. Gerdes, et J. K. Hedrick. Automatic Braking Control for IVHS. Dans *International Symposium on Advanced Vehicle Control (AVEC)*, 1994.
- M. Markou. Novelty Detection : A Review - Part 1 : Statistical Approaches. *Signal Processing*, 83(12) :2481–2497, Déc. 2003a.
- M. Markou. Novelty Detection : A Review - Part 2 : Neural Network Based Approaches. *Signal Processing*, 83(12) :2499–2521, Déc. 2003b.
- N. Martin et C. Doncarli. *Décision et Interprétation Non Stationnaire*. GDR ISIS, 1998.
- C. W. McCullers et A. E. Smith. Predictive diagnostics for motor driven automated doors. Dans *the American Public Transportation Association Rail Transportation Conference*, Juin 2001.
- G. McLachlan et T. Krishnan. *The EM algorithm and extensions*. Wiley Series in Probability and Statistics. Wiley-Interscience, Seconde édition, 2008.
- G. McLachlan et D. Peel. *Finite Mixture Models*. Wiley Series in Probability and Statistics. Wiley-Interscience, Première édition, 2000.

- G. J. McLachlan. Estimating the Linear Discriminant Function from Initial Samples Containing a Small Number of Unclassified Observations. *Journal of the American Statistical Association*, 72(358) :403–406, Juin 1977.
- D. H. McMahon, V. Narendran, D. Swaroop, J. Hedrick, K. Chang, et P. Devlin. Longitudinal vehicle controllers for IVHS : Theory and experiment. Dans *American Control Conference*, 1992, pages 1753–1757. IEEE, 1992.
- Y. Mei. *Asymptotically Optimal Methods for Sequential Change-Point Detection*. PhD thesis, California Institute of Technology Pasadena, California, USA, Mai 2003.
- Y. Mei. Sequential change-point detection when unknown parameters are present in the pre-change distribution. *The Annals of Statistics*, 34(1) :92–122, 2006.
- K. Mengersen, C. Robert, et M. Titterton. *Mixtures : Estimation and Applications*. Wiley Series in Probability and Statistics. Wiley, Première édition, 2011.
- C. Meynet. *Sélection de variables pour la classification non supervisée en grande dimension*. PhD thesis, Université Paris Sud - Paris XI, Nov. 2012.
- E. Miguelanez, K. E. Brown, R. Lewis, C. Roberts, et D. M. Lane. Fault diagnosis of a train door system based on semantic knowledge representation. Dans *Railway Condition Monitoring*, pages 1–6, Juin 2008.
- M. Minoux. *Programmation Mathématique. Théorie et Algorithmes*. Lavoisier, Seconde édition, 2007.
- J. M. Moguerza, A. Muñoz, et S. Psarakis. Monitoring nonlinear profiles using support vector machines. Dans *Progress in Pattern Recognition, Image Analysis and Applications*, pages 574–583. Springer, 2007.
- G. V. Moustakides. Optimal Stopping Times for Detecting Changes in Distributions. *The Annals of Statistics*, 14(4) :1379–1387, 1986.
- G. V. Moustakides. Sequential change detection revisited. *The Annals of Statistics*, 36(2) : 787–807, 2008.
- M. Nadif et G. Govaert. Block clustering of contingency table and mixture model. Dans *Advances in Intelligent Data Analysis VI*, pages 249–259. Springer, 2005.
- S. Natarajan, S. Subramanian, S. Darbha, et K. Rajagopal. A model of the relay valve used in an air brake system. *Nonlinear Analysis : Hybrid Systems*, 1(3) :430–442, Sept. 2007.
- I. V. Nikiforov. A simple recursive algorithm for diagnosis of abrupt changes in random signals. *Information Theory, IEEE Transactions on*, 46(7) :2740–2746, 2000.
- I. V. Nikiforov. A lower bound for the detection/isolation delay in a class of sequential tests. *Information Theory, IEEE Transactions on*, 49(11) :3037–3047, Nov. 2003.

- R. Noorossana, A. Saghaei, et A. Amiri. *Statistical analysis of profile monitoring*, volume 865. Wiley, 2012.
- L. Nouvelière. *Commande robuste pour l'assistance à la conduite : l'automatisation basse vitesse*. PhD thesis, Université de Versailles - Saint Quentin en Yvelines, Déc. 2002.
- L. Nouvelière et M. Braci. Modélisation dynamique d'un bus du réseau TEOR. Technical report, Eurolum - INRETS-LIVIC, Mar. 2007.
- M. Pacella et Q. Semeraro. Monitoring roundness profiles based on an unsupervised neural network algorithm. *Computers & Industrial Engineering*, 60(4) :677–689, Mai 2011.
- E. S. Page. Continuous Inspection Schemes. *Biometrika*, 41(1/2) :100–115, 1954.
- K. Pearson. On lines and planes of closest fit to systems of points in space. *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science*, 2(11) :559–572, 1901.
- J. J. Pignatiello et G. C. Runger. Comparisons of multivariate CUSUM charts. *Journal of quality technology*, 22 :173–186, 1990.
- G. Plassat. *Les technologies des moteurs de véhicules lourds et leurs carburants*, volume 1 of *Données et références*. ADEME, 2005.
- M. Pollak. Average Run Lengths of an Optimal Method of Detecting a Change in Distribution. *The Annals of Statistics*, 15(2), 1987.
- M. Pollak et D. Siegmund. Approximations to the Expected Sample Size of Certain Sequential Tests. *The Annals of Statistics*, 3(6) :1267–1282, 1975.
- A. Polunchenko et A. G. Tartakovsky. State-of-the-Art in Sequential Change-Point Detection. *Methodology and Computing in Applied Probability*, 14(3) :649–684, 2012.
- T. M. Post, J. E. Bernard, et P. Fancher. Torque Characteristics of Commercial Vehicle Brakes. Technical report, 1975.
- F. Puhn. *Brake Handbook*. HP Books, Première édition, 1985.
- L. R. Rabiner. A tutorial on hidden Markov models and selected applications in speech recognition. *Proceedings of the IEEE*, 77(2) :257–286, Fév. 1989.
- J. Ramsay, G. Hooker, et S. Graves. *Functional Data Analysis with R and MATLAB*. Use R ! Springer-Verlag New York, 2009.
- J. O. Ramsay et B. W. Silverman. *Functional Data Analysis*. Springer Series in Statistics. Springer, Seconde édition, Juin 2005.
- S. Raudys. *Statistical and Neural Classifiers : An Integrated Approach to Design*. Advances in Pattern Recognition. Springer, 2001.

- P. Ribot. *Vers l'intégration diagnostic/pronostic pour la maintenance des systèmes complexes*. PhD thesis, LAAS-CNRS, Déc. 2009.
- Y. Ritov. Decision Theoretic Optimality of the Cusum Procedure. *The Annals of Statistics*, 18(3) :1464–1469, 1990.
- C. Roberts, H. P. B. Dassanayake, N. Lehasab, et C. J. Goodman. Distributed quantitative and qualitative fault diagnosis : railway junction case study. *Control Engineering Practice*, 10(4) :419–429, Avr. 2002.
- S. W. Roberts. A Comparison of Some Control Chart Procedures. *Technometrics*, 8 :411–430, 1966.
- A. Roche. Em algorithm and variants : An informal tutorial. *arXiv preprint arXiv :1105.1476*, 2012.
- T. Ryan. *Statistical methods for quality improvement*. Wiley Series in Applied Probability and Statistics Section Series. Wiley, Seconde édition, 2000.
- SAE Technical Standards. J1939 : Serial Vehicle Network for Trucks and Buses. Technical report, SAE International, June 2006.
- A. Samé, G. Govaert, et C. Ambroise. A Mixture Model-Based On-line CEM Algorithm. Dans *Advances in Intelligent Data Analysis VI*, volume 3646 of *Lecture Notes in Computer Science*, pages 739–739. Springer Berlin Heidelberg, 2005.
- A. Samé, C. Ambroise, et G. Govaert. An online classification EM algorithm based on the mixture model. *Statistics and Computing*, 17(3) :209–218, Sept. 2007.
- A. Samé, F. Chamroukhi, G. Govaert, et P. Aknin. Model-based clustering and segmentation of time series with changes in regime. *Advances in Data Analysis and Classification*, 5(4) :301–321, Déc. 2011.
- G. Saporta. *Probabilités, analyses des données et statistique*. Éditions Technip, Seconde édition, Juin 2006.
- M.-A. Sato. Convergence of On-line EM Algorithm. Dans *ICONIP 2000*, volume 1, pages 476–481, 2000.
- B. Schölkopf et A. J. Smola. *Learning with kernels : support vector machines, regularization, optimization and beyond*. the MIT Press, 2002.
- B. Schölkopf, J. C. Platt, J. C. Shawe Taylor, A. J. Smola, et R. C. Williamson. Estimating the Support of a High-Dimensional Distribution. *Neural Computation*, 13(7) :1443–1471, Juil. 2001.
- G. Schwarz. Estimating the dimension of a model. *The annals of statistics*, 6(2) :461–464, 1978.

- F. Sha et F. Pereira. Shallow parsing with conditional random fields. Dans *Proceedings of the 2003 Conference of the North American Chapter of the Association for Computational Linguistics on Human Language Technology-Volume 1*, pages 134–141. Association for Computational Linguistics, 2003.
- J. S. Shaffer et G. H. Alexander. Commercial Vehicle Brake Testing - Part 1 : Visual Inspection Versus Performance-Based Test. Technical report, SAE Technical Paper 952671, Nov. 1995a.
- S. J. Shaffer et G. H. Alexander. Commercial vehicle brake testing - Part 2 : Preliminary results of performance-based test program. Technical report, SAE Technical Paper 952672, Nov. 1995b.
- J. Shawe-Taylor et N. Cristianini. *Kernel Methods for Pattern Analysis*. Cambridge University Press, 2004.
- W. A. Shewhart. *Economic Control of Quality Manufactured Product*. D. Van Nostrand, Princeton, NJ, 1931.
- J.-J. H. Shiau, H.-L. Huang, S.-H. Lin, et M.-Y. Tsai. Monitoring nonlinear profiles with random effects by nonparametric regression. *Communications in Statistics-Theory and Methods*, 38(10) :1664–1679, 2009.
- A. N. Shiryaev. The problem of the most rapid detection of a disturbance of a stationary regime. *Soviet Math. Dokl.*, 2 :795–799, 1961.
- A. N. Shiryaev. On Optimum Methods in Quickest Detection Problems. *Theory of Probability and Its Applications*, 8, 1963.
- D. Siegmund et E. S. Venkatraman. Using the Generalized Likelihood Ratio Statistic for Sequential Detection of a Change-Point. *The Annals of Statistics*, 23(1) :255–271, 1995.
- J. A. Silmon. *Operational industrial fault detection and diagnosis : railway actuator case studies*. PhD thesis, University of Birmingham, Juil. 2009.
- B. Silverman. *Density Estimation for Statistics and Data Analysis*. Monographs on Statistics and Applied Probability. Taylor & Francis, 1986.
- A. E. Smith, D. W. Coit, et C. W. McCullers. Reliability improvement of airport ground transportation vehicles using neural networks to anticipate system failure. Dans *Reliability and Maintainability Symposium*, pages 74–79. Bombardier Transportation, Août 2002.
- A. E. Smith, D. W. Coit, et Y.-C. Liang. Neural Network Models to Anticipate Failures of Airport Ground Transportation Vehicle Doors. *Automation science and engineering, IEEE transactions on*, 7 :183–188, 2010.

- D. B. Sonawane, K. Narayan, V. S. Rao, et S. C. Subramanian. Model-based analysis of the mechanical subsystem of an air brake system. *International Journal of Automotive Technology*, 12(5) :697–704, Oct. 2011.
- B. Song et J. Hedrick. *Dynamic Surface Control of Uncertain Nonlinear Systems : An LMI Approach*. Communications and Control Engineering. Springer, Première édition, 2011.
- B. Song et J. K. Hedrick. Design and Experimental Implementation of Longitudinal Control for Automated Transit Buses. Dans *American Control Conference IEEE*, pages 2751–2756, Juil. 2004.
- S. C. Subramanian, S. Darbha, et K. R. Rajagopal. Modeling the Pneumatic Subsystem of an S-cam Air Brake System. *Journal of Dynamic Systems, Measurement, and Control*, 126(1) :36–46, 2004.
- S. R. C. Subramanian. *A diagnostic system for air brakes in commercial vehicles*. PhD thesis, Texas A&M University, Mai 2006.
- T. Sukchotrat, S. B. Kim, et F. Tsung. One-class classification-based control charts for multivariate process monitoring. *IIE Transactions*, 42(2) :107–120, Nov. 2009.
- R. Sun et F. Tsung. A kernel-distance-based multivariate control chart using support vector methods. *International Journal of Production Research*, 41(13) :2975–2989, 2003.
- A. Tartakovsky et V. Veeravalli. General asymptotic bayesian theory of quickest change detection. *Theory of Probability & Its Applications*, 49(3) :458–497, 2005.
- A. G. Tartakovsky. Multidecision Quickest Change-Point Detection : Previous Achievements and Open Problems. *Sequential Analysis*, 27(2) :201–231, 2008.
- A. G. Tartakovsky et G. V. Moustakides. State-of-the-art in bayesian changepoint detection. *Sequential Analysis*, 29(2) :125–145, 2010.
- A. G. Tartakovsky, B. L. Rozovskii, et K. Shah. A Nonparametric Multichart CUSUM Test for Rapid Intrusion Detection. Dans *Proceedings of Joint Statistical Meetings*, Minneapolis, MN, Août 2005.
- A. G. Tartakovsky, I. V. Nikiforov, et M. Basseville. *Sequential Analysis : Hypothesis Testing and Change-Point Detection*. Chapman & Hall/CRC - Monographs on Statistics & Applied Probability, Première édition, Livre en préparation.
- D. M. Tax. *One-class classification*. PhD thesis, Delft University of Technology, 2001.
- S. Theodoridis et K. Koutroumbas. *Pattern Recognition*. Academic Press, Inc., Orlando, FL, USA, Quatrième édition, 2008.

- C. Thomaz, D. Gillies, et R. Feitosa. Using Mixture Covariance Matrices to Improve Face and Facial Expression Recognitions. Dans *Audio and Video Based Biometric Person Authentication*, volume 2091 of *Lecture Notes in Computer Science*, pages 71–77. Springer Berlin Heidelberg, 2001.
- R. Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, pages 267–288, 1996.
- D. M. Titterton. Recursive Parameter Estimation Using Incomplete Data. *Journal of the Royal Statistical Society. Series B (Methodological)*, 46(2) :257–267, 1984.
- M. A. Traoré. *Supervision adaptative et pronostic de défaillance pour la maintenance prévisionnelle de systèmes évolutifs complexes*. PhD thesis, Université Lille 1, Déc. 2010.
- F. Tsung et K. Wang. *Adaptive Charting Techniques : Literature Review and Extensions*, volume 9. Springer-Verlag Berlin Heidelberg, 2010.
- D. Van Order, D. Skorupski, R. Stinebiser, et R. Kreeb. Fleet Study of Brake Performance and Tire Pressure Sensors. Technical report, US Departement of Transportation, Juil. 2009.
- V. Vapnik. *The Nature of Statistical Learning Theory*. Statistics for Engineering and Information Science. Springer, Première édition, 1999.
- G. Verdier. *Détection Statistique de Rupture de Modèle dans les Systèmes Dynamiques - Application à la Supervision de Procédés de Dépollution Biologique*. PhD thesis, Université Montpellier II, Nov. 2007.
- S. Verron. *Diagnostic et surveillance des processus complexes par réseaux bayésiens*. PhD thesis, Université d’Angers, Déc. 2007.
- P. Vieu et F. Ferraty. *Nonparametric functional data analysis*. Springer Science+ Business Media, 2006.
- L. Vlacic, M. Parent, et F. Harashima. *Intelligent Vehicle Technologies*. Automotive Engineering Series. Butterworth-Heinemann, Première édition, 2001.
- U. von Luxburg. Clustering stability : An overview. *Foundations and Trends® in Machine Learning*, 2(3) :235–274, 2010.
- WABCO. Air Braking Systems for Agricultural and Forestry Vehicles - Error Detection, Seconde édition. Technical report, 2007.
- R. W. Weeks et J. J. Moskwa. Automotive Engine Modeling for Real-Time Control Using MATLAB/SIMULINK. Technical report, University of Wisconsin-Madison, Fév. 1995.
- A. Widodo et B.-S. Yang. Support vector machine in machine condition monitoring and fault diagnosis. *Mechanical Systems and Signal Processing*, 21(6) :2560–2574, Août 2007.

- A. Willsky et H. Jones. A generalized likelihood ratio approach to the detection and estimation of jumps in linear systems. *Automatic Control, IEEE Transactions on*, 21(1) : 108–112, 1976.
- W. H. Woodall. Controversies and contradictions in statistical process control. *Journal of Quality Technology*, 32 :341–350, 2000.
- W. H. Woodall. Profile monitoring. *Encyclopedia of Statistics in Quality and Reliability*, 2007.
- W. H. Woodall, D. J. Spitzner, D. C. Montgomery, et S. Gupta. Using control charts to monitor process and product quality profiles. *Journal of Quality Technology*, 36(3) :309–320, 2004.
- K. Yamanishi et J. I. Takeuchi. A unifying framework for detecting outliers and change points from non-stationary time series data. Dans *Proceedings of the eighth ACM SIGKDD international conference on Knowledge discovery and data mining*, KDD '02, pages 676–681, New York, NY, USA, 2002. ACM.
- T. Y. Yang et L. Kuo. Bayesian Binary Segmentation Procedure for a Poisson Process with Multiple Changepoints. *Journal of Computational and Graphical Statistics*, 10(4) :772–785, 2001.
- G. Yu, C. Zou, et Z. Wang. Outlier detection in functional observations with applications to profile monitoring. *Technometrics*, 54(3) :308–318, 2012.
- X. Yu, T. Shen, G. Li, et K. Hikiri. Regenerative braking torque estimation and control approaches for a hybrid electric truck. Dans *American Control Conference (ACC), 2010*, pages 5832–5837. IEEE, Juin 2010.
- A. Zeileis, C. Kleiber, W. Krämer, et K. Hornik. Testing and dating of structural changes in practice. *Computational Statistics & Data Analysis*, 44(1) :109–123, 2003.
- G. Zwingelstein. *Diagnostic des défaillances*. Traité des nouvelles technologies. Série Diagnostic et maintenance. Hermès Science Publication, Première édition, 2002.

Liste des figures

1.1	Évolution de deux processus au cours du temps avec un changement structurel.	3
1.2	Deux systèmes pneumatiques : frein et porte.	8
1.3	Localisation des principaux boîtiers électroniques embarqués	9
1.4	Architectures télématiques implémentée	10
1.5	Autobus articulé Citelis Euro IV devant le dépôt de la STRAV à Limeil-Brévannes	13
1.6	Diagramme simplifié d'un système de freinage hydraulique	16
1.7	Mécanismes de frein à tambour et à disque	17
2.1	Aperçu général des différentes politiques de maintenance	23
2.2	Principales étapes pour établir un diagnostic automatique à base de modèle analytique	25
2.3	Principales étapes pour établir un diagnostic automatique à base de reconnaissance des formes	27
2.4	Échantillon de 4000 points en deux dimensions. Estimation de la densité par un mélange de gaussiennes.	39
2.5	Échantillon de 25 courbes pendant une ouverture de porte.	40
2.6	Exemple de fonction logistique qui régit les transitions entre 5 segments.	44
2.7	Représentation graphique du modèle probabiliste RHLF	45
2.8	Cours du pétrole de janvier 1974 à octobre 2012	49
2.9	Sélection du modèle de régression univarié à partir des critères BIC et ICL, et le temps CPU	50
2.10	Segmentation des courbes en huit morceaux de degré trois.	51
3.1	Changement brusque à l'instant t_c détecté au temps t_a avec retard	57
3.2	Procédure de Shewhart appliquée au suivi des débits moyens annuels de la rivière du Nil.	59
3.3	Procédure du CUSUM appliquée au suivi des débits moyens annuels de la rivière du Nil.	61

3.4	Exemple de courbes ROC.	70
3.5	Intérêt d'une approche bivariée pour des données en dimension deux. . .	73
3.6	Représentation graphique d'un modèle HMM	82
4.1	Schéma de principe du système de freinage principal et circuit pneumatique d'un autobus Citelis	92
4.2	Validation expérimentale des hypothèses d'embrayage verrouillé et de glissement négligeable entre chaussée et roues	97
4.3	Diagramme du corps libre appliqué à un bus.	99
4.4	Validation expérimentale de l'estimation de R_q , le rapport de démultiplication totale	101
4.5	Modélisation spatiale de la ligne J et estimation de la pente	103
4.6	Justification expérimentale de la relation linéaire entre pression et couple de freinage	105
4.7	Modèle de régression linéaire entre la pression et le couple de freinage . .	106
4.8	Stratégie séquentielle de détection en 3 étapes pour la surveillance du système de freinage	106
4.9	CANversionTool : outil de conversion de données CAN	108
4.10	Extraction du couple de paramètres (K, σ) à partir des données CAN collectées sur l'ATON	109
4.11	Représentation spatiale de quelques plages de freinage significatives et distributions des couples (K, σ) représentant le fonctionnement des véhicules	111
4.12	Processus de dégradation du système de freinage	114
4.13	Illustration des cinq situations de dégradation du système de freinage pour le bus 484	115
4.14	Cartes multivariées de Shewhart et de CUSUM, et les courbes ROC pour les cinq situations de dégradation des plages de freinage du bus 484 . . .	117
5.1	Les principaux systèmes de porte louvoyante : intérieure, extérieure, ou coulissante.	122
5.2	Vérin double-effet	123
5.3	Description du système de porte pneumatique et instrumentation	124
5.4	Réducteur de Débit Unidirectionnel (RDU).	125
5.5	Simulation d'un nuage de courbes et la courbe moyenne.	132
5.6	Segmentation de dix courbes en quatre morceaux de degré 2.	133
5.7	Densité de probabilité du modèle de régression.	135
5.8	Simulation d'un nuage de courbes et la courbe moyenne.	138
5.9	Courbes moyennes estimées et segmentation temporelle de deux modèles appris en mode non/semi-supervisé.	139
5.10	Exemple d'un cycle d'ouverture / fermeture de porte, collecté le 12 juillet 2011	149

5.11	Sélection du modèle de régression multivarié à partir des critères BIC et ICL, et le temps CPU	151
5.12	Segmentation de 11 courbes de fermeture en bon fonctionnement, collectés le 12 juillet 2011	152
5.13	Évaluation graphique du détecteur proposé en fonction des taux de fausses alarmes et des délais moyens à la détection	155
5.14	Autobus standard Citelis et porte surveillée	156
5.15	Échantillon d'un flux de données et zoom sur un cycle d'ouverture/fermeture	158
5.16	Vingt courbes d'ouverture avec zones d'incertitudes et sélection de modèle.	160
5.17	Segmentation de 20 courbes d'ouverture en bon fonctionnement, collectées le 30 janvier 2012	160
5.18	Densité de signatures en bon fonctionnement et quelques dégradations provoquées	162
5.19	Densité de signatures en bon fonctionnement et quelques dégradations provoquées	163
5.20	Statistique du test séquentiel proposé pour la détection de défauts sur des données fonctionnelles. Le test est appliqué à des courbes de porte pneumatique mesurées de janvier à octobre 2012 et 35 anomalies ont été détectées.	165
5.21	Statistiques de localisation des défauts par segment.	167
5.22	Représentation des courbes mesurées entre janvier et octobre 2012.	168

Liste des tableaux

2.1	Surveillance d'un système : les connaissances nécessaires	24
4.1	Comparaison entre tambours et disques de freins	90
4.2	Comparaison d'un système de frein pneumatique/hydraulique	90
4.3	Quatre commandes de freinage pour un bus Citelis (Irisbus)	91
4.4	Notations adoptées pour la modélisation longitudinale d'un autobus	94
4.5	Ratios théoriques d'une boîte de vitesse ZF en fonction du rapport engagé	100
4.6	Collection des plages de freinage pour chaque bus	110
4.7	Cinq situations de dégradations simulées pour le système de freinage	113
4.8	Répartition des plages de bon fonctionnement et des plages de freinage dégradées	114
4.9	AUC des cartes multivariées de Shewhart et CUSUM	118
4.10	Taux de vrais positifs et de faux positifs (en %) de la carte MShewhart et MCUSUM avec seuil théorique	118
4.11	Taux de vrais positifs et de faux positifs (en %) de la carte MShewhart et MCUSUM avec seuil théorique	119
5.1	Quelques indicateurs de performance sur le modèle appris	134
5.2	Paramètres du mélange régressif estimés dans un mode non supervisé sur un ensemble des courbes simulées	135
5.3	Quelques indicateurs sur le modèle appris en mode semi-supervisé	139
5.4	Paramètres des modèles avant et après changement utilisés pour simuler des courbes réalistes	153
5.5	Description des données collectées pour la surveillance des portes	157

Liste des Algorithmes

2.1	Pseudo-code de l'algorithme EM appliqué à un mélange gaussien.	38
2.2	Pseudo-code du modèle RHLP univarié	48
5.1	Pseudo-code du modèle de régression à processus latent dans une version multidimensionnelle	130
5.2	Pseudo-code du modèle de régression à processus latent multivarié dans un mode semi-supervisé	137
5.3	Pseudo-code de la stratégie proposée pour la détection séquentielle sur des données fonctionnelles	143
5.4	Pseudo-code de la stratégie proposée pour la détection séquentielle et la localisation de défauts sur des données fonctionnelles	145
5.5	Pseudo-code de la stratégie proposée pour la détection séquentielle et la localisation de défauts sur des données fonctionnelles, avec estimation des seuils de détection/localisation	148

Liste des publications

CONFÉRENCES INTERNATIONALES

- N. Cheifetz, A. Samé, P. Aknin, et E. de Verdalle. A pattern recognition approach for anomaly detection on buses brake system. Dans *Intelligent Transportation Systems (ITSC), 2011 14th International IEEE Conference*, pages 266 –271, Washington DC, USA, Octobre 2011.
- N. Cheifetz, A. Samé, P. Aknin, et E. de Verdalle. A CUSUM-like approach for online change-point detection on bus door systems. Dans *CM 2012 and MFPT 2012 - The Ninth International Conference on Condition Monitoring and Machinery Failure Prevention Technologies*, London, UK, Juin 2012a.
- N. Cheifetz, A. Samé, P. Aknin, et E. de Verdalle. A sequential testing approach for change-point detection on bus door systems. Dans *Intelligent Transportation Systems (ITSC), 2012 15th International IEEE Conference*, pages 1846–1851, Anchorage, AK, USA, Septembre 2012b.
- N. Cheifetz, A. Samé, P. Aknin, et E. de Verdalle. A CUSUM approach for online change-point detection on curve sequences. Dans *the 20th European Symposium on Artificial Neural Networks, Computational Intelligence and Machine Learning (ESANN)*, pages 399–404, Bruges, BE, Avril 2012c.
- N. Cheifetz, A. Samé, P. Aknin, E. de Verdalle, et D. Chenu. A sequential testing procedure for multiple change-point detection in a stream of pneumatic door signatures. Dans *the 12th International Conference on Machine Learning and Applications (ICMLA'13)*, Miami, Florida, USA, Décembre 2013.

CONFÉRENCE NATIONALE AVEC COMITÉ DE LECTURE

- N. Cheifetz, A. Samé, et P. Aknin. Diagnostic de sous-systèmes d'autobus à base de reconnaissance des formes. Dans Actes IFSTTAR, editor, *Journée Des Doctorants SPI-STIC de l'IFSTTAR 2011*, Villeneuve-d'Ascq, FR, Juin 2011. Prix de la meilleure présentation.

CONFÉRENCES NATIONALES SANS COMITÉ DE LECTURE

- N. Cheifetz. Diagnostic d'un système de freinage d'autobus à base de reconnaissance des formes. Dans *Journée des doctorants ED-MSTIC 2011*, Université Paris-Est - Cité Descartes, Noisy-Le-Grand, FR, Juin 2011.
- N. Cheifetz, A. Samé, P. Aknin, X. Martinez, et E. de Verdalle. Détection séquentielle d'anomalies pour le suivi de portes d'autobus. Dans *Journée GdR S3/SEE/SAFFE-GIS 3SGS*, ENSAM - Paris, FR, Janvier 2012. Poster.

AUTRES CONTRIBUTIONS

- N. Cheifetz. Analyse en facteurs indépendants pour le diagnostic d'un composant de l'infrastructure ferroviaire dans un cadre semi-supervisé. Master's thesis, Paris 6 - Université Pierre et Marie Curie (UPMC), Septembre 2009a.
- N. Cheifetz. *IFA for the Diagnosis of a Railway Infrastructure Device in a Semi-Supervised Learning*. Cagliari, IT, The Analysis of Patterns, Summer school - Troisième édition, Septembre 2009b. Interview on videolectures.net : http://videolectures.net/aop09_cheifetz_iftd/.
- N. Cheifetz et Consortium EBSF. Optimization of alarm thresholds values in predictive maintenance. Technical report, International Association of Public Transport (UITP), Décembre 2012. Livrable Savings du projet EBSF.

Abstract

Detection and classification of temporal CAN signatures to support maintenance of public transportation vehicle subsystems

Nowadays, transportation systems need to be operated with a high level of security and reliability. There is an increasing need for tools to monitor and support the preventive maintenance of some critical subsystems. Within the EBSF (European Bus System of the Future, FP7) project, this PhD work addresses the monitoring issue on two bus subsystems strongly impacting the availability of the vehicles and their maintenance costs: the brake and door systems. The pattern recognition approach, based on the analysis of operating data, is favored in this work. The detection methodology dedicated to the brake systems uses a dynamic model of the vehicle to extract physical features related to the brakes. These features are used as inputs to multivariate control charts which are used to detect any anomalies. The detection strategy dedicated to the door systems directly handles the functional data, which are collected by embedded sensors during opening and closing cycles. A sequential testing of generalized likelihood ratio, based on a parametric regression model with latent segmentation, is proposed to solve the detection issue. An approach to support the diagnosis task, based on local detectors, is also implemented. The results gained from a real dataset collected on a fleet of buses allow to highlight the effectiveness of the proposed methods for the brake study as well as the door study.

Keywords: Industrial diagnosis, change-point detection, control charts, sequential hypothesis testing, curve segmentation, finite mixture models, EM algorithm, longitudinal vehicle modeling, air brake system, pneumatic doors.
