



Surveillance temps-réel des systèmes Homme-Machine. Application à l'assistance à la conduite automobile

Miguel Gonzalez-Mendoza

► To cite this version:

Miguel Gonzalez-Mendoza. Surveillance temps-réel des systèmes Homme-Machine. Application à l'assistance à la conduite automobile. Informatique [cs]. INSA de Toulouse, 2004. Français. NNT : . tel-00336732

HAL Id: tel-00336732

<https://theses.hal.science/tel-00336732>

Submitted on 5 Nov 2008

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

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

THÈSE

Présentée au

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

En vue de l'obtention du titre de

Docteur de l'Institut National des Sciences Appliquées de Toulouse

Spécialité

Systèmes Automatiques

par

Miguel GONZALEZ MENDOZA

Surveillance temps-réel des systèmes Homme–
Machine. Application à l'assistance à la conduite
automobile

Soutenue le 16 Juillet 2004, devant le jury :

Directeurs de thèse :

André TITLI

Professeur à l'INSA de Toulouse

Neil HERNANDEZ-GRESS

Professeur à l'ITESM CEM

Président :

Daniel ESTEVE

Directeur de recherche LAAS-CNRS

Rapporteurs :

Dominique MEIZEL

Professeur à l'ENSIL

José Armando HERRERA-CORRAL

Professeur à l'ITESM CT

Examineur :

Bruno JAMMES

Maître de conférences à l'UPS

Invités :

Serge BOVERIE

Dr. Siemens VDO Automotive S.A.S.

Jérôme THOMAS

Dr. ACTIA

A mon épouse Berenice et à nos futurs Enfants

A mes parents Mayis et Miguel

A mes frères Mauricio et Milton

A toute ma Famille et à mes Amis...

Avant Propos

Les travaux présentés dans cette thèse sont le résultat des recherches réalisées dans le Laboratoire d'Analyse et d'Architecture des Systèmes du Centre National de la Recherche Scientifique, LAAS-CNRS, dans les groupes de recherche Diagnostic, Supervision et Conduite qualitatifs, DISCO, et Microsystèmes et Intégration des Systèmes, MIS. Je remercie à Monsieur Jean-Claude LAPRIE et à Monsieur Malik GHALLAB, ancien et actuel directeur du LAAS, de m'avoir accueilli dans cet établissement.

J'exprime toute ma gratitude et reconnaissance à mon directeur de thèse Monsieur André TITLI, Professeur à l'Institut National des Sciences Appliquées de Toulouse, INSA, et cadre Scientifique au LAAS-CNRS, pour sa confiance et son amitié toute au long de ces années de travail. Sa forme très particulière d'écouter, de discuter et de proposer des idées claires et précises, a su diriger sagement les mouvements de ce jeune chercheur, en me donnant, en même temps, un grand niveau d'autonomie et une manière plus élargie d'apprécier le monde. Je lui remercie également de m'avoir maintenu sur la voie de la représentation à partir de la connaissance, notamment sur les systèmes d'inférence flous, pour lesquels je voyais moins d'intérêt que pour la représentation à partir des données.

J'adresse tous mes remerciements à Monsieur Neil HERNANDEZ GRESS, Professeur à l'Institut Technologique et d'Etudes Supérieures de Monterrey, ITESM, campus Estado de Mexico et codirecteur de thèse, de m'avoir soutenu tant au niveau scientifique comme au niveau personnel. Nos interminables discussions sur l'apprentissage, les réseaux de neurones, ainsi que sur la surveillance du conducteur automobile et d'autres innombrables thèmes scientifiques et personnelles, ont beaucoup apporté au déroulement de ces travaux. J'espère que nos travaux conjoints ouvriront une vaste voie de recherche et que l'on pourra développer, de manière toujours conjointe, au Mexique.

Je remercie tout particulièrement Monsieur Daniel ESTEVE, Directeur de Recherche CNRS, pour avoir accepté la présidence de mon jury de thèse. Je lui remercie également pour nos innombrables discussions sur l'hypovigilance du conducteur et de m'avoir témoigné toute sa confiance dans le déroulement du projet Européen AWAKE et du projet National PREDIT. J'admire sa versatilité pour aborder n'importe quel sujet scientifique et pour gérer une quantité considérable de travail... « Chapeau ! ».

Je suis reconnaissant envers Monsieur Dominique MEIZEL, Professeur à l'Ecole Nationale Supérieure d'Ingénieurs de Limoges, ENSIL, ancien responsable du Groupe de Recherche Coopération Homme Machine pour l'Aide à la Conduite, CHMAC, pour le professionnalisme avec lequel il a jugé mon travail de thèse, en tant que rapporteur. Ses remarques et commentaires, lors du workshop AWAKE à Paris et du manuscrit lui-même, m'ont aidé à l'amélioration des travaux.

Mes remerciements vont aussi à Monsieur José Armando HERRERA CORRAL, Professeur à l'Institut Technologique et d'Etudes Supérieures de Monterrey, ITESM, campus Toluca, pour avoir accepté la lourde tâche de rapporteur. Ses remarques et nos discussions sur l'Intelligence Artificielle, ont contribué à l'amélioration du manuscrit. C'est grâce à lui que j'ai eu un premier aperçu de la Recherche Scientifique, du LAAS et de la France en générale. Je lui remercie pour son amitié, ses conseils et son

soutien dans toutes les démarches administratives devant l'ITESM, pour le Mastaire, et le CONACyT, pour le DEA et le Doctorat.

Je suis très reconnaissant envers Monsieur Bruno JAMMES, Maître de conférences à l'Université Paul Sabatier à Toulouse, pour l'ensemble des travaux conjoints afin de mettre au point le module de diagnostic de l'hypovigilance du conducteur automobile LAAS. Nos journées entières à vérifier le fonctionnement de l'analyse en ondelettes et la construction des caractéristiques pertinentes au diagnostic de l'hypovigilance, ont donné ses fruits. Je lui remercie pour son amitié, son soutien, sa confiance et l'intérêt qu'il a porté à juger ce manuscrit.

Un grand merci à Monsieur Serge BOVERIE, responsable des Développements Avancés chez SIEMENS VDO Automotive et responsable du projet PREDIT, pour le temps qu'il a passé à étudier mon manuscrit et pour les activités en commun qu'il a géré pour les projets AWAKE et PREDIT. Il a su toujours guider les développements pour qu'ils soient les plus compatibles possible avec les contraintes industrielles.

Je remercie infiniment à Monsieur Jérôme THOMAS, responsable des Développements des Diagnostics chez ACTIA S. A., d'avoir accepté mon invitation pour faire partie de ce jury de thèse. Je lui remercie également pour toutes ses remarques lors du développement pour le fonctionnement temps réel du module de diagnostic LAAS. Son aide a été précieuse, particulièrement lors des quelques essais sur le camion Actros, dans le cadre du projet AWAKE, auxquels il m'a invité, afin de vérifier le fonctionnement de notre algorithme de diagnostic en ligne et l'améliorer.

Au LAAS j'exprime ma reconnaissance envers :

- Monsieur Joseph AGUILAR-MARTIN, directeur du Laboratoire Européen LEA-SICA et ancien directeur du groupe DISCO, de m'avoir toujours aidé et supporté pour réaliser les séjours en Catalogne et participer aux différentes conférences. Son grand enthousiasme, bonne humeur et disponibilité a créé une bonne ambiance entre nos camarades.
- Madame Anne-Marie GUE, directeur du groupe MIS, pour m'avoir toujours aidé à effectuer les missions au cours de ces années pour les projets.
- Madame Louise TRAVE-MASSUYES, directeur du groupe DISCO, de m'avoir également soutenue lors des différentes missions et de m'avoir invité à participer au réseau MONET. Je lui remercie également pour sa disponibilité et ces commentaires par rapport au diagnostic.
- Monsieur Alfredo SANTANA DIAZ, pour nos collaborations sur la détection de la baisse de la vigilance du conducteur et pour m'avoir introduit au monde des ondelettes.
- Monsieur Sébastien BEYOU, stagiaire de DEA, pour sa contribution à ce travail de thèse.
- Mes collègues du groupe DISCO et MIS qui ont contribué, avec son bon humeur à rendre plus convivial mon séjour au LAAS : Tatiana KEMPOWSKY, Juan Carlos HAMMON, Codruta BILEGAN, Thierry MIQUEL, Victor Hugo GRISALES, Claudia ISAZA, Antonio ORANTES, Héctor HERNANDEZ, Jean Charles ATINE, Aïmed MOKHTARI, Sébastien REGIS, Patrick ABGRALL, Edouard DIEZ, ...
- Un grand merci au service de documentation, au service de reproduction et au magasin pour leur constante disponibilité et amabilité. Je tiens également à remercier la gestion du personnel et la gestion financière pour leur amabilité et efficacité.
- A mes camarades et amis du DEA Julien PETTRE et Patrice LANGOUET.
- A tous ceux dont le nom m'échappe et que j'ai pu côtoyer dans ce grand laboratoire et dont la liste serait trop longue.

Ce travail a été effectué dans le cadre du programme Européen AWAKE et du projet national PREDIT. Les idées émanées des longues journées de discussion sont intégrées dans cette thèse. Je tiens à remercier le consortium Européen et à nos partenaires de m'avoir permis de participer à ces projets. Une pensée spéciale pour :

- Alain GIRALT, Dr. SIEMENS VDO Automotive, pour nos travaux communs et nos longues discussions sur le diagnostic de l'hypovigilance du conducteur automobile. Je lui remercie également pour l'amitié et la confiance qu'il m'a témoignée.
- Evangelos BEKIARIS, chercheur de l'Institut Hellénique du Transport en Grèce et responsable du projet Européen AWAKE, pour la confiance et son support au cours de ces trois années de travail conjoint.
- Alain MUZET, directeur du Centre de Physiologie Appliquée du CNRS, CEPA-CNRS, pour son amitié, sa confiance et ses remarques afin de construire une référence physiologique fiable pour valider le système de diagnostic.
- Stella NIKOLAU, Aris POLYCHRONOPOULUS, Björn PETERS, Anna ANAUD, Claudio ANTONELLO, Karel BROOKHUIS, Dick DEWAARD, Rino BROUWER et tous nos autres partenaires qui ont rendu le travail plus agréable et efficace. Sans eux, AWAKE n'aurais jamais été possible.

J'exprime toute ma gratitude à mes amis et collègues à Toulouse. Ils ont été un fort soutien pour Berenice et pour moi pendant ces années en France ; Elvia PALACIOS, Elsa RUBIO, Efraín LOPEZ, Claudia JARAMILLO, Antonio MARIN et Andrea GALLARDO, César ZAMILPA et Denise Da LUZ, Gerardo Alejandro VELAZQUEZ et Ana HERNANDEZ, Juan Carlos ZUÑIGA et Sofía BERUMEN, Gustavo ARECHAVALETA et Margarita GUTIERREZ, Alfredo SANTANA et Marquidea PACHECO, Mario Alain GONZALEZ et Nazareth TELLEZ, Saul POMARES et Lietta VALDES, Roberto REYNA et Daniela DRAGOMIRESCU, Neil HERNANDEZ et Teresa OSORIO, Martín MOLINA et Yusdivia ROMAN, Efraín JAIME et Martha HINOJOSA.

Merci aux millions de mexicains qui, par l'intermédiaire du CONACyT et de la SEP, ont participé avec leur soutien financier.

Je ne pourrai jamais finir de remercier ma *chiquita*, Bere pour m'avoir soutenu et supporté, particulièrement dans les moments difficiles. Sa détermination, son courage et sa compréhension ont été pour moi une source inépuisable de motivation pour continuer. Nous venons d'aboutir une première étape dans notre vie et nous commençons une autre avec l'espoir de cimenter un avenir plus solide pour nos futurs enfants et de pouvoir contribuer, dans la mesure du possible, avec notre petit grain de sable pour un monde meilleur.

Finalement, merci à mes parents Mayis et Miguel de m'avoir soutenu et guidé dans tous les moments cruciaux dans ma vie. Même si la vie est dure, ils nous ont fait voir, à mes frères et moi, qu'elle est pleine de satisfactions, si on travail dure et honnêtement. Ils sont un exemple à suivre et je leur en remercie infiniment. Merci à mes frères Mauricio et Milton. Ils ont toujours été un grand support pour moi, une compagnie dans la vie que j'apprécie énormément et que j'espère pouvoir partager avec nos propres familles.

Table des Matières

Table de Matières.....	v
Table de Figures.....	ix
Abréviations.....	xiii
Introduction.....	1
Chapitre I	
<i>Facteurs impliqués dans la surveillance de l'opérateur humain dans les systèmes homme-machine</i>	
I.1 Introduction.....	6
I.2 Les systèmes complexes.....	6
I.2.1 L'impact de l'automatisation dans les systèmes complexes.....	7
I.2.1.1 Avantages de l'automatisation.....	7
I.2.1.2 Inconvénients de l'automatisation.....	7
I.2.2 Les interactions homme-système.....	8
I.2.2.1 Coopération homme-système.....	8
I.2.2.2 Partage de tâches.....	10
I.3 Défaillances dans les systèmes homme-machine.....	11
I.3.1 La performance humaine dans l'interaction homme-machine.....	11
I.3.2 Charge de travail.....	12
I.3.3 Conscience de la situation.....	12
I.3.4 Surveillance de système.....	13
I.3.5 Travail d'équipe.....	14
I.3.6 Confiance.....	14
I.3.7 Utilisabilité et acceptation.....	15
I.3.8 Erreur humaine.....	15
I.4 Métriques sur la performance humaine.....	16
I.4.1 Mesures physiologiques.....	16
I.4.1.1 L'activité cérébrale.....	17
I.4.1.2 L'activité cardiaque.....	18
I.4.1.3 L'activité oculaire.....	18

I.4.2	Mesures subjectives.....	19
I.4.2.1	Autoévaluation instantanée.....	20
I.4.2.2	Technique subjective d'évaluation de charge de travail.....	20
I.4.2.3	Indice de charge de la tâche.....	20
I.4.2.4	Echelle d'endormissement Karolinska	20
I.4.3	Mesures de performance	21
I.4.3.1	Métriques de performance de la charge de travail.....	21
I.4.3.2	Mesures de la charge de tâche	22
I.5.	Surveillance des systèmes homme-machine	23
I.5.1	Les systèmes d'assistance intelligents.....	24
I.6	Conclusion.....	24

Chapitre II

Le prétraitement des signaux, l'analyse temporelle et fréquentielle et extraction de caractéristiques

II.1	L'intérêt du prétraitement.....	28
II.2	Analyse en temps continu.....	28
II.2.1	Filtrage linéaire stationnaire.....	29
II.2.2	La transformée de Fourier	29
II.3	La révolution discrète	30
II.3.1	L'échantillonnage des signaux analogiques	31
II.3.2	Filtres discrets	32
II.3.2.1	Filtrage linéaire stationnaire discret.....	32
II.3.2.2	Série de Fourier	33
II.3.3	Signaux finis.....	34
II.3.3.1	Convolutions circulaires	34
II.3.3.2	Transformée de Fourier discrète.....	34
II.3.3.3	Transformée de Fourier rapide	35
II.3.3.4	Convolutions rapides	36
II.4	L'analyse temps-fréquence	36
II.4.1	Atomes temps-fréquence	36
II.4.2	Transformée de Fourier à fenêtre	37
II.4.2.1	Choix de la fenêtre.....	38
II.4.2.2	Transformée de Fourier à fenêtre rapide	39
II.4.3	Transformée en ondelettes.....	39
II.4.3.1	Transformée en ondelettes discrète	41
II.4.3.2	Transformée en ondelettes dyadiques.....	41
II.4.4	Analyse en ondelettes.....	42
II.4.4.1	Régularité	42
II.4.4.2	Maxima de la transformée en ondelettes et détection de singularités.....	43
II.4.5	Bases d'ondelettes	44
II.5	Extraction de caractéristiques	45
II.5.1	Analyse statistique.....	45
II.5.1.1	Statistiques du premier ordre	45
II.5.1.2	Statistiques du deuxième ordre.....	46
II.5.2	Analyse spectrale.....	46
II.5.3	Fenêtres d'analyse	47
II.5	Conclusion	47

Chapitre III

Techniques d'apprentissage et d'intelligence artificielle pour la modélisation de systèmes complexes

III.1	Introduction	51
III.2	Modélisation à partir des données	51
III.2.1	L'apprentissage statistique	51
III.2.1.1	Les machines d'apprentissage	52
III.2.1.2	Les fonctions de perte et la minimisation du risque	52
III.2.1.3	Les trois problèmes principaux d'apprentissage	52
III.2.1.4	Les principes d'induction	52
III.2.2	Les machines d'apprentissage classiques : l'approche neuronale	55
III.2.2.1	Les machines linéaires	55
III.2.2.2	Les machines non-linéaires	59
III.2.2.3	L'estimation de la densité	63
III.2.2.4	Remarques	66
III.2.3	Les machines à vecteurs de support SVM	67
III.2.3.1	Règles de décision non linéaires	67
III.2.3.2	SVM pour la reconnaissance de formes	69
III.2.3.3	SVM pour la régression	71
III.2.3.4	SVM pour l'estimation de la densité	73
III.2.3.5	Remarques	74
III.2.4	Le mécanisme d'apprentissage des SVM	74
III.2.4.1	Caractéristiques du problème d'optimisation quadratique SVM	74
III.2.4.2	Les algorithmes d'optimisation des problèmes quadratiques	75
III.2.4.3	Implémentation des SVM	78
III.2.4.4	Application des SVM sur des bases de données	80
III.3	Modélisation à partir de la connaissance	81
III.3.1	Généralités sur l'intelligence artificielle	81
III.3.1.1	Agents intelligents	82
III.3.1.2	Systèmes Experts	82
III.3.2	Les systèmes d'inférence flous	83
III.3.3.1	Eléments d'un système d'inférence flou	83
III.3.3.2	Généralités sur la logique floue	84
III.3.3.3	Le mécanisme de fuzzification, d'inférence et de défuzzification	86
III.3.3.4	Types de modèles flous	89
III.3.3.5	Génération de règles à partir de l'expertise	94
III.4	Modélisation à partir de la connaissance et les données	94
III.4.1	Modèles flous conventionnels avec apprentissage	94
III.4.1.1	Choix du nombre d'ensembles flous	95
III.4.1.2	Paramétrisation des règles	95
III.4.2	Modèles neuro-flous	96
III.4.2.1	Les algorithmes de groupage par données (clustering)	96
III.4.2.3	Le groupage SVM	98
III.4.2.4	Initialisation des algorithmes de groupage	99
III.4.3	Identification des modèles TS pour des systèmes MIMO	100
III.4.3.1	Structure d'un modèle flou TS pour des systèmes MIMO	100
III.4.3.2	Méthode d'identification	101
III.4.3.3	Exemple d'identification	102
III.5	Conclusion	104

Chapitre IV

Le problème de la surveillance du conducteur automobile. Les projets AWAKE et PREDIT

IV.1	Introduction	107
IV.1.1	Les enjeux de la sécurité routière.....	107
IV.2	Les systèmes d'assistance à la conduite	109
IV.2.1	Taxonomie des systèmes d'assistance à la conduite	109
IV.2.3	La tâche de conduite	110
IV.2.3.1	Les capteurs de position latérale et le contrôle associé	111
IV.2.3.2	Les capteurs de distance longitudinale et le contrôle associé.....	112
IV.3	La surveillance du conducteur automobile	112
IV.3.1	Généralités	112
IV.3.1.1	Mesures directes du niveau de vigilance	112
IV.3.1.2	Mesures directes de la performance de la conduite.....	114
IV.3.2	Bilan sur les systèmes de surveillance du véhicule/conducteur.....	116
IV.3.2.1	Quelques systèmes de surveillance de la conduite et de l'état du conducteur	117
IV.3.2.2	L'analyse synthétique.....	124
IV.4	Les projets de recherche	125
IV.4.1	Le programme Européen AWAKE (Septembre 2001 – Septembre 2004)	126
IV.4.2	Le projet national Hypovigilance du Conducteur – PREDIT (Mars 2002 – Mars 2004).....	126
IV.4.3	Architectures de fonctionnement	127
IV.4.3.1	L'architecture AWAKE	127
IV.4.3.2	L'architecture PREDIT	128
IV.5	Les essais	129
IV.5.1	Moyens expérimentaux	129
IV.5.2	Les campagnes d'essais	131
IV.5.2.1	Les essais sur démonstrateurs	131
IV.5.2.2	Les essais sur simulateurs	133
IV.6	Développement du module de diagnostic de situations à risque	134
IV.6.1	Architecture du module de diagnostic de situations à risque.....	134
IV.6.2	Module de Diagnostic Événementielle basé sur le temps de sortie de voie.....	135
IV.6.3	EDM basé sur vibreur de bord de route adaptatif.....	135
IV.6.4	Comparaisons des différentes approches EDM	138
IV.7	Détection de l'hypovigilance du conducteur automobile	141
IV.7.1	Architecture du module de détection de l'Hypovigilance.....	142
IV.7.1.1	Architecture du HDM performance	142
IV.7.1.2	Analyse temporelle et fréquentielle des signaux	143
IV.7.1.3	Fusion des caractéristiques.....	146
IV.7.1.3	Diagnostic cumulé.....	147
IV.7.2	Evaluation du HDM performance.....	148
IV.7.2.1	Mesures de référence de la vigilance	148
IV.7.2.2	Validation.....	149
IV.7.2.3	Résultats	151
IV.8	Conclusion	162
Conclusion Générale.....		163
Bibliographie		167

Table de Figures

Chapitre I

Facteurs impliqués dans la surveillance de l'opérateur humain dans les systèmes homme-machine

Figure I.1 : Performance d'un système homme-machine comme une fonction des niveaux d'autonomie.	10
Figure I.2 : Métriques sur la performance humaine : mesures physiologiques, mesures subjectives et mesures de performance.	16
Figure I.3 : Charge de travail et performance en 6 régions. Dans la région D (désactivation) l'état de l'opérateur est affecté. Dans la région A2, la performance est optimale, l'opérateur peut se « débrouiller » facilement avec les requêtes de la tâche. Dans les régions A1 et A3, la performance n'est pas affectée mais l'opérateur doit exercer plus d'effort pour maintenir le niveau de performance. Dans la région B, ce n'est plus possible et la performance se dégrade, pendant que dans la région C la performance est dans un niveau minimum : l'opérateur est surchargé.	23
Figure I.4 : Structure d'un système homme-machine surveillé.	23

Chapitre II

Le prétraitement des signaux, l'analyse temporelle et fréquentielle et extraction de caractéristiques

Figure II.1 : Schéma d'un système de perception et de prétraitement du signal.	28
Figure II.2 : (a) Signal x et sa transformée de Fourier $\hat{x}$. (b) Un échantillonnage uniforme de rend sa transformée de Fourier périodique. (c) Passe bas idéal. (d) Le filtrage de (b) par (c) reconstitue x .	31
Figure II.3 : (a) Signal x et sa transformée de Fourier $\hat{x}$. (b) Repliement spectral dû à un recouvrement des $\hat{x}(\omega - 2k\pi/T)$ pour différentes valeurs k (en pointillés). (d) Le filtrage de (b) par (c) génère un signal à base fréquence qui diffère de x .	32
Figure II.4 : Boîte de Heisenberg représentant un atome temps-fréquence.	37
Figure II.5 : Boîte de Heisenberg de deux atomes de Fourier à fenêtre $g_{u,\xi}$ et $g_{v,\gamma}$.	38
Figure II.6 : a) L'étalement de g est mesuré par sa largeur de bande $\Delta\omega$ et par l'amplitude de ses lobes latéraux, situées en $\omega = \pm \omega_0$. b) Fenêtres g de support $[-1/2, 1/2]$: Hamming, Gaussienne, Hanning et Blackman.	39
Figure II.7 : Boîte de Heisenberg de deux atomes de Fourier à fenêtre $g_{u,\xi}$ et $g_{v,\gamma}$.	40
Figure II.8 : Ondelette « chapeau mexicain » pour $\sigma=1$ et sa transformée de Fourier.	41
Figure II.9 : Transformée en ondelettes réelle calculée avec une ondelette en chapeau mexicain. Les axes verticaux et horizontaux représentent respectivement $\log_2 s$ et u . Les échelles les plus fines sont en haut. Les coefficients nuls correspondent à du noir, donc à des parties régulières.	41
Figure II.10 : Modules maximaux de $Wx(u,s)$ de l'exemple de la Figure II.9.	43

Chapitre III

Techniques d'apprentissage et d'intelligence artificielle pour la modélisation de systèmes complexes

Figure III.1 : a) Estimation d'une fonction par un hyperplan, décrivant seulement les dépendances linéaires, et par une courbe sinus de haute fréquence, approximant la fonction en ses échantillons mais sans la décrire réellement. b) La borne réelle du risque est un compromis entre le risque empirique (erreur d'apprentissage) et l'intervalle de confiance (capacité de l'ensemble de fonctions S_k utilisées)	54
Figure III.2 : a) Le neurone biologique et b) le perceptron.	55
Figure III.3 : Les fonctions d'activation les plus utilisées. De gauche à droite : sigmoïdale, gaussienne, tangente hyperbolique et sinusoidale.	56
Figure III.4 : Des fonctions de perte et leurs modèles de densité correspondants. De gauche à droite : gaussienne, Laplacienne, robuste de Huber et ε -insensible.	57
Figure III.5 : Exemple de perceptron multicouche (trois couches cachées).	60
Figure III.6 : La construction d'hyperplans séparateurs dans l'espace des caractéristiques est équivalente à la construction de fonctions de décision non linéaires dans l'espace d'entrée. Donc, des données linéairement inséparables dans l'espace d'entrée peuvent être linéairement séparées dans l'espace des caractéristiques.	68
Figure III.7 : L'hyperplan séparateur optimal (avec la plus grande marge) est unique et peut être défini à l'aide des vecteurs de support (entourés). Dans le cas non linéairement séparable, le $i^{\text{ème}}$ point a une variable de relaxation associée ξ_i représentant l'amplitude de l'erreur de classification.	69
Figure III.8 : L'estimation de la régression à vecteurs de support utilisant la fonction de perte ε -insensible. Les vecteurs de support (avec $0 < \alpha_i < C$), entourés, sont dans $\pm \varepsilon$ à partir de l'hyperplan. Les autres vecteurs sont dans la zone (avec $\alpha_i = 0$) ou hors zone (avec $\alpha_i = C$).	71
Figure III.9 : L'estimation de densité construite à partir de l'hyperplan séparateur optimal, défini à l'aide des vecteurs de support (entourés).	73
Figure III.10 : Schéma général de l'architecture d'un système expert.	82
Figure III.11 : Composants basiques d'un système d'inférence flou.	83
Figure III.12 : Fonctions d'appartenance : singleton, triangulaire, trapézoïdale, gaussienne.	84
Figure III.13 : Modificateurs linguistiques : plus ou moins petit, pas très petite, plutôt grande.	85
Figure III.14 : Partitionnement à l'aide des fonctions trapézoïdales sur $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . Le partitionnement de l'espace d'entrée est composé de neuf règles. Les aires grisées de la partie basse de la figure montrent les régions de recouvrement des ensembles flous.	87
Figure III.15 : Partitionnement par groupage (clustering) à l'aide de trois ensembles flous multidimensionnels A_1 , A_2 , et A_3 , et leur projection respective sur chaque axe pour leur interprétation linguistique. Seulement trois règles existent.	87
Figure III.16 : Représentation du mécanisme d'inférence max-min (Mamdani). A' est l'ensemble flou d'entrée et B' est l'ensemble flou de sortie. La défuzzification de B' par COG est y_g et par MOM est y_m	90
Figure III.17 : Modèle flou relationnel dans $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . La relation R du partitionnement de l'espace d'entrée est composée d'un total de neuf règles.	92
Figure III.18 : Modèle flou Takagi-Sugeno dans $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . La partition de l'espace d'entrée est composée de neuf règles (bas). La sortie y est une approximation linéaire par morceaux d'une fonction non linéaire (haut).	93
Figure III.19 : Normes des distances : Euclidienne (cercle), Mahalanobis (ellipsoïde orienté) et SVM (ensemble non convexe).	98
Figure III.20 : La groupage soustractif (de gauche à droite) : l'ensemble aléatoire de 100 points, les groupes identifiés et leurs centres (entourés), les fonctions d'appartenance gaussiennes.	99
Figure III.21 : Entrée pour l'identification : relation de dilution des eaux usées. Sorties pour l'identification : biomasse, substrat énergétique et substrat xénobiotique.	102
Figure III.22 : Comparaison entre la sortie du processus et les sorties FCM-Flou, GK-Flou et SVM-Flou.	103

Chapitre IV

Le problème de la surveillance du conducteur automobile. Les projets AWAKE et PREDIT

Figure IV.1 : L'automobile et quelques sous-systèmes.....	109
Figure IV.2 : Modèle fonctionnel de conduite automobile.	111
Figure IV.3 : Vue générale du ELS.....	114
Figure IV.4 : Extraction de l'indice VHALL. SH représente la zone de mouvements lents et MH représente la zone de mouvements rapides.....	115
Figure IV.5 : Le temps de sortie de voie, <i>tlc</i> . haut) <i>tlc</i> gauche, milieu) position latérale, bas) <i>tlc</i> droit. ..	116
Figure IV.6 : Système de détection de somnolence de Toyota.	117
Figure IV.7 : Système de détection de somnolence et d'inattention et système de support de maintien du véhicule sur la voie de Nissan.....	118
Figure IV.8 : Le système SafeTRACK.	118
Figure IV.9 : Système de détection de somnolence de DC.....	119
Figure IV.10 : Capteur de suivi du regard de DC.	119
Figure IV.11 : Le système intelligent d'aide au conducteur de Honda.	120
Figure IV.12 : Système de détection PERCLOS.	120
Figure IV.13 : Simulateur de conduite VIRTEX de Ford et le premier système LDW commercialisé massivement en Amérique.	121
Figure IV.14 : Système de suivi du regard de Delphi.	122
Figure IV.15 : Schéma fonctionnel AWAKE.	128
Figure IV.16 : Schéma fonctionnel PREDIT.	128
Figure IV.17 : Véhicule d'expérimentation CopiTech, LAAS-CNRS.	129
Figure IV.18 : Véhicule d'expérimentation Laguna, Siemens-VDO.....	130
Figure IV.19 : Simulateur de conduite PAVCAS, CEPA-CNRS.	130
Figure IV.20 : Simulateurs de conduite automobile également participant dans le projet AWAKE.....	131
Figure IV.21 : Les véhicules démonstrateur AWAKE.	131
Figure IV.22 : Schéma fonctionnel du module de détection de situations à risque.	134
Figure IV.23 : a) Fonctions d'appartenance d'entrée pour la distance latérale et y (μ_{close} et μ_{far}). b) Fonctions d'appartenance d'entrée pour la première dérivée de la distance latérale $\dot{y}$ (μ_{small} , μ_{med} et μ_{large}). c) Fonctions d'appartenance d'entrée pour le temps de correction t_{nc} (μ_{short} et μ_{long}).....	137
Figure IV.24 : a) Fonctions d'appartenance de sortie pour VRBS _{adj} ($\mu_{tighten}$, μ_{none} et μ_{widen}). b) Exemple de réglage du VRBS utilisant y et $\dot{y}$ seulement.....	138
Figure IV.25 : Exemple de la trajectoire du véhicule, montrant l'effet du VRBS ₁ (ligne en tiré long) et du VRBS ₂ (ligne en tiré court) pour le conducteur C3. Le carré représente un avertissement du vibreur placé à 0.15 m. et le triangle un avertissement du <i>tlc</i>	140
Figure IV.26 : Le VRBS ₂ (ligne en tiré court), comparé à l'évolution du VRBS ₁ (ligne en tiré long), génère un avertissement (cercle) pour le conducteur C3. Il n'y a pas d'avertissement du vibreur à 0.15 m. mais une détection du <i>tlc</i> (triangle).	141
Figure IV.27 : Schéma fonctionnel du module de détection de l'Hypovigilance.	142
Figure IV.28 : Schéma fonctionnel du module HDM performance.....	143
Figure IV.29 : a) Spectre en fréquence de la position latérale pour un conducteur typique. b) Echelles d'analyse définies par l'ondelette B-spline	144
Figure IV.30 : a) Spectre en fréquence de l'angle du volant pour un conducteur typique. b) Echelles d'analyse définies par l'ondelette B-spline.	144
Figure IV.31 : Détection de ruptures dans la position latérale (LP) et l'angle du volant (SWA) utilisant l'ondelette B-Spline.	145
Figure IV.32 : Exemple à deux dimensions de la construction de la classe normale. a) Borne apprise par les SVM b) Modèle de densité associée.....	146

Figure IV.33 : Hypothèse de récupération de la vigilance lors de une phase de baisse vigilance observée.	147
Figure IV.34 : Influence du nombre de non détections par rapport à la fréquence.	149
Figure IV.35 : Coefficient de corrélation de l'ensemble formé par les caractéristiques de conduite et les références (bornes extrêmes). La ligne diagonale représente le niveau 1 (corrélation d'une variable par rapport à elle-même).	149
Figure IV.36 : Résultats pour le conducteur C01.	152
Figure IV.37 : Résultats pour le conducteur C02.	153
Figure IV.38 : Résultats pour le conducteur C03.	154
Figure IV.39 : Résultats pour le conducteur C04.	155
Figure IV.40 : Résultats pour le conducteur C05.	156
Figure IV.41 : Résultats pour le conducteur C06.	157
Figure IV.42 : Résultats pour le conducteur C07.	158
Figure IV.43 : Résultats pour le conducteur C08.	159
Figure IV.44 : Résultats pour le conducteur C09.	160
Figure IV.45 : Résultats pour le conducteur C10.	161

Abréviations

Chapitre I

Facteurs impliqués dans la surveillance de l'opérateur humain dans les systèmes homme-machine

ANS	système nerveux autonome	autonomic nervous system
CNS	système nerveux central	central nervous system
ECG	électrocardiogramme	electrocardiogram
EEG	électroencéphalogramme	electroencephalogram
EOG	électrooculogramme	electrooculogram
HMI	interaction homme-machine	human-machine interaction
HP	période du cœur	heart period
HR	fréquence cardiaque	heart rate
IBI	l'intervalle intra-battement	inter-beat interval
PNS	système nerveux parasympathique	parasympathetic nervous system
KSS	échelle d'endormissement Karolinska	Karolinska sleepiness scale
PNS	système nerveux périphérique	peripheral nervous system
SNS	système nerveux sympathique	sympathetic nervous system

Chapitre II

Le prétraitement des signaux, l'analyse temporelle et fréquentielle et extraction de caractéristiques

DFT	transformée de Fourier discrète	discret Fourier transform
DWT	transformée en ondelettes discrète	discret wavelet transform
F	transformée de Fourier	Fourier transform
FFT	transformée de Fourier rapide	fast Fourier transform
FIR	filtres à réponse finie	finite impulse response
IIR	filtres à réponse infinie	infinite impulse response
STFT	transformée de Fourier à court terme	short-time Fourier transform
W	transformée en ondelettes	wavelet transform

Chapitre III

Techniques d'apprentissage et d'intelligence artificielle pour la modélisation de systèmes complexes

ANFIS	système d'inférence flou basé sur des réseaux adaptatifs	adaptive-network-based fuzzy inference system
ANN	réseaux de neurones artificiels	artificial neural networks
COG	centre de gravité	center of gravity
DF	fonction de décision ou fonction discriminante	decision function or discriminant function
EM	maximisation de l'espérance	expectation maximization
EQP	problème d'optimisation quadratique avec des contraintes égalité	equally constrained quadratic programs

FCM	C-moyennes floues	fuzzy c-means
FCRM	modèles de c-régression floue	fuzzy c-regression model
FIS	systèmes d'inférence flou	fuzzy inference systems
GG	Gath–Geva	Gath–Geva
GK	Gustafson–Kessel	Gustafson-Kessel
GMM	modèles de combinaison gaussienne	gaussian mixture models
iid	indépendant et identiquement distribué	independent and identically distributed
KDE	estimation de la densité par kernels	kernel density estimation
kNN	k plus proche voisin	k nearest neighbor
LMS	moindres carrées	least mean squares
LVQ	apprentissage à quantification vectorielle	learning vector quantization
ML	maximum de vraisemblance	maximum likelihood
MLP	perceptron multicouches	multi layer perceptron
MOM	moyenne des maximums	mean of maxima
PCM	c-moyennes floues possibilistes	possibilistic c-means
PDF	fonction de densité de probabilité	probability density function
RBF	réseaux de fonctions à base radiale	radial basis functions
RDA	analyse discriminante régularisée	regularized discriminant analysis
SLP	perceptron monocouche	single layer perceptron
SSDP	symétrique semi-définie positive	symmetric semi-definite positive
SVC	classification à vecteurs de support	support vector classification
SVDE	estimation de densité à vecteurs de support	support vector density estimation
SVM	machines à vecteurs de support	support vector machines
SVR	régression à vecteurs de support	support vector regresion

Chapitre IV

Le problème de la surveillance du conducteur automobile. Les projets AWAKE et PREDIT

ADAS	système avancé pour l'aide à la conduite	advanced driving assistance system
ACC	régulateur de vitesse adaptatif	adaptive cruise control
CAS	systèmes d'évitement de collision	collision avoidance systems
ELS	capteur de mouvement des paupières	eye lid sensor
GPS	système de positionnement global	US global positioning system
IHCC	régulateur de vitesse intelligent pour autoroute	intelligent highway cruise control
ITS	systèmes de transport intelligents	intelligent transportation systems
LDW	avertisseur de sortie de voie	lane departure warning
LKAS	système de maintien du véhicule sur la voie	lane-keeping assistance system
PERCLOS	pourcentage de fermeture des yeux	percentage of closure of eyes
RBS	vibreurs des bords latéraux de la route	rumble strips
RDW	avertissement de sortie de route	road departure warning
SVRD	sorties de route mono-véhicule	single-vehicle road departure
TLC	temps de croisement	time to lane crossing
VRBS	vibreux virtuel du bord de route	virtual rumble strip

Introduction

Analyser puis maîtriser la complexité sont des objectifs permanents de l'action de Recherche dans toutes les disciplines. L'ambition est d'abord de comprendre, de trouver des modèles représentatifs pour pouvoir anticiper la survenue de phénomènes particuliers.

La complexité peut être appréhendée dans des approches très diverses en fonction des disciplines et des objectifs que l'on vise : le nombre de variables mises en jeu et les caractères non linéaires des lois qui les associent donnent une certaine mesure de cette complexité :

- Dans les systèmes réels, qu'il s'agisse de l'homme ou de son environnement, les modélisations sont toujours partielles
- Dans les systèmes artificiels, conçus par le génie humain, les systèmes sont en principe maîtrisés dans la limite où ils peuvent être soumis et résister à des influences externes incontrôlées et éventuellement être sous le « contrôle » d'opérateurs humains qui vont apporter, dans le fonctionnement, toute leur complexité propre.

Dans les systèmes artificiels, il faut toutefois ne parler que d'une maîtrise de principe car il est bien connu aujourd'hui que dans les développements algorithmiques, on ne peut pas valider toutes les configurations d'utilisation de ces algorithmes.

Dans la surveillance des systèmes complexes, il paraît intéressant d'adopter une approche de la complexité qui distingue le fonctionnement normal d'un système et l'apparition de fonctionnements anormaux, définis comme des écarts à l'enveloppe des états normaux. Avec cette approche, où on ne cherche pas à analyser finement le fonctionnement des systèmes complexes, on limite l'ambition à le caractériser, selon une exigence précise, par des valeurs de ses variables pertinentes correspondant à son fonctionnement normal et à détecter les écarts à cette normalité. Si l'on considère des systèmes strictement technologiques, on va parler de diagnostic de fautes. Si l'on considère un système homme-machine, il faudra faire intervenir la notion d'erreur humaine ou quelquefois de défaillance humaine (cas de l'hypovigilance du conducteur automobile).

Notre travail se situe dans le cadre de ***la surveillance de systèmes où l'homme est dans la boucle, où il joue un rôle actif***. C'est probablement le cas des systèmes complexes les plus difficiles à appréhender puisque les systèmes de mesure que l'on va installer vont donner des informations qui, certes, vont caractériser le fonctionnement technologique, mais aussi les choix réalisés par l'opérateur humain. La difficulté évidente est le niveau de bruit généré sur les variables de mesure par ces choix successifs :

- Si les choix sont justifiés, que le système fonctionne bien, il faut bien admettre que l'homme est un bon opérateur reproductible et qu'il peut être, en quelque sorte, modélisé partiellement par ses ***bonnes habitudes***.
- Si ce sont des mauvais choix, espérons en petit nombre, ils pourront être repérés par la détection du comportement anormal du système tout entier.

Ce raisonnement se fonde sur l'étude des modes opératoires naturels de l'homme : leur existence doit permettre de modéliser la normalité de fonctionnement. Il suffira, alors, de détecter les écarts à cette normalité pour ensuite repérer et, dans les meilleurs cas, anticiper des défaillances du système global.

C'est cette approche que nous nous proposons d'explorer dans notre thèse. Dans une formulation générale, on pourrait dire qu'il s'agit d'évaluer la probabilité de défaillance par la création d'un modèle de bon fonctionnement et la détection d'écarts significatifs à ce bon fonctionnement. Bien sûr, nous nous appuierons sur le traitement d'un exemple précis qui est la détection de l'hypovigilance du conducteur. Ce n'est qu'après avoir conclu sur cet exemple que nous reviendrons sur les perspectives générales, ouvertes par notre approche de la surveillance et du diagnostic des systèmes complexes.

Notre travail fait suite à plusieurs thèses [77], [79], [134], sur la détection d'hypovigilance. Elles ont successivement montré :

- La présence de modes opératoires personnalisés relatifs à la conduite automobile, détectables par des techniques d'attente, [79].
- La fiabilité d'un diagnostic temps réel basé sur la sélection de variables permettant de réaliser une classification normale-anormale des modes de conduite, [77].
- Les limites de l'approche dans les étapes de validation réalisée par la comparaison entre le diagnostic du système de surveillance (mesures de performance) celui du physiologiste basé sur l'enregistrement EEG et l'avis du conducteur (mesures physiologiques et mesures subjectives), [134].

Notre objectif applicatif est d'améliorer la performance du diagnostic temps réel :

- Par la fusion de modèles, résultant d'une analyse plus fine d'évolutions lentes (vigilance) dans les modes de conduite et de l'analyse d'événements instantanés (performance et niveau de risque).
- Par la prise en compte des conditions environnementales pouvant conduire à multiplier les modèles de plus faible complexité, mais de meilleure précision dans un domaine de validité (décomposition et fusion de modèles).

Notre travail est organisé en quatre chapitres :

Dans le chapitre un, nous discutons l'impact des systèmes homme-machine, notamment leurs différents niveaux de coopération et de partage de tâches, pour aborder leurs différents types de défaillances en nous concentrant sur la partie humaine ; les causes de l'erreur humaine et les différents domaines de la performance humaine impliqués dans l'interaction homme-machine, tout en mentionnant que les différentes métriques que nous pouvons exploiter : mesures physiologiques, mesures subjectives et mesures de performance. C'est ici que l'on exploite toute la connaissance *a priori* du système HM. Les objectifs des systèmes de surveillance et commande supervisée seront présentés pour la construction de systèmes d'assistance intelligents.

Le chapitre deux est consacré à la description des différentes méthodes d'analyse de signaux, pour leur prétraitement et l'extraction d'informations pertinentes. Nous nous sommes limité à l'étude de signaux monodimensionnels pour l'analyse de leur dynamique en temps et en fréquence et de leurs caractéristiques statistiques et fréquentielles ; l'implémentation des techniques sous-jacentes pour une analyse temps réel est provisoirement évoquée. C'est dans cette étape que l'on produit un ensemble « optimal » d'entrées avec le maximum d'information en un minimum de variables.

Le chapitre trois décrit les différentes approches de modélisation que nous avons abordées, commençant par la modélisation utilisant des techniques d'apprentissage, des systèmes experts ou les deux. Du côté des méthodes d'apprentissage, nous détaillerons tout particulièrement le développement des machines à vecteurs de support pour lesquelles nous avons fait une étude du processus d'apprentissage. Nous avons proposé des résolutions adaptées du problème d'optimisation quadratique, QP, concerné, ainsi qu'une stratégie optimale d'implémentation de la méthode de décomposition du QP pour la résolution de problèmes grande échelle. Concernant les systèmes à base de connaissance, nous avons exposé les différents types de systèmes d'inférence floue, abordant chacun de leurs composants, et

leur construction. Nous traitons également les schémas de construction de systèmes profitant des observations et de l'expertise, afin de développer un système de modélisation hybride SVM-floue. Ces techniques vont nous permettre de construire un modèle pour pouvoir fournir un diagnostic convenable à partir des données d'entrée.

Le chapitre quatre situe le contexte général de l'application des méthodes introduites dans les chapitres deux et trois pour le développement des systèmes d'assistance au conducteur automobile. Nous présenterons les objectifs du projet de recherche Européen AWAKE (System for effective Assessment of driver vigilance and Warning According to traffic risK Estimation) et ceux du projet de recherche Français « Facteurs Biologiques et Environnementaux de la Dégradation de la Vigilance chez les Conducteurs et Système de Détection Automatique des Baisses de Vigilance »¹, action du PREDIT, ainsi que les objectifs généraux de travail. De cette manière, nous proposons une stratégie générale de système temps-réel pour la surveillance du niveau de vigilance du conducteur (à dynamique lente) et la surveillance du niveau de risque lié à la situation actuelle de conduite (dynamique instantanée). Finalement, nous présenterons les différentes expériences et analyserons et discuterons les résultats.

Nous terminons, enfin, par des conclusions et une prospective ouverte sur le problème de maîtrise des systèmes Homme-Machine qui reste très délicat. Notre travail s'est déroulé au LAAS dans le cadre d'une collaboration entre les groupes MIS et DISCO ce qui nous a fait bénéficier d'un partenariat privilégié.

¹ Nous utilisons l'acronyme PREDIT pour faire référence à ce projet en particulier

Chapitre I

Facteurs impliqués dans la surveillance de l'opérateur humain dans les systèmes Homme–Machine.

Sommaire

I.1	Introduction.....	6
I.2	Les systèmes complexes	6
I.2.1	L'impact de l'automatisation dans les systèmes complexes.....	7
I.2.1.1	Avantages de l'automatisation.....	7
I.2.1.2	Inconvénients de l'automatisation	7
I.2.2	Les interactions homme–système.....	8
I.2.2.1	Coopération homme–système.....	8
I.2.2.2	Partage de tâches.....	10
I.3	Défaillances dans les systèmes homme–machine.....	11
I.3.1	La performance humaine dans l'interaction homme–machine.....	11
I.3.2	Charge de travail.....	12
I.3.3	Conscience de la situation	12
I.3.4	Surveillance de système	13
I.3.5	Travail d'équipe	14
I.3.6	Confiance.....	14
I.3.7	Utilisabilité et acceptation	15
I.3.8	Erreur humaine	15
I.4	Métriques sur la performance humaine	16
I.4.1	Mesures physiologiques	16
I.4.1.1	L'activité cérébrale	17
I.4.1.2	L'activité cardiaque	18
I.4.1.3	L'activité oculaire	18
I.4.2	Mesures subjectives.....	19
I.4.2.1	Autoévaluation instantanée.....	20
I.4.2.2	Technique subjective d'évaluation de charge de travail	20
I.4.2.3	Indice de charge de la tâche.....	20
I.4.2.4	Echelle d'endormissement Karolinska	20
I.4.3	Mesures de performance.....	21
I.4.3.1	Métriques de performance de la charge de travail	21
I.4.3.2	Mesures de la charge de tâche	22
I.5.	Surveillance des systèmes homme–machine	23
I.5.1	Les systèmes d'assistance intelligents.....	24
I.6.	Conclusion	24

I.1 Introduction

Nous nous intéressons, dans ce premier chapitre, à présenter de manière générale la notion de système complexe, considéré du point de vue de l'automaticien concerné par l'analyse et le contrôle des systèmes.

Nous mettons l'accent sur les systèmes homme-machine qui nécessitent une surveillance automatisée globale, incluant la détection d'incidents techniques et de défaillances humaines. Nous tentons de faire une analyse critique de l'existant : Quelles approches ont déjà été conduites dans la surveillance des systèmes homme-machine ? De quelle manière peut-on envisager une modélisation du comportement humain ? Comment cette modélisation peut-elle s'intégrer dans un outil de diagnostic ?

Partant de cette réflexion générale, nous présentons un cas particulier, celui de la conduite automobile, qui nous servira à la fois de sujet de recherche et de support de généralisation des résultats obtenus. C'est un sujet qui a fait l'objet d'efforts importants, compte tenu des risques encourus par les utilisateurs du transport routier. Nous donnerons notre point de vue sur l'état des connaissances et des insuffisances que nous avons identifiées.

Sur cette base, nous formulons notre problématique :

- Recherche d'une méthodologie plus approfondie et plus efficace dans l'application à la détection de l'hypovigilance du conducteur automobile.
- Tentative de généralisation de l'utilisation de nos méthodes et outils à la supervision de systèmes homme-machine plus généraux.

I.2 Les systèmes complexes

Les sciences de la nature ont connu un véritable succès en réussissant à expliquer les phénomènes du monde entourant de l'homme. Une approche très efficace a été celle de décomposer en plusieurs classes de phénomènes et d'étudier chacun d'entre eux isolément. De cette manière, on peut expliquer la physique atomique à partir des particules élémentaires, la chimie à partir des molécules et les organismes à partir des cellules, etc. Cette approche, appelé **réductionnisme**, a connue un succès considérable [5].

Les sciences de l'ingénieur ont comme but de développer la connaissance nécessaire pour concevoir des systèmes de production d'énergie, de production d'instruments et de biens d'équipement, des systèmes de transports, de communications, etc. Quand les sciences de l'ingénieur ont émergé, l'utilisation de l'approche réductionniste a été naturelle. Ceci a conduit à une subdivision en génie civil, génie mécanique, génie électrique et génie chimique, qui a fonctionné au XIX^{ème} siècle et au début du XX^{ème} siècle. Cependant, au fur à mesure que la complexité des systèmes augmentait, il a paru évident que les problèmes de conception ne pouvaient pas s'aborder à travers une seule discipline spécifique, qu'il fallait adopter un point de vue holistique, cela dit, considérer l'interaction entre les différentes parties d'un système au lieu des unités isolées. Ceci a amené à un grand développement de plusieurs disciplines nouvelles dont l'**automatique**, démontrant ainsi qu'il existe des principes d'interaction dans les systèmes, comme celui de la **rétroaction**, essentiels pour aborder l'analyse et la commande de systèmes complexes technologiques.

Ainsi, dans les années 1970–1980, les systèmes complexes sont caractérisés par leur grande dimension et l'on s'efforce à les structurer en sous-systèmes faiblement interconnectés. La résolution des problèmes afférant à leur maîtrise conduit alors à développer des méthodes basées sur le principe de décomposition (« diviser pour mieux régler ») :

$$\{\text{Sol } Pb \text{ global}\} \Leftrightarrow \{\text{Sol } subPb_1(\lambda), \dots, \text{Sol } subPb_N(\lambda)\}_{\lambda=\lambda^*}, \quad (\text{I.1})$$

pouvant aller jusqu'à une décomposition ultime (système atomique), où il n'est pas envisagé d'aller au delà, soit par nature physique, soit parce que c'est dénué d'intérêt.

La décomposition peut prendre deux aspects [5]:

- Décomposition *horizontale* visant à hiérarchiser les systèmes complexes en plusieurs sous-systèmes interconnectés.
- Décomposition *verticale* basée sur la complexité de la tâche de commande (systèmes à réaliser : régulateur, optimiseur, adaptateur).

Actuellement, la notion de complexité intègre d'autres aspects, notamment la prise en compte de l'homme en tant qu'acteur de décision ou élément dans la boucle, qui amènent à traiter des informations de nature plus variées : variables quantitatives, variables qualitatives, variables linguistiques, entourées souvent d'incertitude et imprécision, d'où le recours à des nouvelles méthodes et techniques (intelligence artificielle, systèmes experts, etc.).

I.2.1 L'impact de l'automatisation dans les systèmes complexes

L'automatisation peut être considérée comme un processus de substitution d'une certaine activité humaine par un dispositif ou machine, mais également comme un état de développement technologique. Cependant, d'autres, [166], estiment que l'automatisation devrait être vue comme la substitution d'un **agent** par un autre. Néanmoins, la présence de l'automatisation a infiltré chaque aspect de la vie moderne. Les machines non seulement facilitent, rendent plus sûr et efficace le travail, mais également nous donnent plus de temps libre. L'arrivée de l'automatisation nous a permis d'atteindre ces buts. Avec l'automatisation, les machines peuvent maintenant exécuter plusieurs activités que nous devons faire autrefois. Maintenant, les portes automatiques s'ouvrent pour nous, les thermostats règlent la température dans nos maisons pour nous, les transmissions automobile « passent les vitesses » pour nous.

I.2.1.1 Avantages de l'automatisation

Il existe un certain nombre d'avantages dans l'automatisation des systèmes homme-machine. Parmi eux existent l'augmentation de la capacité de production et de la productivité, la réduction des petites erreurs, la réduction de la charge de travail manuel et de la fatigue, l'assistance dans des opérations courantes, la manipulation plus précise des opérations courantes et l'utilisation économique des machines.

Même si, dans certaines applications, on peut (on doit) faire le choix de systèmes complètement automatisés, autonomes (i.e. la conquête spatiale lointaine, l'exploration de fonds marins, le démantèlement de centrales nucléaires, etc.), l'opérateur humain est souvent introduit dans les systèmes en tant qu'élément de décision dans la boucle (systèmes de pilotage variés : avion, train, métro, navire, automobile,...) ou en tant qu'élément de la hiérarchie de gestion, supervision des systèmes de production (contrôle de trafic aérien, conduite de procédés industriels, surveillance des centrales nucléaires, des raffineries, etc.).

Cette présence humaine intervient, souvent, pour des raisons de complexité moindre, d'acceptabilité sociale, mais surtout parce que l'homme n'est pas encore capable, lors de la conception d'un système complexe, d'imaginer des schémas d'action préprogrammés pour des situations inconnues. L'opérateur humain a un ensemble de références qui lui permettent rapidement d'identifier la situation présente. En particulier, il est capable de faire des analogies, même si celles-ci ne sont pas identiques aux « modèles » qu'il connaît déjà. Ces modèles de situations sont le résultat d'années d'expérience. Ce type de connaissance est très difficile, est bien souvent impossible, à faire expliciter [32].

I.2.1.2 Inconvénients de l'automatisation

L'automatisation a également un certain nombre d'inconvénients. L'automatisation augmente et complexifie les charges des responsables du fonctionnement, du dépannage et de la gestion des systèmes.

Dans, [166], Woods déclaré que l'automatisation est « ... un ensemble intégré – un ensemble composé de différentes dimensions intégrées conjointement comme un système composé de matériel-logiciel. Quand un nouveau système automatisé est introduit dans une activité, le changement est réparti au long de ces dimensions ». Certains de ces changements incluent :

- l'addition ou le changement d'une tâche, telle que l'installation et l'initialisation d'un dispositif, le contrôle de la configuration et les séquences d'opération ;
- le changement des demandes cognitives, telles que la baisse de la conscience de la situation ;
- le changement du rôle actuel des opérateurs dans le système, souvent reléguant des personnes aux tâches de surveillance ;
- l'incrément de l'intégration des différentes parties d'un système, souvent ayant pour résultat la surcharge et la « transparence » des données ;
- et l'incrément de la satisfaction de ceux qui emploient la technologie.

Ces changements peuvent avoir comme conséquence une satisfaction professionnelle pauvre (automatisation vue comme déshumanisante), la vigilance affaiblie, des systèmes intolérants aux fautes, des défaillances silencieuses, une augmentation de la charge de travail cognitif, des défaillances induites par automatisation, l'excès de dépendance, l'ennui accru, la confiance diminuée, la diminution de la compétence manuelle, des fausses alarmes et une diminution de la conscience de mode opératoire, [123].

I.2.2 Les interactions homme-système

Un *système homme-machine*, HM, est composé d'un opérateur en interaction avec un système technologique. Ce concept est généralement associé à un poste de travail et s'applique quelles que soient l'étendue et la complexité du système homme-machine (un pilote d'avion ou de voiture, par exemple).

En plus des deux éléments d'un système homme-machine, le *système technologique* et l'*homme*, il est important de s'intéresser aux interactions homme-système régies par l'*organisation du travail*. D'une part, le système technologique peut être un processus de fabrication, une voiture, un ordinateur, etc. et regroupe, en ergonomie, les systèmes matériels avec lesquels l'homme est en interaction à travers des interfaces homme-système. Ces dernières peuvent être définies comme l'ensemble de moyens mis à disposition de l'opérateur. D'autre part, l'opérateur, concept central en ergonomie, renvoie à tout homme réalisant une tâche en interaction avec un système technologique dans une situation de travail. Le terme homme est lui-même utilisé de façon plus large pour faire référence à l'opérateur isolé ou au collectif de travail. Ce collectif est constitué d'une collection, groupe ou équipe d'opérateurs coopérants, chargés de manière permanente ou temporaire de l'accomplissement d'une fonction. Finalement, l'organisation du travail couvre l'ensemble des règles régissant la relation entre les hommes et les systèmes technologiques. Elle inclut les spécifications du rôle des hommes et des systèmes technologiques (fonctions, tâches et procédures), la répartition de tâches entre les différents composants du système homme-machine, les aspects réglementaires et contractuels (horaires, organisation en équipes, relations hiérarchiques, règlement de sécurité, ...).

I.2.2.1 Coopération homme-système

La notion de coopération est introduite au sein des interactions homme-système, en distinguant différents niveaux : co-activité, collaboration, coopération.

La co-activité homme-système

Il s'agit d'une relation simple dans laquelle l'homme et le système technologique n'ont pas à interagir, mais à réaliser des tâches qui interfèrent, positivement ou négativement. L'interférence résulte d'une relation conjoncturelle de nature *temporelle* (les composants interviennent successivement en un même endroit), *spatiale* (les composants interviennent simultanément en un même endroit), ou *causale*

(les composants interviennent en des endroits différents mais sur des supports reliés par une relation causale).

Dans la co-activité, l'interférence est éventuellement connue par l'un ou l'autre des composants en relation ; malgré cela, elle n'est pas maîtrisée par eux. C'est le cas d'un calculateur et d'un homme utilisant simultanément les capacités de traitement d'un calculateur central, où le lancement d'un traitement lourd par l'homme alors que le système monopolise les capacités du calculateur et interfère, négativement dans ce cas, avec l'activité de l'autre utilisateur. Dans le domaine homme-homme, cette co-activité se rencontre lors des interventions de la maintenance sur un procédé en exploitation ; même s'il existe des procédures pour maîtriser les interférences, il n'est pas rare que des défaillances, parfois catastrophiques, succèdent aux interventions des équipes de maintenance (On peut bien entretenir une voiture et ne pas être à l'abri d'une défaillance mécanique à n'importe quel moment).

Inversement, la co-activité peut permettre de détecter des erreurs, les interférences entre les activités permettant par exemple de détecter une mauvaise configuration du système. La co-activité peut être connue ou non, maîtrisée ou non, coordonnée ou non, au niveau du système incluant les composants humains et technologiques. Ces co-acteurs n'ont pas de relations fonctionnelles explicites, ni de relations hiérarchiques.

La collaboration homme-système

Celle-ci est une forme d'interaction dans laquelle l'homme et le système s'assistent ou s'utilisent pour atteindre leurs buts spécifiques. Parmi l'ensemble de leurs fonctions, les composants ont la tâche (explicite ou non) de satisfaire la demande d'interaction de l'autre. Ils disposent de moyens nécessaires pour mettre en œuvre cette interaction (en l'occurrence, les interfaces homme-système). L'assistance peut être toujours unidirectionnelle (l'interaction entre l'homme et le système se résume à la demande, par le système des valeurs de variables non instrumentées, par exemple) ou bidirectionnelle (l'homme peut demander un relevé de valeurs et le système peut solliciter l'homme pour « renseigner » certains paramètres).

Il existe de nombreuses formes de collaboration, de la plus simple (demande ponctuelle d'information, sans explication de la finalité d'utilisation de celle-ci), à la plus riche (offre volontaire de collaboration pour assister l'autre composant dans la réalisation de sa tâche ; prise en charge par le système de tâches habituellement réalisées par l'homme pour alléger sa charge de travail, détection, diagnostic et résolution d'un incident du système par l'homme, ...).

La plupart des interactions que l'on peut observer entre un opérateur et un système de conduite d'un procédé relève de ce niveau d'interactions où l'on peut considérer les composants humains et technologiques comme des collaborateurs. Les collaborateurs sont inscrits dans une relation hiérarchique et/ou fonctionnelle.

La coopération homme-système

Dans cette interaction, l'homme et le système sont en charge, collectivement, de l'atteinte d'un objectif commun, souvent inaccessible à l'un ou à l'autre. Pour qu'un dialogue instauré entre un homme et un système inscrive l'interaction dans un contexte de coopération, il faut qu'un objectif commun soit clairement identifiable (à défaut d'être identifié), partagé de façon explicite entre les deux composants.

De même que pour la collaboration, il existe de nombreuses formes de coopération ; le niveau plus élevé se rencontre dans les coopérations homme-homme et repose sur une reconnaissance mutuelle d'intention. Si la collaboration peut s'inscrire dans la relation homme-homme ou homme-système hiérarchique, la coopération s'inscrit nécessairement dans une relation exclusivement fonctionnelle, en dehors des relations hiérarchiques pouvant exister entre les coopérants. Une divergence sur les actions, les moyens pour l'atteinte de l'objectif ou sur les objectifs intermédiaires induit une *rupture de la coopération* accompagnée d'une redistribution des rôles. Cette rupture replace au mieux l'interaction au

niveau inférieur de collaboration. Celui qui prend la main, prend du même coup la charge de l'atteinte de l'objectif, l'autre lui apporte *assistance* (une voiture équipée d'un système ACC peut assister le conducteur en l'informant de la distance frontale quand ce module n'est pas activé).

Ainsi, la coopération homme-système passe par la mise en jeu de mécanismes coopératifs utilisant les fonctionnalités de l'interface destinées à faciliter les activités collectives entre les opérateurs d'une équipe ou entre l'opérateur et un système technologique. Les fonctionnalités doivent, bien entendu, respecter des critères de convivialité et de pertinence pour être totalement adaptées aux besoins des opérateurs.

I.2.2.2 Partage de tâches

Le niveau de coopération entre les hommes et les systèmes dépend de la logique de répartition des tâches. Ainsi, on peut la définir comme l'allocation de rôles entre les différents composants du système homme-machine (humains, organisationnels, technologiques) pour permettre à ce dernier de fournir le service attendu à ses systèmes utilisateurs.

Pour exploiter au mieux les qualités respectives des hommes et des systèmes, on peut mesurer la performance d'un simple système homme-machine comme une fonction des différents niveaux d'autonomie et des différents niveaux de compréhension ou de connaissance autour de l'environnement du système [32], Figure I.1. L'axe horizontal représente les niveaux d'autonomie allant du contrôle manuel jusqu'à l'automatisation complète. L'axe vertical représente la performance du système homme-machine. Celle-ci peut être mesurée par le temps d'exécution, la précision des résultats, le coût de la solution ou un autre critère heuristique approprié, pour un problème donné. Chaque courbe de la Figure I.1 correspond à un certain niveau de connaissance ou compréhension, la plus basse étant celle qui représente un niveau pauvre de connaissance.

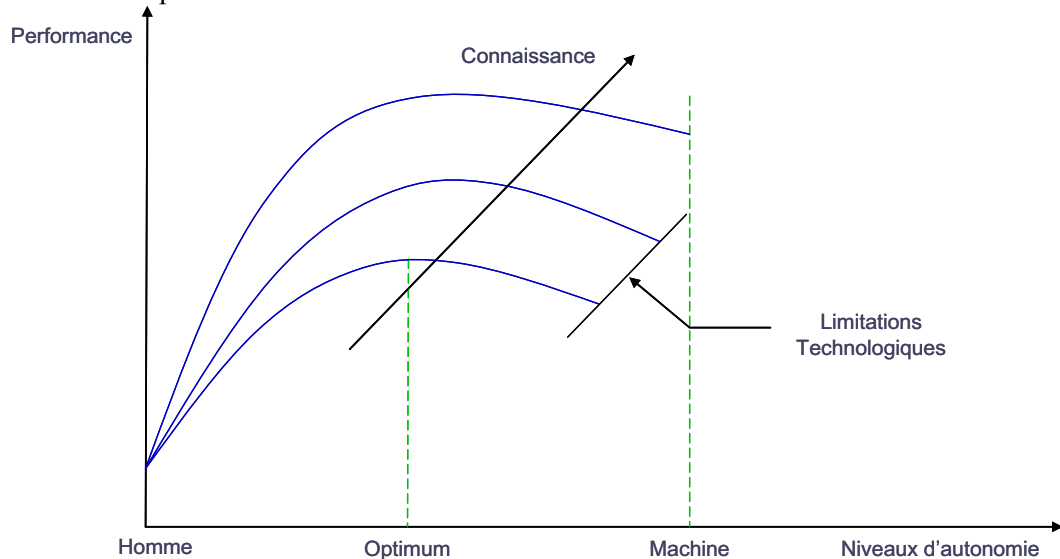


Figure I.1 : Performance d'un système homme-machine comme une fonction des niveaux d'autonomie.

En accord avec ce modèle, la performance du système homme-machine augmente avec l'autonomie de la machine, mais seulement jusqu'à un certain optimum après lequel elle commence à décliner. Si l'autonomie de cette machine est trop importante pour que l'opérateur humain comprenne ce qu'elle fait, il peut perdre le contrôle de la situation. Cette perte peut entraîner également une diminution de la vigilance due à une charge de travail insuffisante ou excessive. Il faut bien préciser que, pour un niveau de connaissance donné, le niveau d'autonomie ne peut pas excéder une certaine valeur fixée par les limitations technologiques.

I.3 Défaillances dans les systèmes homme-machine

Une défaillance du système homme-machine survient lorsque la fonction délivrée s'*écarte* de la satisfaction de la fonction de ce système. Celle-ci peut directement résulter de la défaillance de l'un des composants (humain ou technologique) ou encore provenir des erreurs de spécification du système homme-machine.

Suivant le point de vue considéré, la perception des défaillances peut être *cohérente* (tous les éléments du système ont le même jugement de défaillance) ou *incohérente*, qualifiée également de *byzantine* (les utilisateurs peuvent avoir des jugements différents d'elle), [96]. Les limites de tolérance fixées pour apprécier la défaillance ne sont pas absolues mais relatives à chacune des références pouvant estimer le caractère valide ou défaillant du comportement du système.

La défaillance d'un système homme-machine résulte d'une **erreur** affectant son état. La chaîne causale entre les causes de l'erreur, l'erreur elle-même et ses conséquences s'applique donc aussi bien au système homme-machine qu'à ses composants humains et technologiques, mais le point d'entrée diffère selon que l'on cherche à agir sur un composant technique ou humain. Dans le cadre des **composantes technologiques**, un concept important est celui de **faute** (cause adjugée ou supposée de l'erreur).

Pour ce qui concerne le **composant humain**, l'obtention de la sûreté de fonctionnement passe par des ressources *internes* à l'opérateur (connaissance, représentations mentales, diversité des modes de traitement, ...) ainsi que par des ressources *externes* (détection de ses erreurs par l'interface homme-machine ou **assistance externe** à ses mécanismes de régulation internes). Evidemment, on ne peut agir sur les causes internes de l'erreur humaine et sur les moyens internes de prévention et tolérance que de l'extérieur par des supports tels que la formation, les modifications d'interfaces et la documentation, etc. Ces mesures sont indirectes et présentent un long délai de réponse.

Actuellement, la recherche des causes de l'erreur (internes ou externes à l'homme) n'est que rarement possible et ne donne pas forcément les éléments nécessaires pour établir une stratégie de recouvrement correspondante. On ne dispose que d'une quantité limitée (ou nulle) d'informations sur les gestes de l'opérateur, sa motivation d'action, etc. qui sont des éléments nécessaires pour déterminer de façon certaine le caractère adapté ou erroné de son comportement. La vérification du caractère erroné permet la déduction des techniques et moyens à mettre en œuvre pour optimiser sa contribution dans le système homme-machine.

I.3.1 La performance humaine dans l'interaction homme-machine

L'étude de la performance humaine fait référence à sept grands domaines, chacun identifié comme critique aux interactions des systèmes homme-machine, [52] :

- charge de travail
- conscience de la situation
- surveillance de système
- travail d'équipe
- confiance
- utilisation et acceptation par les utilisateurs
- erreur humaine

En plus de ces sept domaines généraux de performance humaine, un domaine humain « spécifique » additionnel de performance est la **performance de la tâche** (c.-à-d. de quelle façon l'opérateur exécute sa tâche, par rapport aux mesures concernant la performance globale du système). Puisqu'une telle mesure dépend fortement de son application, elle est omise de la liste ci-dessus et est donc adressée au cas par cas. Pour chacun des sept secteurs, il existe un certain nombre de théories traitant des mécanismes sous-jacents et de la signification opérationnelle des résultats.

I.3.2 Charge de travail

L'intérêt de définir et développer des métriques sur la charge de travail s'est considérablement développé à partir milieu des années 1970. Il existe une discussion continue dans la communauté scientifique au sujet de la définition de la charge de travail (une discussion qui est également reflétée dans la variété de moyens disponibles pour évaluer la charge de travail). Néanmoins, il y a un accord général qui établit que la charge de travail mentale n'est pas unitaire, mais un concept multidimensionnel englobant la difficulté d'une tâche et l'effort appliqué associé (physique et mental).

Inhérent à la notion de charge de travail mental a été l'idée que l'opérateur humain a une capacité limitée pour traiter l'information. Les modèles de traitement de l'information des années 1950 se sont développés hors du champ de la technologie des communications. L'un des points de vue empiriques les plus supportés était le regroupement des multiples ressources spécifiques. Les tâches diffèrent entre elles par rapport à leurs demandes en termes de : modalité de l'entrée (visuelle contre auditif), de perception des entrées (spatiale/verbale), d'interprétation de l'information (codage/interprétation/réponse) et du type de réponse (manuelle contre verbale).

Les techniques pour mesurer la charge de travail sont typiquement classées par catégories :

- *physiologiques*,
- *subjectives*, (incluant les psychologiques)
- ou de *performance*.

Dans chaque catégorie, il y a un certain nombre d'indicateurs spécifiques disponibles pour le chercheur. Les critères appropriés pour juger l'utilité de divers indicateurs de charge de travail sont en général : sensibilité, diagnosticabilité, coût, acceptabilité de l'opérateur, requêtes d'implémentation, fiabilité, et intrusion, [41]. Cependant, rare est l'indicateur de charge de travail mental pouvant satisfaire tous ces critères à la fois.

I.3.3 Conscience de la situation

La conscience de la situation (situation awareness, SA) est un concept actuellement populaire, employé pour décrire la compréhension d'un opérateur des caractéristiques dynamiques et complexes d'un système, [53]. Actuellement, le rôle primordial de la SA dans le domaine des facteurs humains est dû, en grande partie, à l'incrément de la nature cognitive des tâches homme-machine, [50]. La littérature sur la performance humaine distingue généralement la SA et la construction de la surveillance de système, comme discuté dans la section suivante. Bien qu'il y ait un certain chevauchement potentiel entre les deux (mais le même peut être dit de tous les domaines identifiés de la performance humaine), le dernier concerne généralement la détection et la réponse initiale à un certain état discret non nominal (une alarme soudaine de dépassement des limites d'opération). La notion de la SA se prolonge bien loin au delà de ceci, pour entourer la compréhension de l'opérateur des états du système complexe et la capacité de prévoir le futur comportement du système. Néanmoins, certaines des techniques comportementales d'évaluation de la SA pourraient employer des mesures de surveillance typiques (par exemple, *temps de réponse, taux de détections, fausses alarmes*).

La carence courante d'un consensus sur la définition de la SA est propagée dans la gamme des techniques de mesure employées pour l'évaluer. Les techniques se sont étendues, commençant par des mesures d'autoévaluation, des évaluations « au dessus de l'épaule », jusqu'à l'utilisation des mesures physiologiques telles que la *direction du regard*, [52]. Dans, [50], les auteurs ont passé en revue les techniques disponibles de mesure sur la SA et ont distingué trois types populaires de mesure de la SA :

- *méthodes subjectives*, telles que l'autoévaluation,
- *méthodes à base de questionnaires*, soit avec la situation originale absente, soit avec la situation continûment présente,
- *méthodes de performance*, semblable aux mesures de charge de travail, les tâches opérationnelles relevantes pouvant fournir des mesures implicites sur la SA.

Comme avec les mesures « charge de travail mentale » il y a plusieurs années, des discussions sur comment évaluer la SA reviennent souvent sur un ensemble fondamental de critères évaluatifs (sensibilité, diagnosticabilité, coût, acceptation de l'opérateur, requêtes d'implémentation, fiabilité et intrusion), [41], [52]. Souvent ces critères doivent être mesurés l'un contre l'autre. Un exemple particulier sont les mesures subjectives de la SA, qui sont contraintes au même ensemble d'avantages potentiels (prix réduit, facilité de gestion) et d'inconvénients (limitations de mémoire, susceptibilité aux caractéristiques de demande) en tant que techniques subjectives en général. Suivant la classification présentée ci-dessous, les méthodes à base de questionnaires et de performance, peuvent être nommées techniques « objectives », contrairement au groupe de techniques subjectives. Les techniques objectives et les techniques subjectives de la SA ont des avantages relatifs.

Les mesures subjectives de la SA, comme d'autres types de mesures subjectives, sont susceptibles d'erreur. Les évaluations objectives, d'autre part, sont souvent coûteuses ou difficiles à gérer (les mesures physiologiques) ou excessivement intrusives (les méthodes à base de questionnaires), [41]. Une solution semblerait se situer dans la conception de sous-tâches appropriées qui permettent la collection naturelle de données comportementales, avec un minimum d'interruption de tâche. Si de telles données peuvent être enregistrées en tant qu'élément de la routine normale de l'opérateur (ou légèrement modifiée), l'acceptation est susceptible d'être beaucoup plus importante. Ceci est particulièrement critique dans des scénarios plus réalistes, tels que les simulations de grande fidélité ou l'environnement opérationnel. Les mesures de performance peuvent éviter plusieurs difficultés théoriques, actuellement abordées par la communauté scientifique. C'est le point de départ pour les sous-tâches, ou « la performance implicite », des mesures de la SA.

Finalement, il est important de réitérer qu'aucune technique n'est susceptible d'être appropriée pour tous les domaines de la SA. On discute que les mesures « objectives » (i.e. temps de réponse) et subjectives (i.e. autoévaluations) de la SA pourraient influencer fondamentalement différents aspects de conscience. Considérant que l'ancienne tendance se relie à la conscience, la dernière tendance relie plus à la « conscience de cette conscience », ou la réaction subjective de l'individu à sa propre conscience. Le cas dans lequel deux mesures de la SA (valides et fiables) peuvent être dissociées, met en évidence un point important : dans le développement d'un système, ou dans n'importe quel cas où l'acceptation de l'utilisateur est essentielle, l'évaluation subjective des opérateurs peut influencer l'acceptation et, finalement, l'utilité du système. Si un aiguilleur du ciel, par exemple, n'identifie pas (subjectivement) les avantages de SA d'un nouvel outil, il utilisera très probablement des moyens raffinés pour éviter son utilisation.

I.3.4 Surveillance de système

La croissante conscience de la situation implique une évaluation de la performance de la surveillance humaine des systèmes complexes, notamment dans le domaine de l'aviation et des réseaux ferroviaires. Deux facteurs conduisent cette tendance. Premièrement, l'acceptation du fait que les humains sont, par nature, des moniteurs peu performants. Les tâches de vigilance passive ne sont pas les activités pour lesquelles l'opérateur humain est bien adapté. Deuxièmement, l'utilisation croissante de l'automatisation signifie que la présence de telles tâches est susceptible d'augmenter à l'avenir. C'est-à-dire, l'automatisation limite l'opérateur, de plus en plus, dans le rôle de moniteur passif des systèmes automatisés.

L'étude expérimentale de la **vigilance** s'est développée pendant les années 50 sur la performance des moniteurs des sonars maritimes. La performance sur des tâches de vigilance peut être caractérisée par le niveau absolu de vigilance (en termes, par exemple, de taux de détection) aussi bien que par sa dégradation qui accompagne typiquement la **durée de la tâche**. Cette deuxième caractéristique de la vigilance de la performance, le changement de la surveillance de la performance dans le temps, est appelée **la baisse de la vigilance**. Il est connu depuis quelques décennies que la performance se dégrade, dans de telles conditions, pendant les premières dix minutes, [52].

La motivation derrière ce domaine d'étude a été, évidemment, la préoccupation que les opérateurs (par exemple des pilotes ou des contrôleurs aériens) pourraient manquer un signal critique (tel qu'un signal d'avertissement de probabilité de collision) avec des conséquences catastrophiques. Malheureusement, les humains n'ont pas encore produit un dispositif mécanique avec une fiabilité de 1.0 (c.-à-d., fiabilité parfaite). Néanmoins, les défaillances des systèmes automatisés sont inévitables. Si de tels échecs se produisent une fois que la performance de l'opérateur s'est dégradée pendant un long trajet, la capacité de l'opérateur de détecter et répondre rapidement à un défaut de fonctionnement sera fortement compromise. Les demandes de surveillance imposées, en partie, par l'automatisation sont en pleine croissance due à la prolifération de composantes accompagnant cet automatisation. Des recherches, [52], ont précisé que l'automatisation de n'importe quelle fonction augmente par trois le nombre de fonctions à surveiller : la fonction elle-même, le système automatisé, et l'indicateur du système automatisé sont maintenant des sources possible de dysfonctionnement. Ce fait augmente la probabilité de dysfonctionnement et augmente le nombre de dispositifs à surveiller (un feu vert en panne indique qu'il y a un problème avec un capteur, l'ampoule ou le système lui-même).

En cas de dysfonctionnement du système, l'augmentation du nombre de composants complique le diagnostic. Le problème potentiel de la surcharge de surveillance est aggravé par l'utilisation de la microélectronique qui permet la mise en place de moyens peu coûteux et compacts assurant une richesse d'information à l'opérateur. Pour empêcher que l'opérateur humain devienne lui-même surchargé d'information, le développement des systèmes doit prévoir quelle information exhiber et quelle information retenir.

L'évaluation de la surveillance de la performance dans un cadre de validation est paradoxale. Bien qu'il soit relativement facile d'identifier des métriques appropriées pour surveiller la performance, l'utilisation des méthodes de validation de façon appropriée est excessivement difficile. Ce phénomène est dû à la nature même de la tâche car les « signaux » critiques sont très rares. Un pilote d'avion peut passer une carrière entière sans rencontrer une extinction de moteur, ou un train d'atterrissage coincé. C'est pour cette raison que la validation peut ne jamais présenter de tels événements rares, avec une fréquence suffisante pour permettre une analyse statistique.

I.3.5 Travail d'équipe¹

Le travail d'équipe est un aspect essentiel de la gestion des systèmes complexes. Dans le cadre du trafic aérien, [52], le travail en équipe est défini comme :

« Un groupe de deux personnes ou plus qui interagissent dynamiquement et interdépendant dans des rôles, des fonctions et des responsabilités spécifiques assignés. Ils doivent s'adapter continuellement entre eux pour assurer l'établissement d'un fonctionnement sûr, ordonné et agile du trafic aérien ».

I.3.6 Confiance

La confiance est un terme bien connu dans la vie quotidienne. Nous parlons de la confiance que nous avons envers nos proches (famille, amis et collègues), la façon dans laquelle nous croyons en ce que nous voyons ou ce qui est dit (par exemple dans un journal), ou à quel point nous sommes confiants de que quelque chose travaille correctement (par exemple une automobile). Clairement, la confiance a plusieurs significations, mais elle peut être définie simplement comme confiance placée en une personne ou une chose, ou de façon plus précise, le degré de croyance dans la force, la capacité, la vérité ou la fiabilité d'une personne ou d'une chose. Dans le contexte des systèmes complexes homme-machine, [52], la confiance est définie comme suit :

¹ Par complétude, nous énonçons le cadre multi-opérateur.

« La confiance est la situation dans laquelle un utilisateur est disposé à agir sur la base des recommandations, des actions et des décisions fournis par une aide à la décision artificiellement intelligente ».

C'est une définition utile. Cependant, le terme « artificiellement intelligent » suggère que le but principal est l'utilisation de systèmes experts et de systèmes informatiques. C'est la raison pour laquelle le terme « outils basés sur l'ordinateur » est préféré.

Dans le domaine de la psychologie, la confiance joue un rôle décisif car elle participe en l'interaction entre les conditions particulières de stimulus et les comportements particuliers. En d'autres termes, il s'agit d'un état interne qui ne peut pas être mesuré directement mais est impliqué dans la base de certaines observations et mesures. Le degré de confiance dans l'automatisation pourrait, théoriquement au moins, être impliqué dans des mesures objectives d'exécution de contrôleur (par exemple fréquence, exactitude ou vitesse d'interaction), si le rapport entre ces mesures et l'automation peut être établi sans équivoque.

Une variable intervenante telle que la confiance peut être mesurée *subjectivement* en demandant à un opérateur de dire simplement comment il se sent. En effet, l'utilisation des échelles d'évaluation subjectives est le moyen le plus courant de mesurer la confiance. Il faut noter que si l'origine des estimations subjectives peut être modélisée, on peut convertir des variables intervenantes en mesures objectives.

La confiance est donc une construction composée de plusieurs éléments, dont les principaux identifiés dans la littérature scientifique sont la prévisibilité, la fiabilité, la foi (Trust), la sûreté, la robustesse, la familiarité, la compréhension, l'explication de l'intention, l'utilité, la compétence, la confiance en soi et la réputation.

I.3.7 Utilisabilité et acceptation

L'utilisabilité et l'acceptation d'un système de surveillance sont fondamentales à la sûreté et à l'efficacité des systèmes commandés, aussi bien qu'au confort des opérateurs eux-mêmes. De cette manière, le système devrait être instinctif et compatible avec des interactions quotidiennes. Dans [122], les auteurs définissent l'utilisabilité comme :

« Une mesure de la facilité avec laquelle un système peut être appris ou employé, de sa sûreté, de son effectivité et efficacité, et de l'attitude de ses utilisateurs envers elle ».

I.3.8 Erreur humaine

Le travail sous des conditions de responsabilité élevée, d'environnement potentiellement stressant, fait que la probabilité d'erreur humaine augmente, ainsi que la probabilité d'erreur système, malgré les techniques employées pour essayer de les éviter. Néanmoins, la compréhension des causes de l'erreur humaine à travers des techniques de mesure peut aider à souligner les problématiques envers lesquelles il faut être plus attentif.

Il est difficile de définir l'impact de l'erreur humaine d'une manière acceptable à différents praticiens dans différents domaines d'expertise, mais la définition simple peut être la suivante :

« Toute action ou inaction qui mène réellement ou potentiellement à des conséquences négatives de système, où plus d'une action possible est disponible »

Cette définition souligne l'action et l'omission et inclut les échecs de perception et d'attention, de mémoire, de processus décisionnels et de réponse d'exécution. La prérogative sur la définition d'erreur humaine peut sembler plus importante pour les autorités abordant des aspects légaux que pour les

spécialistes en facteurs humains, où l'appropriation du blâme porte peu de signification. Cependant, la définition de l'erreur humaine a un effet profond sur l'identification et la classification de l'erreur.

En réalité, l'erreur humaine tend à être multi-causale. L'erreur peut être une fonction du choix, de la formation et de l'expérience. Cependant, et d'une manière plus particulière pour cette thèse, la qualité de l'équipement et les conditions de fonctionnement ont un impact important sur la probabilité de l'erreur humaine. Il est souvent difficile de séparer les erreurs de conception, les erreurs système et les erreurs humaines, et plus particulièrement les erreurs liées à l'interaction homme-machine (HMI).

I.4 Métriques sur la performance humaine

Au long la section I.3 nous avons évoqué l'importance de trois groupes de métriques pour mesurer les facteurs humains, notamment pour la charge de travail : mesures physiologiques, subjectives, ou de performance. La Figure I.2 illustre les types de métriques issues de la performance humaine. Historiquement, ces méthodes (encore sujet d'étude) ont été appliquées dans le cadre de l'aviation pour permettre l'étude des éventuelles défaillances en conditions extrêmes, [52].



Figure I.2 : Métriques sur la performance humaine : mesures physiologiques, mesures subjectives et mesures de performance.

I.4.1 Mesures physiologiques

L'utilisation des mesures physiologiques repose sur l'hypothèse que des changements de la charge de travail provoquent des différences mesurables dans certains processus physiologiques de l'opérateur humain (généralement involontaires). L'avantage principal de (certains) indicateurs physiologiques est qu'ils restent disponibles même en l'absence de comportement manifeste. En conséquence, ils sont potentiellement utiles dans l'évaluation de la charge de travail dans des scénarios de validation dans lesquels les demandes de réponse sont lentes (comme dans les cockpits avion fortement automatisés). Par ailleurs, la plupart des mesures peuvent être collectées de façon continue et relativement non intrusive, grâce à la miniaturisation des technologies. Les inconvénients de ce type de mesures sont l'équipement spécialisé requis, la grande quantité de données nécessitant une analyse sophistiquée, la nécessité d'une expertise technique, les rapports signal-bruit critiques, etc. [7], [41].

Les indicateurs physiologiques les plus employés sont les mesures provenant de deux types de structures anatomiques distinctes, [94] : le système nerveux central (central nervous system, CNS) et le système nerveux périphérique (peripheral nervous system, PNS). Le CNS comprend l'encéphale et la

moelle épinière. Le PNS peut être divisé en le système nerveux somatique, contrôlant l'activation des muscles volontaires, et le système nerveux autonome (autonomic nervous system, ANS), qui contrôle les organes internes et est autonome dans le sens que les muscles associés ne sont pas sous un contrôle volontaire. Le ANS est encore subdivisé en le système nerveux sympathique (sympathetic nervous system, SNS) et le système nerveux parasympathique (parasympathetic nervous system, PNS). Tandis que le PNS maintient les fonctions corporelles, le SNS est dirigé vers les réactions de secours. Le SNS provoque des réactions de fuite ou de combat, en réaction à un stress ou à un stimulus, comme, par exemple, une augmentation de la fréquence cardiaque, de la sécrétion de salive et de sueur. Le PNS contrebalance ces effets en ralentissant la fréquence cardiaque, en dilatant les vaisseaux sanguins et en relâchant les fibres des muscles lisses involontaires.

I.4.1.1 L'activité cérébrale

Les mesures CNS incluent l'activité électrique, magnétique et métabolique du cerveau. L'électroencéphalogramme (electroencephalogram, EEG) est un ensemble de mesures provenant de l'activité électrique cérébrale.

L'analyse fréquentielle des EEG permet de déterminer le niveau d'endormissement d'un individu. On distingue typiquement quatre zones de fréquence de ce signal :

< 4 Hz.	ondes delta, présents durant un sommeil profond,
4 Hz. – 8 Hz.	ondes thêta, diminution importante de la vigilance,
8 Hz. – 13 Hz.	ondes alpha, diminution modérée de la vigilance,
13Hz. <	ondes bêta, prédominant pendant un état éveillé actif.

Cependant, les différences entre des différents individus peuvent être grandes. Des caractéristiques associées à ce signal EEG sont :

Sûreté	L'effet apparaît comme reproductible
Validité	La fréquence est proportionnelle au niveau de vigilance
Sensitivité	Très sensible, car la méthode es basée sur des mesures précises de l'activité cérébrale.
Diagnosticabilité	La fréquence des EEG augmente avec la charge de travail.
Mise en œuvre	Difficulté pour l'analyse des données ; relation signal-bruit important ; calibration personnalisée, coût élevé d'instrumentation et de support humain spécialisé.
Intrusivité	Electrodes EEG posées sur le cuir chevelu

D'autres métriques issues de l'EEG sont les potentiels relatifs à des événements (event-related potentials, ERP). En particulier, la métrique P_{300} (appelé de cette manière car due aux changements électriques qui se produisent environ 300 millisecondes après la présentation d'un stimulus) est utilisée pour mesurer la charge de travail. Néanmoins, cette méthode est réservée pour des recherches de base très détaillées de la charge de travail, [53]. Caractéristiques de la métrique P_{300} :

Sûreté	L'effet apparaît comme reproductible
Validité	L'amplitude du P_{300} est inversement proportionnelle à la charge de travail
Sensitivité	Très sensible, car la méthode est basée sur des mesures précises de l'activité cérébrale.
Diagnosticabilité	L'amplitude du P_{300} augmente avec la difficulté de la tâche.
Mise en œuvre	Difficulté pour l'analyse des données ; relation signal-bruit important ; calibration personnalisée, coût élevé d'instrumentation et de support humain spécialisé.
Intrusivité	Electrodes EEG posées sur le cuir chevelu

I.4.1.2 L'activité cardiaque

Le fonctionnement du cœur est associé au PNS et au SNS. Chaque contraction du cœur pompe le sang à travers tout le système circulatoire. La contraction est produite par des impulsions électriques qui peuvent être mesurées par un électrocardiogramme (electrocardiogram, ECG), [41], [53].

Des analyses en temps, en fréquence et d'amplitude du ECG permettent l'extraction de plusieurs mesures. Dans le domaine du temps, la détection des ondes *R* permet de mesurer le temps entre les pics donnant l'intervalle intra-battement (inter-beat interval, IBI). La **fréquence cardiaque** (heart rate, HR) est directement liée à la période du cœur (heart period, HP) ou à l'IBI. Elle permet de mesurer l'effort physique impliqué dans la réalisation d'une tâche. Des caractéristiques associées à l'activité cardiaque sont :

Sûreté	L'effet apparaît comme reproductible
Validité	le HR varie avec la charge de travail, mais aussi avec d'autres facteurs. La validité reste discutable.
Sensitivité	Sensible aux variations des demandes de la tâche, mais aussi à l'effort physique, émotions et stress
Diagnosticabilité	Diagnosticabilité limitée. Il s'agit d'un indice intégrant l'effet de toutes les demandes de la tâche et la réponse émotionnelle de l'opérateur.
Mise en œuvre	Enregistrement portable possible. L'analyse des données est sophistiquée et la relation avec la performance est complexe.
Intrusivité	Faible à moyenne. Les électrodes ne nécessitent pas une localisation précise.

Trois bandes fréquentielles caractérisent l'activité cardiaque :

0.02 Hz. – 0.06 Hz.	fréquences impliquées dans la régulation de température du corps,
0.07 Hz. – 0.14 Hz.	fréquences impliquées dans la régulation sanguine court terme,
0.15 Hz. – 0.50 Hz.	fréquences influencées par les fluctuations respiratoires.

La **puissance à 0.1 Hz.**, déterminée par une décomposition spectrale de la variation de la fréquence cardiaque (HRV), semble une bonne mesure de l'effort mental, [53]. Des caractéristiques sont :

Sûreté	Une mesure sûre, appliquée avec succès dans des environnements variés
Validité	En général, une variable valide.
Sensitivité	Grande sensibilité, mais vulnérable à la contamination par le stress et l'environnement.
Diagnosticabilité	Une mesure globale de l'effort.
Mise en œuvre	L'une des mesures physiologiques les plus pratiques. Non recommandé pour des tâches de courte durée, à cause de son analyse sur deux minutes d'enregistrement.
Intrusivité	Faible à moyenne. Les électrodes ne nécessitent pas une localisation précise.

I.4.1.3 L'activité oculaire

L'activité oculaire est difficile à classer entre les mesures physiologiques et les mesures de performance mais, traditionnellement, elle est considérée parmi les paramètres physiologiques, probablement due à l'une des techniques de mesure, l'électrooculogramme (electrooculogram, EOG).

La stratégie de recherche visuelle est indicative de la recherche d'informations. La **durée de fixation du regard** est liée à la difficulté de l'obtention/interprétation des informations, tandis que la **fréquence** l'est à l'importance de ces informations. Des caractéristiques des mesures de l'activité oculaire sont :

Sûreté	Existence de petites différences pouvant détecter la charge de travail dans différentes situations
Validité	Bonne validité en tant que mesure globale de la charge de travail

Sensitivité	Grande sensibilité, aussi aux facteurs émotionnels et environnementaux (lumière, etc.), ce qui ramène à un contrôle expérimental précis et difficile à maintenir.
Diagnosticabilité	Diagnosticabilité limitée.
Mise en œuvre	Équipement variée : EOG, enregistrement vidéo ou vision artificielle. Coût élevé d'instrumentation et besoin de calibration.
Intrusivité	Les électrodes du EOG sont placées autour des yeux ; intrusivité moyenne. Les techniques de vision sont d'intrusivité faible, voire nulle. Bonne acceptation de la part de l'opérateur.

Les clignements des yeux ont comme paramètres **le taux de clignements**, **la durée du clignement** et **la latence des clignements**, cette dernière liée au stimulus de l'occurrence d'une situation. Des recherches ont montré que les résultats sur le taux de clignement sont mitigés, tandis que la latence augmente et la durée de fermeture suite à une augmentation des demandes de la tâche, [94]. La **fréquence des clignements** est un indicateur de la fatigue. Les caractéristiques des mesures des clignements des yeux sont :

Sûreté	Des relations constantes entre les demandes de la tâche et la latence/durée des clignements.
Validité	Mesures valides de la charge de travail et de la fatigue.
Sensitivité	Sensibilité aux variations des demandes de la tâche.
Diagnosticabilité	Diagnostic de la charge de travail visuelle
Mise en œuvre	Équipement variée : EOG, enregistrement vidéo ou vision artificielle. Besoin de calibration.
Intrusivité	Les électrodes du EOG sont placées autour des yeux ; intrusivité moyenne. Les techniques de vision sont d'intrusivité faible, voire nulle. Bonne acceptation de la part de l'opérateur.

La réponse pupillaire est une mesure importante de la charge de travail, même si les changements les plus importants du **diamètre de la pupille** se produisent sous des facteurs émotionnels ou environnementaux (illumination). Caractéristiques :

Sûreté	Des différences petites mais sûres, permettant l'évaluation de la charge de travail dans des situations différentes.
Validité	Grand niveau de validité comme mesure globale de la charge de travail.
Sensitivité	Grande sensibilité aussi aux facteurs émotionnels et environnementaux.
Diagnosticabilité	Diagnosticabilité limitée.
Mise en œuvre	Pratique seulement dans des conditions de laboratoire. Équipement commercialement disponible mais coûteux.
Intrusivité	Peut être intrusive, selon la tâche.

I.4.2 Mesures subjectives

La classe des méthodes d'évaluation subjectives de la charge de travail repose sur l'auto-appréciation de l'opérateur de son effort dans l'exécution d'une certaine tâche. Leur utilisation fréquente est due au faible coût et à la facilité d'emploi (principalement sous la forme papier/crayon), à la validité potentiellement élevée pour les applications et à la facilité d'analyse. Pourtant, leur plus grand inconvénient est, peut-être, les seuils inhérents aux évaluations subjectives, [41], [52], [53]. Les sujets pourraient être incapables de s'évaluer d'une façon précise, à cause des limitations de mémoire, des perceptions imprécises, ou tout simplement le désir de dire au chercheur « ce qu'il veut entendre ». Malgré cela, les métriques subjectives sont demeurées populaires dans le domaine de la recherche, du développement et de la validation.

I.4.2.1 Autoévaluation instantanée

L'autoévaluation instantanée (instantaneous self assessment, ISA) a été développée comme simple outil pour permettre à l'opérateur d'estimer la charge de travail perçue durant l'exécution de tâches. L'opérateur donne une évaluation de la charge de travail constatée sur une échelle de 1 (très faible) à 5 (très forte). Les caractéristiques de cette mesure sont :

Sûreté	Sûre si les évaluations de l'opérateur sont précises.
Validité	Fortement corrélée avec la NASA TLX et avec d'autres mesures de la charge de travail.
Sensitivité	Sensibilité modérée, limité à une échelle à 5 niveaux.
Diagnosticabilité	Pas de diagnosticabilité.
Mise en œuvre	Une mesure très pratique, avec une demande d'équipement minimal.
Intrusivité	Faible.

I.4.2.2 Technique subjective d'évaluation de charge de travail

La technique subjective d'évaluation de charge de travail (subjective workload assessment technique, SWAT) inclut des échelles pour la charge du temps, de l'effort mental et du stress physiologique, chacune à trois niveaux, [53]. Les sujets doivent coder les 27 combinaisons possibles des niveaux sur trois échelles, avant de pouvoir évaluer les événements ou tâches particulières. Des caractéristiques associées à cette technique sont :

Sûreté	Sûre, même avec des retards de plus de 30 minutes.
Validité	Dimensions non validées empiriquement.
Sensitivité	En général, moins sensible que le TLX, mais testée extensivement
Diagnosticabilité	Multidimensionnelle : trois échelles (charge du temps, de l'effort mental et du stress physiologique)
Mise en œuvre	Deux temps : développement d'une échelle (27 combinaisons) et évaluation des événements.
Intrusivité	Légèrement plus « demandante » que d'autres mesures subjectives (pas 1 peut demander 45 min.). Acceptabilité plus faible que le TLX.

I.4.2.3 Indice de charge de la tâche

L'indice de charge de la tâche de la NASA (task load index, TLX) se mesure sur une échelle multidimensionnelle abordant six grandes dimensions de charge de travail (mentale, physique, temporelle, effort, performance et frustration), [43]. On peut mentionner comme caractéristiques de cet indice :

Sûreté	Fort niveau de sûreté
Validité	Validé extensivement
Sensitivité	Grande sensibilité
Diagnosticabilité	Multidimensionnelle : six échelles : mentale, physique, temporelle, effort, performance et frustration. Diagnostic différentiel. Une évaluation globale est possible.
Mise en œuvre	Deux étapes : évaluation de l'événement et comparaison couplée au processus pour déterminer l'importance de chaque facteur de la tâche.
Intrusivité	1 à 2 minutes pour compléter, hors ligne.

I.4.2.4 Echelle d'endormissement Karolinska

Une métrique subjective du niveau d'endormissement à dix niveaux est l'échelle d'endormissement Karolinska (Karolinska sleepiness scale, KSS), [7]. Il s'agit d'une échelle à dix niveaux où 0 est associé à un état complètement éveillé et 9 correspond à l'endormissement total. Des caractéristiques sont :

Sûreté	Sûre, mais de dynamique lente (évaluation toutes les 2 à 10 minutes d'intervalle)
Validité	Bon niveau de validité.
Sensitivité	Sensibilité modérée, mais il peut y avoir des seuils.
Diagnosticabilité	Pas de diagnosticabilité
Mise en œuvre	Evaluation du niveau d'endormissement en ligne, à intervalles réguliers
Intrusivité	Faible.

I.4.3 Mesures de performance

I.4.3.1 Métriques de performance de la charge de travail

Les mesures de performance essaient d'inférer la charge de travail par la mesure directe de l'exécution de tâche. Bien qu'il existe quelques légères différences de terminologie, cette classe est généralement acceptée pour englober deux types de méthodes : *mesures de tâche primaire* et *mesures de tâche secondaire*. Les deux méthodes reposent sur l'influence de l'incrément de la charge de travail sur la performance d'une certaine tâche.

Mesures de tâche primaire

Les méthodes de tâche primaire mesurent directement la performance sur la tâche considérée et se concentrent sur des paramètres tels que l'intervalle inter-stimulus (ISI), la complexité de commande et le nombre de sources d'information. Les mesures de tâche primaire impliquent typiquement la variation d'un certain paramètre (par exemple, complexité du tracking de vol en aviation, ou maintien du véhicule sur la voie en conduite automobile) qui affectera des demandes de tâche au point que l'exécution tombe au-dessous d'un certain critère de performance, fournissant de ce fait une mesure de capacité résiduelle aux attributions de ressource. Les mesures de tâche primaire ont l'avantage de pouvoir directement relier la charge de travail à la performance du système.

Pour des tâches manipulant des événements discrets qui nécessitent des réponses bien définies, les mesures plus typiques sont :

- **temps de réaction** (reaction time, RT) : qui est le temps entre la présentation d'un stimulus et l'exécution d'une réponse. Pour des tâches manipulant des RT plus petits que quelques secondes, il faut mesurer le RT à la milliseconde près.
- **précision** : souvent représentée sous forme de pourcentage ou proportion d'erreurs.

D'habitude, il est conseillé d'enregistrer les deux types de mesures, car le compromis vitesse/précision est une caractéristique présente dans plusieurs tâches. Un opérateur peut augmenter la vitesse en sacrifiant l'exactitude. Si seulement des mesures de temps de réaction sont disponibles, on pourrait incorrectement conclure que la performance s'est améliorée.

Une difficulté d'interprétation émerge parfois quand il existe un échange entre les mesures de vitesse et d'erreur. Sans la connaissance de la fonction d'échange, il est difficile de déterminer si la performance s'est simplement déplacée en un point différent sur la même fonction vitesse/précision ou s'est réellement améliorée ou dégradée. Dans la pratique, les mesures de précision et de vitesse sont consistantes.

Pour des tâches impliquées en mouvement continu (commande du volant, – steering –), une mesure courante est :

- **racine du carré de l'erreur moyenne** (root mean square, RMS, error) : L'erreur sous forme de distance entre le point actuel et le pont désiré. L'utilisation du RMS, plutôt que la moyenne arithmétique, pénalise l'inconsistance.

Pour les tâches de vigilance, caractérisées par la présentation de signaux faibles et irréguliers, la théorie de la détection des signaux (signal detection theory, SDT) est souvent utilisée. Les réponses sont codées comme montré dans le Tableau I.1.

TABLEAU I.1
RESULTATS DE LA REPONSE-STIMULUS POUR UNE TACHE DE DETECTION.

		stimulus	
		Non	Oui
réponse	Non	rejet correct	perte
	Oui	fausse alarme	bonne détection

Mesures de tâche secondaire

Quand une tâche est ajoutée à la tâche primaire, on peut utiliser les mesures pour les tâches secondaires. On peut appliquer deux paradigmes à la performance de tâche double, [41]. Dans le « paradigme de chargement de tâche », la performance de la tâche secondaire est stable, même si la performance de la tâche primaire se dégrade. Dans le « paradigme subsidiaire de tâche », la performance de la tâche primaire a la priorité. Conséquemment, la performance de la tâche secondaire varie avec la difficulté.

Les mesures de tâche secondaire les plus utilisées sont, [53] :

- **Production d'intervalle** : le sujet doit préciser un taux spécifique. A mesure que la charge de travail augmente, les intervalles entre les taux deviennent de plus en plus variables.
- **Estimation du temps** : le sujet estime combien de temps s'est écoulée (depuis le début de la tâche). En général, des intervalles de temps sont progressivement sous-estimés à mesure que la charge de travail augmente.
- **Génération aléatoire de nombres** : à mesure que la charge de travail augmente, l'individu peut recourir aux ordres bien connus (par exemple, 1, 2, 3, ...) où la diminution de l'aspect aléatoire peut être mesurée mathématiquement.
- **Inspection du temps de réaction** : Un stimulus indépendant de la tâche primaire apparaît périodiquement ; le temps de réaction à ce stimulus est considéré comme un reflet des demandes de la tâche primaire.

I.4.3.2 Mesures de la charge de tâche

La « communauté de facteurs humains » fait souvent la distinction entre la *charge de tâche* (la tâche imposée à un opérateur) et la *charge de travail* (la réponse subjective de l'opérateur). La charge de travail est, en général, vue comme plus importante que la charge de tâche et incorpore des influences telles que la pression du temps, l'effort, la formation, les interrogatives, les stratégies, etc. Les constructions de la charge de tâches et de la charge de travail sont parfois vues comme analogues au stress et à l'effort dans le monde physique, dans lequel l'un résulte de l'autre. Bien qu'il y ait un certain désaccord au sein de la communauté scientifique au sujet du degré d'inférence que les mesures de charge de tâche permettent, elles ont été traditionnellement des sources importantes de données dans la charge de travail.

L'objectif de la tâche (ou l'interprétation subjective de la tâche) détermine le but à accomplir (en termes de précision ou vitesse) et en conséquence affecte la *demande* de la tâche. La demande de la tâche peut être définie en termes d'étapes d'opération qui déterminent la complexité de la tâche. À quel point la tâche est accomplie est une mesure objective, à savoir le *niveau de performance* achevé. Cependant, comment la tâche est expérimentée (la difficulté de tâche) n'est pas une propriété objective. La difficulté de la tâche dépend de la complexité de la tâche, des possibilités de l'opérateur (capacité), de son état et de la stratégie appliquée. La *charge de travail* mentale est déterminée directement par la difficulté de la tâche. L'assignation des ressources de traitement dépend de la difficulté de la tâche et la charge de travail mental est reflétée par la quantité de ressources assignées. La Figure I.3 montre la performance de la tâche et la charge de travail en fonction de la demande, [41].

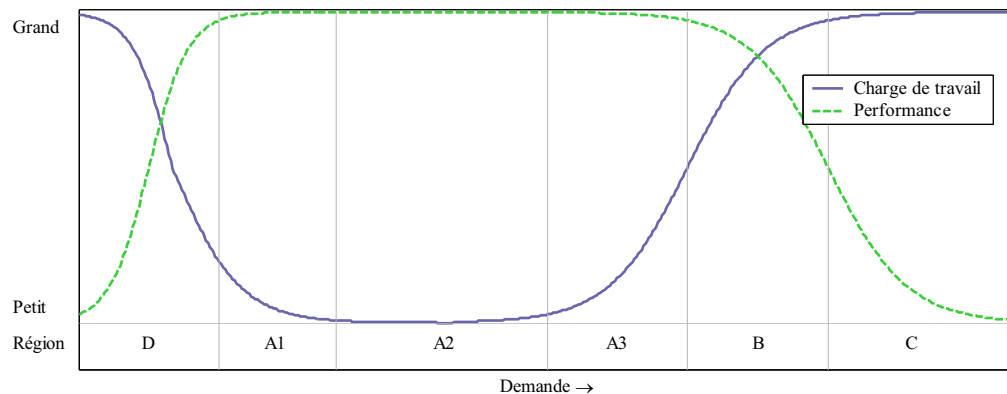


Figure I.3 : Charge de travail et performance en 6 régions. Dans la région D (désactivation) l'état de l'opérateur est affecté. Dans la région A2, la performance est optimale, l'opérateur peut se « débrouiller » facilement avec les requêtes de la tâche. Dans les régions A1 et A3, la performance n'est pas affectée mais l'opérateur doit exercer plus d'effort pour maintenir le niveau de performance. Dans la région B, ce n'est plus possible et la performance se dégrade, pendant que dans la région C la performance est dans un niveau minimum : l'opérateur est surchargé.

I.5. Surveillance des systèmes homme-machine

L'objectif de la **surveillance** des systèmes complexes est de mettre en œuvre des fonctions de détection, de localisation et de diagnostic de défaillances, hors ligne ou en ligne, pouvant dans ce dernier cas inclure des fonctions de détection des changements de mode de fonctionnement et des variations des performance du système, [149]. Dans le cas des systèmes homme-machine, la surveillance doit permettre d'obtenir des informations concernant l'évolution de la performance du système technologique, les actions de l'opérateur par rapport à la situation et l'interaction du système homme-machine lui-même avec l'environnement.

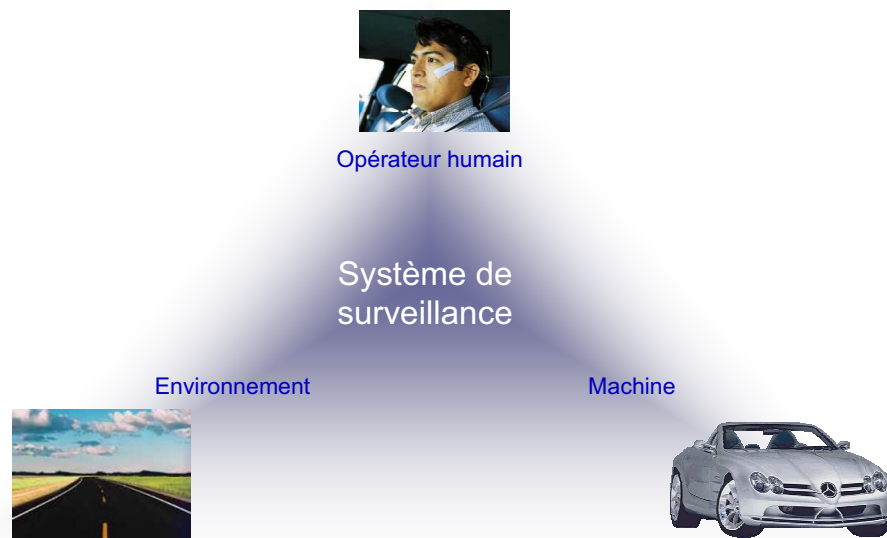


Figure I.4 : Structure d'un système homme-machine surveillé.

En ce qui concerne le système technologique, on peut profiter des informations issues de la surveillance et les utiliser dans un cadre de supervision visant à maintenir le système dans un mode opératoire optimal, que ce soit en situation normale ou en situation de faute. Elle doit conduire la couche

de commande (*commande supervisée*) et procéder en cas de faute à une accommodation, soit par la reconfiguration de la commande, soit par reconfiguration de ses objectifs, soit par la reconfiguration du système lui-même. Ce contrôle peut s'effectuer de manière partiellement ou entièrement automatique.

La conception d'un système de surveillance temps réel d'un système homme-machine doit exploiter toutes les mesures disponibles (décrites dans la dernière section) : mesures physiologiques, subjectives, ou de performance. Néanmoins, il faut être très attentif à la validité de chaque mesure, car le degré de représentativité des mesures sur un phénomène particulier est différent. En général, il existe une dissociation entre les mesures de performance et les autoévaluations. Néanmoins, différents travaux ont montré une dissociation entre les autoévaluations et les paramètres physiologiques, [7], [41]. Ces dissociations existent quand elles ne correspondent pas aux changements du comportement ou si une mesure indique une tendance positive tandis que l'autre indique une tendance négative.

Une fois les informations pertinentes identifiées, il faut savoir lesquelles pourrons être directement exploitées par le système de surveillance temps réel. En général, la totalité des mesures de performance peut être exploitée, car l'accès à ces informations est relativement facile et requiert un système d'acquisition des données. Par contre, seules les mesures physiologiques non intrusives peuvent être utilisées, tout en satisfaisant les contraintes de coût d'équipement. L'ensemble des métriques écarté pourra être exploité dans un cadre de validation du système. C'est de toutes les mesures subjectives et de l'ensemble des mesures physiologiques disponibles, non utilisés par le système de surveillance temps réel.

I.5.1 Les systèmes d'assistance intelligents

Les systèmes d'assistance intelligents (intelligent assistance systems, IAS) peuvent agir en médiateur entre l'opérateur humain et le système technologique. Ils constituent une interface entre le système de surveillance ou de supervision et l'opérateur humain. Du côté de la supervision, les IAS peuvent fournir des informations pertinentes de l'évolution de la performance du système HM entier et alerter en cas de risque (de dysfonctionnement ou de danger). Du côté de la commande supervisée, les IAS peuvent exécuter des fonctions du système qui peuvent être partagés par l'opérateur humain

L'idée d'un système d'aide intelligent peut s'expliquer avec l'exemple suivant. Dans un cockpit d'avion, un copilote humain partage le travail, mais pas la dernière responsabilité, avec le commandant. Ce dernier est le seul maître à bord : il peut consulter son copilote chaque fois qu'il le souhaite durant le vol ; cependant, il prendra la décision ultime. Si le commandant délègue une partie de sa responsabilité au copilote, alors il prendra cette délégation comme une tâche à exécuter. D'ailleurs, le commandant peut, à n'importe quel moment, décider d'arrêter l'exécution d'une tâche par le copilote, s'il le juge nécessaire. Malgré cela, le copilote peut avoir des initiatives personnelles, par exemple, vérifier des paramètres, être attentif à la situation actuelle, prédire des fautes déductibles, etc. Un copilote peut traiter la connaissance incluse dans le manuel d'opération, à sa propre initiative ou à la requête du commandant. Il doit être capable d'expliquer, au détail précis, les résultats de ce traitement.

I.6. Conclusion

Dans ce chapitre nous avons discuté l'impact de l'automatisation dans les systèmes complexes technologiques, rappelant ses avantages et ses inconvénients vis-à-vis de l'interaction de la partie technologique avec la partie humaine. Cela nous a permis d'aborder les aspects de l'interaction homme-machine, notamment leurs différents niveaux de coopération et de partage de tâches.

Ensuite, nous avons énoncé les différents types de défaillances dans les systèmes homme-machine, en nous concentrant sur la partie humaine. Cela nous a amené à étudier les causes de l'erreur humaine, abordant les différents domaines de la performance humaine impliqués dans l'interaction homme-

machine. Ceci nous permet d'aborder les différentes métriques que nous pouvons exploiter : mesures physiologiques, mesures subjectives et mesures de performance.

Nous finissons ce chapitre avec une introduction aux systèmes de surveillance des systèmes homme-machine, visant pour la surveillance de la performance du système HM et les actions de l'opérateur humain et, dans la mesure du possible, pour la supervision. De cette manière, nous présentons les éléments nécessaires pour la construction de systèmes d'assistance intelligents, commençant pour identifier les objectifs du système de surveillance à réaliser, l'identification des informations nécessaires et leur utilisation pour pouvoir inférer une décision (diagnostic ou commande).

Chapitre II

Le prétraitement des signaux, l'analyse temporelle et fréquentielle et extraction de caractéristiques

Sommaire

II.1	L'intérêt du prétraitement	28
II.2	Analyse en temps continu	28
II.2.1	Filtrage linéaire stationnaire	29
II.2.2	La transformée de Fourier	29
II.3	La révolution discrète	30
II.3.1	L'échantillonnage des signaux analogiques	31
II.3.2	Filtres discrets.....	32
II.3.2.1	Filtrage linéaire stationnaire discret.....	32
II.3.2.2	Série de Fourier.....	33
II.3.3	Signaux finis.....	34
II.3.3.1	Convolutions circulaires	34
II.3.3.2	Transformée de Fourier discrète	34
II.3.3.3	Transformée de Fourier rapide	35
II.3.3.4	Convolutions rapides	36
II.4	L'analyse temps-fréquence	36
II.4.1	Atomes temps-fréquence	36
II.4.2	Transformée de Fourier à fenêtre	37
II.4.2.1	Choix de la fenêtre.....	38
II.4.2.2	Transformée de Fourier à fenêtre rapide.....	39
II.4.3	Transformée en ondelettes.....	39
II.4.3.1	Transformée en ondelettes discrète.....	41
II.4.3.2	Transformée en ondelettes dyadiques.....	41
II.4.4	Analyse en ondelettes	42
II.4.4.1	Régularité.....	42
II.4.4.2	Maxima de la transformée en ondelettes et détection de singularités.....	43
II.4.5	Bases d'ondelettes	44
II.5	Extraction de caractéristiques	45
II.5.1	Analyse statistique.....	45
II.5.1.1	Statistiques du premier ordre	45
II.5.1.2	Statistiques du deuxième ordre	46
II.5.2	Analyse spectrale.....	46
II.5.3	Fenêtres d'analyse	47
II.5	Conclusion	47

II.1 L'intérêt du prétraitement

Traiter un signal, c'est extraire de l'information de mesures effectuées par des capteurs en vue d'atteindre un but spécifique. Ce but peut aller de la compréhension du monde physique (les physiciens, les géologues, les chimistes, les biologistes, etc.) à l'action sur ce monde (en robotique, dans les applications militaires, etc.), en passant par la reconstruction d'une information transmise au moyen d'un médium physique (c'est le cas des sons, des signaux de télécommunications, etc.).

Dès que l'on utilise un capteur pour mesurer une quantité, on est amené à effectuer un prétraitement. Un schéma relativement général est donné par la Figure II.1. Le phénomène physique qui fait réagir le capteur a, en général, été émis par une ou plusieurs sources et présente des variations temporelles (comme un signal sonore) ou spatiales (une scène en vision). Ces sources émettent ou réémettent des signaux qui sont transmis par un milieu, par exemple sous la forme d'ondes électromagnétiques, de sons, etc. Ces ondes qui portent l'information vers les capteurs peuvent être déformées : étalées, atténuées ou retardées dans le temps, en fonction de la fréquence, réfléchies par des obstacles, etc. L'analyse et la compensation des déformations de formes très variables apportées par le milieu de transmission forment, sans doute, l'un des problèmes centraux du traitement du signal.

Le signal mesuré par les capteurs est, la plupart du temps, entaché d'un bruit de mesure. Il présente des fluctuations qui ne sont pas uniquement déterminées par le phénomène étudié et qui peuvent en compliquer considérablement l'analyse. La présence de ces perturbations et les tentatives pour en atténuer les effets, en particulier dans le domaine des télécommunications, est à la base du développement de la théorie du signal. Dans d'autres applications, c'est la grande variabilité du signal émis qui a suscité le développement de techniques de traitement numérique du signal, en particulier les techniques de reconnaissances de formes.

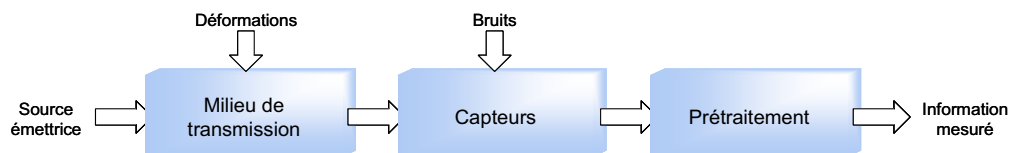


Figure II.1 : Schéma d'un système de perception et de prétraitement du signal.

Dans le prétraitement du signal, on dispose idéalement d'une certaine connaissance sur le signal étudié pour en extraire les informations utiles, ce qui n'est pas toujours le cas. Une autre difficulté au traitement des signaux mesurés dans le monde réel est que leurs caractéristiques, même en termes de probabilités, ne sont en général pas parfaitement connues et un traitement ne s'avèrera utile que si son efficacité est démontrée sur des données réelles. En particulier, la réflexion sur les informations dont dispose le « traiteur de signaux » est trop souvent éludée par le concepteur de systèmes de traitement, qui se contente d'en donner une formalisation imprécise ou de les ramener à des modèles connus et manipulables mais qui sont excessivement simplificateurs. La recherche de cette connaissance sur les signaux et la manière dont ils ont été émis ou perturbés est probablement la partie fondamentale de l'étude et de la mise au point d'un système de traitement de signaux.

II.2 Analyse en temps continu

C'est en 1807 que Fourier présente à l'Institut de France un mémoire où il prétend que toute fonction périodique peut être représentée comme une somme de sinusoides. Cette idée a eu un impact profond en analyse mathématique, en physique et en sciences de l'ingénieur, mais il aura fallu un siècle et demi pour comprendre la convergence des séries de Fourier et compléter la théorie des intégrales de Fourier.

II.2.1 Filtrage linéaire stationnaire

Dans de nombreuses applications fondées sur la propagation des ondes, on simplifie considérablement les problèmes étudiés en faisant des hypothèses sur la manière dont un système déforme un signal. Deux des hypothèses les plus importantes sont la **linéarité** et l'**invariance dans le temps** (stationnarité). A cette échelle, elles semblent bien représenter le comportement de nombreux systèmes physiques. Lorsqu'un système est linéaire et invariant dans le temps (linear time-invariant, LTI), on a les propriétés suivantes:

- **Linéarité** : Le principe de superposition spécifie que les entrées $x_1(t)$ et $x_2(t)$, ayant des sorties $y_1(t)$ et $y_2(t)$ respectivement, l'entrée $x(t)=x_1(t)+x_2(t)$ produit une sortie $y(t)=ay_1(t)+by_2(t)$, où a et b sont des constantes.
- **Stationnarité** : La stationnarité d'un opérateur L signifie que, si une entrée $x(t)$ fournit la sortie $y(t)$, le signal de sortie associé à l'entrée décalée $x(t-\tau)$ est aussi décalé de la même valeur de τ :

$$y(t-\tau)=Lx(t-\tau). \quad (\text{II.1})$$

- **Réponse impulsionnelle** : Une fois les hypothèses de linéarité et d'invariance temporelle vérifiées, on peut caractériser le système par sa réponse impulsionnelle $h(t)$. Par définition, $h(t)$ est la réponse du système à une impulsion de Dirac, $h(t)=L\delta(t)$. La connaissance de cette réponse permet le calcul du signal de sortie $y(t)$, quelle que soit la nature du signal d'entrée $x(t)$, à l'aide du produit de convolution :

$$y(t)=Lx(t)=\int_{-\infty}^{\infty}x(v)h(t-v)dv=\int_{-\infty}^{\infty}h(v)x(t-v)dv=h*x(t). \quad (\text{II.2})$$

- **Causalité** : Un système linéaire est physiquement réalisable si sa réponse impulsionnelle est nulle pour $t<0$. Cela signifie que le système n'anticipe pas la réponse à une sollicitation. Dans ce cas, le produit de convolution devient :

$$y(t)=\int_{-t}^{\infty}x(v)h(t-v)dv=\int_{-\infty}^{\infty}h(v)x(t-v)dv. \quad (\text{II.3})$$

- **Stabilité** : Un système linéaire est stable si pour une entrée $x(t)$ bornée la sortie $y(t)$ est également bornée. Le système linéaire sera stable si :

$$\int_0^{\infty}|h(v)|dv<+\infty. \quad (\text{II.4})$$

- **Réponse fréquentielle (fonctions de transfert)** : Les exponentielles complexes $e^{j\omega t}$ sont les vecteurs propres des opérateurs de convolution. En effet, la valeur propre :

$$\hat{h}(\omega)=\int_{-\infty}^{\infty}h(v)e^{-j\omega v}dv$$

est la valeur de la transformée de Fourier de h à la fréquence ω . Comme les sinusoïdes sont les vecteurs propres des systèmes LTI, il est alors possible de représenter l'opérateur Lx par ses valeurs propres $\hat{h}(\omega)$. L'analyse de Fourier montre que, lorsque x satisfait des conditions assez faibles, on peut l'écrire comme une intégrale de Fourier.

II.2.2 La transformée de Fourier

La transformée de Fourier, ou intégrale de Fourier, a été développée initialement pour étudier les fonctions de durée finie, et étendue aux fonctions périodiques. Nous donnerons les formules les plus utiles en traitement des signaux à temps continu. Notons aussi que la transformée de Fourier est un cas particulier de la transformée de Laplace, qui est plus utilisée par les automaticiens, en particulier pour caractériser plus facilement la stabilité des systèmes linéaires.

L'intégrale de Fourier (**transformée de Fourier**) :

$$\hat{x}(\omega)=\int_{-\infty}^{\infty}x(t)e^{-j\omega t}dt \quad (\text{II.5})$$

mesure « combien » il y a d'oscillation dans x à la fréquence ω . Si $x \in \mathbf{L}^1(\mathbb{R})$ cette intégrale converge, et

$$|\hat{x}(\omega)| \leq \int_{-\infty}^{\infty}|x(t)|dt < +\infty. \quad (\text{II.6})$$

La transformée de Fourier est une fonction bornée et continue en ω . Si $\hat{x} \in \mathbf{L}^1(\mathfrak{R})$ et est également intégrable, alors la **transformée de Fourier inverse** est :

$$x(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{x}(\omega) e^{-i\omega t} d\omega, \quad (\text{II.7})$$

qui est une somme de sinusoides $e^{i\omega t}$ d'amplitude $\hat{x}(\omega)$. La transformée de Fourier inverse (II.7), implique que x soit continue. Son extension à l'espace $\mathbf{L}^2(\mathfrak{R})$ démontre son existence pour les fonctions discontinues, [98].

En traitement du signal, la propriété la plus importante de la transformée de Fourier est le théorème de **convolution** (Il s'agit d'une autre manière d'exprimer le fait que les sinusoides $e^{i\omega t}$ sont des vecteurs propres des opérateurs de convolution). Si $x \in \mathbf{L}^1(\mathfrak{R})$ et $h \in \mathbf{L}^1(\mathfrak{R})$, la fonction convolution $g = h * x$ a comme transformée de Fourier :

$$\hat{g}(\omega) = \hat{h}(\omega) \hat{x}(\omega). \quad (\text{II.8})$$

En conséquence, la réponse $Lx = g = h * x$ d'un système LTI peut être calculé par

$$Lx(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{h}(\omega) \hat{x}(\omega) e^{i\omega t} d\omega. \quad (\text{II.9})$$

Chaque composante fréquentielle d'amplitude $\hat{x}(\omega)$ est amplifiée ou atténué d'un facteur $\hat{h}(\omega)$. C'est la raison pour laquelle une telle convolution s'appelle un **filtrage fréquentiel**, et $\hat{h}$ la **fonction de transfert** du filtre.

Le Tableau II.1 résume les propriétés les plus importantes de la transformée de Fourier, qui sont souvent utilisées dans les calculs.

TABLEAU II.1
Propriétés les plus importantes de la transformée de Fourier.

PROPRIÉTÉ	FONCTION	TRANSFORMÉE DE FOURIER	
	$x(t)$	$\hat{x}(\omega)$	
Inverse	$\hat{x}(t)$	$2\pi x(-\omega)$	(II.10)
Convolution	$x_1(t) * x_2(t)$	$\hat{x}_1(\omega) \hat{x}_2(\omega)$	(II.11)
Multiplication	$x_1(t) x_2(t)$	$\frac{1}{2\pi} \hat{x}_1(\omega) * \hat{x}_2(\omega)$	(II.12)
Translation	$x(t-v)$	$e^{-jv\omega} \hat{x}(\omega)$	(II.13)
Modulation	$e^{-j\xi t} x(t)$	$\hat{x}(\omega - \xi)$	(II.14)
Changement d'échelle	$x(t/s)$	$ s \hat{x}(s\omega)$	(II.15)
Dérivations en temps	$x^{(p)}(t)$	$(j\omega)^p \hat{x}(\omega)$	(II.16)
Dérivations en fréquence	$(-jt)^p x(t)$	$\hat{x}^{(p)}(\omega)$	(II.17)
Conjugaison complexe	$x^*(t)$	$\hat{x}^*(-\omega)$	(II.18)
Symétrie Hermitienne	$x(t) \in \mathfrak{R}$	$\hat{x}(-\omega) = \hat{x}^*(\omega)$	(II.19)

II.3 La révolution discrète

Pendant les années 1950, le traitement du signal digital a été utilisé pour simuler des circuits électroniques. Aujourd'hui, les algorithmes digitaux se sont insérés dans la plupart des forteresses traditionnelles du traitement du signal analogique (image, son, transmission, etc.). Il faut reconnaître que les calculs analogiques des circuits électroniques sont plus rapides que ceux des circuits digitaux implémentés sur des microprocesseurs, mais ils sont moins souples et moins précis. Ainsi, les microprocesseurs remplacent les circuits analogiques dès que leur rapidité de calcul est suffisante pour les applications temps réel.

II.3.1 L'échantillonnage des signaux analogiques

Qu'il s'agisse de signaux sonores ou d'images, l'échantillonnage des signaux analogiques permet l'obtention de signaux discrets. La façon la plus simple de discrétiser un signal analogique x est d'enregistrer ses valeurs $\{x(nT)\}_{n \in \mathbb{Z}}$ à intervalles régulières de T . L'interpolation de ces échantillons peut donner une approximation de $x(t)$ à n'importe quel instant de $t \in \mathbb{R}$.

Un signal discret peut se représenter comme une somme d'impulsions de Dirac¹. A tout échantillon $x(nT)$ on associe un Dirac $x(nT)\delta(t-nT)$ situé en $t=nT$. Ainsi, un échantillonnage uniforme de x correspond à une somme pondérée de Diracs :

$$x_d(t) = \sum_{n=-\infty}^{\infty} x(nT)\delta(t-nT). \quad (\text{II.20})$$

Comme la transformée de Fourier de $\delta(t-nT)$ vaut $e^{-inT\omega}$, la transformée de Fourier de x_d est une série de Fourier :

$$\hat{x}_d(\omega) = \sum_{n=-\infty}^{\infty} x(nT)e^{-inT\omega}. \quad (\text{II.21})$$

Pour comprendre comment calculer $x(t)$ à partir de ses échantillons $f(nT)$, et donc x à partir de x_d , il faut relier les transformées de Fourier $\hat{x}$ et $\hat{x}_d$. La transformée de Fourier d'un signal discret x_d obtenu par échantillonnage de x est :

$$\hat{x}_d(\omega) = \frac{1}{T} \sum_{k=-\infty}^{\infty} \hat{x}\left(\omega - \frac{2k\pi}{T}\right), \quad (\text{II.22})$$

ce qui rend $\hat{x}$ $2\pi/T$ périodique, par sommation de ses translatées $\hat{x}(\omega - 2k\pi/T)$.

C'est à partir de ce résultat que Whittaker a démontré le **théorème d'échantillonnage** en 1935 dans un livre sur l'interpolation, pour être redécouverte par Shannon en 1949 pour l'appliquer à la théorie de l'information, [98]. Le théorème d'échantillonnage impose que le support de $\hat{x}$ soit dans $[-\pi/T, \pi/T]$, ce qui garantit que x n'ait pas de variation brutale entre deux échantillons consécutifs ; alors, la formule d'interpolation résultante est :

$$x(t) = \sum_{n=-\infty}^{\infty} x(nT)h_T(t-nT), \text{ avec } h_T(t) = \frac{\sin(\pi t/T)}{\pi t/T}. \quad (\text{II.23})$$

La Figure II.2 illustre les différentes étapes d'un échantillonnage puis d'une reconstruction à partir des échantillons, dans le domaine temporel et le domaine fréquentiel.

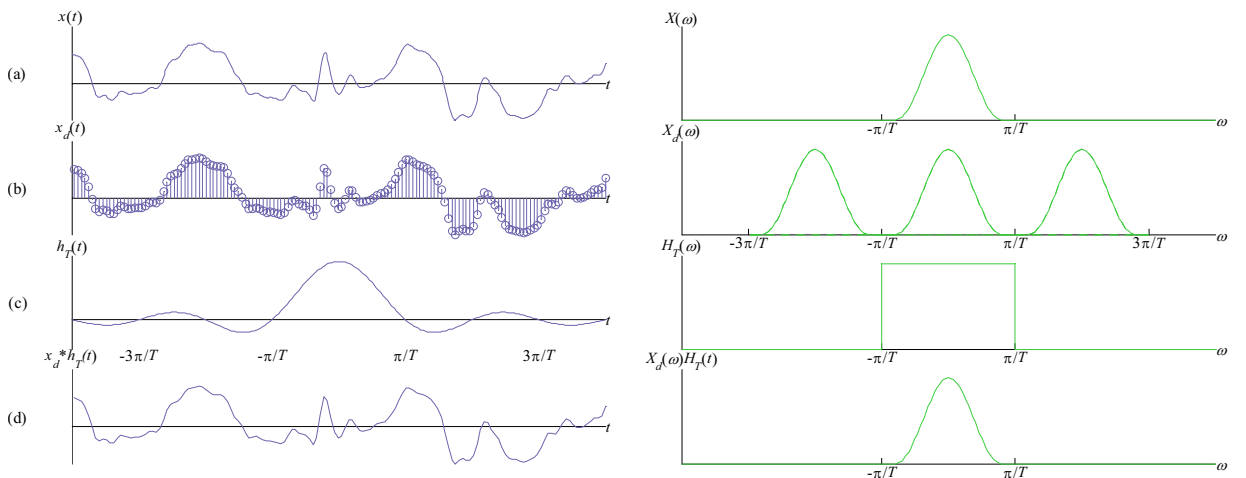


Figure II.2 : (a) Signal x et sa transformée de Fourier $\hat{x}$. (b) Un échantillonnage uniforme de x rend sa transformée de Fourier périodique. (c) Passe bas idéal. (d) Le filtrage de (b) par (c) reconstitue x .

¹ Par simplicité, on abrégera « impulsion de Dirac » par « Dirac ».

Souvent, des contraintes sur les capacités de calcul et de mémoire imposent le pas d'échantillonnage et, en général, le support de $\hat{x}$ n'est pas dans $[-\pi/T, \pi/T]$. Dans ce cas, la formule d'interpolation (II.23) ne redonne pas $x(t)$. Si le support de $\hat{x}$ déborde de $[-\pi/T, \pi/T]$, le support de $\hat{x}(\omega - 2k\pi/T)$ intersecte $[-\pi/T, \pi/T]$ pour plusieurs $k \neq 0$, Figure II.3. Ce repliement des composantes de haute fréquence sur un intervalle à basse fréquence est appelé **repliement spectral** (aliasing), [110].

Pour utiliser le théorème de l'échantillonnage en présence de repliement spectral, x est approximé par le plus proche signal $\tilde{x}$ dont la transformée de Fourier est à support dans $[-\pi/T, \pi/T]$. Le filtrage de x par h_T évite le repliement spectral en supprimant toutes les fréquences au-delà de π/T . Un convertisseur analogique-digital est donc composé d'un filtre qui limite le support fréquentiel à $[-\pi/T, \pi/T]$, suivi d'un échantillonnage uniforme de période T .

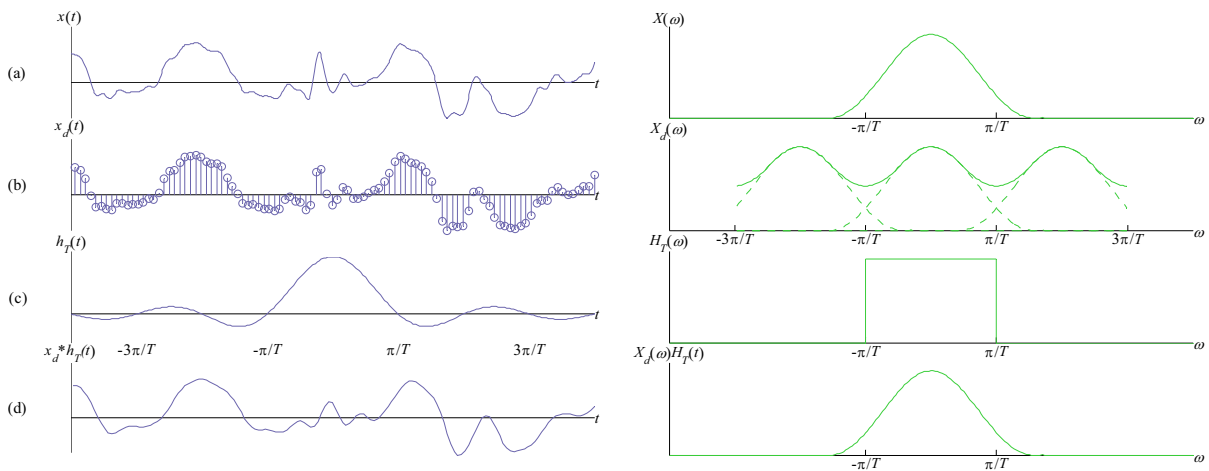


Figure II.3 : (a) Signal x et sa transformée de Fourier $\hat{x}$. (b) Repliement spectral dû à un recouvrement des $\hat{x}(\omega - 2k\pi/T)$ pour différentes valeurs k (en pointillés). (d) Le filtrage de (b) par (c) génère un signal à base fréquence qui diffère de x .

Le théorème d'échantillonnage donne une condition suffisante pour reconstituer un signal à partir de ces échantillons, mais l'on peut établir d'autres conditions suffisantes pour d'autres schémas d'interpolation, [110].

II.3.2 Filtres discrets

Le filtrage linéaire effectué par des dispositifs électroniques est l'une des fonctions fondamentales en traitement du signal à temps continu. L'avènement du traitement numérique a conduit à une amélioration importante de ces dispositifs en termes de fiabilité, de reproductibilité, de souplesse et de complexité des fonctions réalisables.

La précision peut être améliorée indéfiniment en augmentant la taille des mémoires. Ainsi deux circuits de filtrage identiques donneront exactement le même résultat, ce qui n'est pas le cas pour les traitements analogiques, notamment par exemple à cause du vieillissement des composants.

II.3.2.1 Filtrage linéaire stationnaire discret

Les algorithmes classiques de traitement du signal discret reposent essentiellement sur des opérateurs linéaires stationnaires, [110]. La stationnarité est limitée à la grille d'échantillonnage. Pour simplifier les notations, supposons $T=1$ et $x[n]$ les valeurs des échantillons. Alors, les systèmes LTI discrets satisfont :

- *Linéarité au principe de superposition* : si pour les entrées $x_1[n]$ et $x_2[n]$, ayant des sorties $y_1[n]$ et $y_2[n]$ respectivement, l'entrée $x[n] = x_1[n] + x_2[n]$ produit une sortie $y[n] = ay_1[n] + by_2[n]$, où a et b sont des constantes.
- *Stationnarité* : Un opérateur L est stationnaire si pour une entrée $x[n]$ retardé de $p \in \mathbb{Z}$, génère une sortie également retardée de p :

$$y[t-p] = Lx[t-p]. \quad (\text{II.24})$$

- *Réponse impulsionnelle discrète* : Soit $h[n] = L \delta[n]$. La linéarité et la stationnarité impliquent :

$$y[n] = Lx[n] = \sum_{p=-\infty}^{\infty} x[p]h[n-p] = x * h[n]. \quad (\text{II.25})$$

Un opérateur linéaire stationnaire discret se calcule donc par convolution discrète. Si $h[n]$ est à support fini, la somme (II.25) se calcule en un nombre fini d'opérations, donnant lieu aux **filtres à réponse finie** (finite impulse response, FIR). Certaines convolutions avec des **filtres à réponse infinie** (infinite impulse response, IIR) peuvent se calculer avec un nombre fini d'opérations, à condition qu'elles puissent se réécrire sous forme récurrente (II.30).

- *Causabilité* : Un filtre discret L est causal si $L \delta[n]$ ne dépend que des valeurs $x[n]$ pour $n \leq p$. Donc, la formule de convolution (II.25) implique $h[n] = 0$ pour $n < 0$.
- *Stabilité* : Le filtre est stable si toute entrée bornée $h[n]$ engendre une sortie bornée $L \delta[n]$:

$$|Lx[n]| \leq \sup_{n \in \mathbb{Z}} |z[n]| \sum_{k=-\infty}^{+\infty} |h[k]|. \quad (\text{II.26})$$

cela implique que $\sum_{n=-\infty}^{+\infty} |h[n]| < +\infty$, ce qui signifie que $h \in \ell^1(\mathbb{Z})$. On peut vérifier que cette condition nécessaire est également suffisante.

- *Fonction de transfert* : La transformée de Fourier joue un rôle fondamental dans l'analyse des opérateurs stationnaires discrets, car les sinusoides discrètes $e[n-\omega] = e^{i\omega n}$ en sont les vecteurs propres :

$$\hat{h}(\omega) = \sum_{p=-\infty}^{+\infty} h[p]e^{-i\omega p}.$$

II.3.2.2 Série de Fourier

La série de Fourier est un cas particulier de la transformée de Fourier pour des signaux qui sont des sommes de Dirac. Pour tout $n \in \mathbb{Z}$, comme $e^{-j\omega n}$ est de période 2π , les séries de Fourier sont aussi de période 2π . Une fonction périodique 2π est entièrement déterminée par sa restriction à $[-\pi, \pi]$, ce qui implique que la famille de fonctions $\{e^{-ik\omega}\}_{k \in \mathbb{Z}}$, est une base orthonormée de $L^2[-\pi, \pi]$. En conséquence, si $x \in \ell^2(\mathbb{Z})$, alors la série de Fourier :

$$\hat{x}(\omega) = \sum_{n=-\infty}^{\infty} x[n]e^{-i\omega n} \quad (\text{II.27})$$

peut s'interpréter comme la décomposition de $\hat{x} \in L^2[-\pi, \pi]$ sur une base orthonormée. Les coefficients de la série de Fourier s'écrivent donc comme des produits scalaires :

$$x[n] = \langle \hat{x}(\omega), e^{-j\omega n} \rangle = \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{x}(\omega) e^{-i\omega n} d\omega. \quad (\text{II.28})$$

La série de Fourier satisfait les propriétés suivantes, [98] :

- *Convergence ponctuelle* : L'identité (II.27) s'entend au sens de la convergence en moyenne quadratique :

$$\lim_{N \rightarrow +\infty} \left\| \hat{x}(\omega) - \sum_{k=-N}^N x[k]e^{-i\omega k} \right\| = 0.$$

- *Convolutions* : Si $x \in \ell^1(\mathbb{Z})$ et $h \in \ell^1(\mathbb{Z})$, alors $g = h * x$ et $\hat{g}(\omega) = \hat{h}(\omega)\hat{x}(\omega)$. La formule de reconstruction (II.28) montre que le signal filtré peut s'écrire :

$$x * h[n] = \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{h}(\omega)\hat{x}(\omega)e^{i\omega n} d\omega. \quad (\text{II.29})$$

La fonction de transfert $\hat{h}$ amplifie ou atténue les composantes fréquentielles $\hat{x}(\omega)$ de $x[n]$. Comme exemple, un **filtre récursif** calcule $h = Lx$ en résolvant l'équation récurrente :

$$\sum_{k=0}^K a_k x[n-k] = \sum_{k=0}^M b_k g[n-k], \quad (\text{II.30})$$

avec $b_0 \neq 0$. Si $g[n]=0$ et $x[n]=0$ pour $n < 0$, alors la dépendance de g par rapport à x est linéaire et stationnaire, donc $g=h*x$. La fonction de transfert s'obtient en calculant la transformée de Fourier de (II.30) donc :

$$\hat{h}(\omega) = \frac{\sum_{k=0}^K a_k e^{-ik}}{\sum_{k=0}^M b_k e^{-ik}}. \quad (\text{II.31})$$

Un calcul direct de la convolution $g[n]=h[n]*x$ à partir des $x[n-p]$ nécessiterait une infinité d'opérations, alors que (II.31) permet le calcul de $g[n]$ par $K+M$ additions et multiplications sur des valeurs passées de x et de g .

- *Multiplication par une fenêtre* : Une réponse impulsionnelle à support infini h peut s'approximer par une réponse impulsionnelle à support fini $\tilde{h}$ en multipliant h par une fenêtre g de support fini : $\tilde{h}[n]=g[n]*h[n]$. Il est possible de vérifier que ce produit revient à une convolution dans le domaine fréquentiel :

$$\hat{\tilde{h}}(\omega) = \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{h}(\xi) \hat{g}(\omega - \xi) d\xi = \frac{1}{2\pi} \hat{h} * \hat{g}(\omega). \quad (\text{II.32})$$

L'utilisation d'une fenêtre rectangulaire $g=1_{[-N,N]}$ présente d'importants lobes latéraux loin de $\omega=0$. Le correspondant $\hat{\tilde{h}}$ est, donc, une mauvaise approximation de $\hat{h}$. La fenêtre de Hanning :

$$g[n] = \cos^2\left(\frac{\pi n}{2N}\right) 1_{[-N,N]}[n].$$

est plus régulière et a donc une transformée de Fourier plus concentrée dans les basses fréquences. En [98], sont présentées les propriétés spectrales d'autres fenêtres.

II.3.3 Signaux finis

Dans la pratique, x n'est connu que sur un support fini, $0 \leq n < N$. Il faut donc modifier les convolutions pour tenir compte des effets de bord en $n=0$ et $n=N-1$. Il faut également redéfinir la transformée de Fourier pour des suites finies en vue des calculs numériques.

II.3.3.1 Convolutions circulaires

Soient x et h des signaux constitués de N échantillons. Pour calculer le produit de convolution :

$$\tilde{x} * \tilde{h}[n] = \sum_{p=-\infty}^{\infty} \tilde{x}[p] \tilde{h}[n-p], \text{ pour } 0 \leq n < N,$$

Il est nécessaire de connaître $\tilde{x}[n]$ et $\tilde{h}[n]$ au-delà de $0 \leq n < N$. Une approche consiste à étendre $\tilde{x}$ et $\tilde{h}$ par périodisation sur N échantillons : $x[n]=\tilde{x}[n \bmod N]$, $h[n]=\tilde{h}[n \bmod N]$. La convolution circulaire de ces deux signaux de période N est une somme restreinte à leur période :

$$x \circledast h[n] = \sum_{p=0}^{N-1} x[p] h[n-p] = \sum_{p=0}^{N-1} x[n-p] h[p], \quad (\text{II.33})$$

signal de période N .

II.3.3.2 Transformée de Fourier discrète

La transformée de Fourier discrète (discret Fourier transform, DFT) établit que tout signal de période N peut se décomposer comme une somme de sinusoïdes discrètes. De cette manière, la famille :

$$\left\{ e_k[n] = \exp\left(\frac{i2\pi kn}{N}\right) \right\}_{0 \leq k < N}$$

est une base orthogonale de l'espace des signaux de période N . Comme l'espace est de dimension N , toute famille orthogonale de N vecteurs en est une base orthogonale. On peut décomposer, sur cette base, tout signal x de période N :

$$x = \sum_{k=0}^{N-1} \frac{\langle x, e_k \rangle}{\|e_k\|^2} e_k. \quad (\text{II.34})$$

Par définition, la transformée de Fourier discrète de x est :

$$\hat{x}[k] = \langle x, e_k \rangle = \sum_{n=0}^{N-1} x[n] \exp\left(\frac{-i2\pi kn}{N}\right). \quad (\text{II.35})$$

Comme $\|e_k\|^2 = N$, (II.34) fournit une formule d'inversion de Fourier discrète :

$$x[n] = \frac{1}{N} \sum_{k=0}^{N-1} \hat{x}[k] \exp\left(\frac{-i2\pi kn}{N}\right). \quad (\text{II.36})$$

L'orthogonalité de la base implique également la formule de Plancherel :

$$\|x\|^2 = \sum_{n=0}^{N-1} |x[n]|^2 = \frac{1}{N} \sum_{k=0}^{N-1} |\hat{x}[k]|^2. \quad (\text{II.37})$$

Comme les $\{\exp(i2\pi kn)\}_{0 \leq n < N}$ sont des vecteurs propres de la convolution circulaire, le théorème de la convolution établit que si x et h sont de période N , la transformée de Fourier discrète de $g = x \otimes h$ est :

$$\hat{g}[k] = \hat{f}[k] \hat{h}[k]. \quad (\text{II.38})$$

Ceci montre qu'une convolution circulaire peut s'interpréter comme un filtrage fréquentiel discret. Ce résultat permet d'introduire les calculs rapides de convolution en utilisant la transformée de Fourier rapide.

II.3.3.3 Transformée de Fourier rapide

Pour un signal x à N échantillons, le calcul direct des N sommes de Fourier discrètes (II.35) nécessite N^2 additions et multiplications complexes. L'algorithme de la transformée de Fourier rapide (fast Fourier transform, FFT) réduit la complexité numérique à $C(N) = K(N \log_2 N)$ par une réorganisation de calculs.

Supposons que N est pair. Pour les fréquences paires, on regroupe des termes n et $n + N/2$:

$$\hat{x}[2k] = \sum_{n=0}^{N/2-1} (x[n] + x[n + N/2]) \exp\left(\frac{-i2\pi kn}{N/2}\right). \quad (\text{II.39})$$

Aux fréquences impaires, ce regroupement devient :

$$\hat{x}[2k+1] = \sum_{n=0}^{N/2-1} \exp\left(\frac{-i2\pi kn}{N}\right) (x[n] + x[n + N/2]) \exp\left(\frac{-i2\pi kn}{N/2}\right). \quad (\text{II.40})$$

L'équation (II.39) montre que les coefficients des fréquences paires s'obtiennent en calculant la transformée de Fourier du signal $N/2$ périodique :

$$x_p[n] = x[n] + x[n + N/2].$$

Les coefficients des fréquences impaires se déduisent de (II.40) en calculant la transformée de Fourier du signal $N/2$ périodique :

$$x_i[n] = \exp\left(\frac{-i2\pi n}{N}\right) (x[n] + x[n + N/2]).$$

Une transformée de Fourier inverse rapide se déduit de la transformée de Fourier rapide de son complexe conjugué $\hat{x}^*$ en remarquant que :

$$x^*[n] = \frac{1}{N} \sum_{k=0}^{N-1} \hat{x}^*[k] \exp\left(\frac{-i2\pi kn}{N}\right). \quad (\text{II.41})$$

Il existe plusieurs variantes de cet algorithme de calcul rapide. Le but est de diminuer la constante K . A ce jour, la transformée de Fourier rapide la plus efficace est la FFT par l'algorithme de *split-radix*, légèrement plus compliquée que la procédure précédente, [98].

II.3.3.4 Convolutions rapides

La faible complexité des algorithmes de transformée de Fourier rapide en font un moyen efficace de calcul des convolutions discrètes, grâce à la convolution circulaire (II.38). Alors, pour les signaux discrets x et g , non nuls pour $0 \leq n < M$, le signal causal :

$$g[n] = x * h[n] = \sum_{k=-\infty}^{+\infty} x[k]h[n-k] \quad (\text{II.42})$$

est non nul pour $0 \leq n < 2M$. Le calcul de ce produit de convolution nécessite $M(M+1)$ multiplications et additions. Pour calculer la FFT en utilisant la convolution circulaire (II.38), on écrit (II.42) comme suit :

$$a[n] = \begin{cases} x[n] & \text{si } 0 \leq n < M \\ 0 & \text{si } M \leq n < 2M \end{cases} \quad (\text{II.43})$$

$$b[n] = \begin{cases} h[n] & \text{si } 0 \leq n < M \\ 0 & \text{si } M \leq n < 2M \end{cases}. \quad (\text{II.44})$$

Définissant $c = a \otimes b$, on peut vérifier que $c[n] = g[n]$ pour $0 \leq n < 2M$. Les $2M$ coefficients non nuls de g s'obtiennent en calculant $\hat{a}$ et $\hat{b}$ à partir de a et b , puis en calculant la transformée de Fourier discrète inverse de $\hat{c} = \hat{a}\hat{b}$. Le calcul de la convolution par FFT se fait donc en un total de $3M \log_2 M + 11M$ multiplications réelles. Donc, pour $M \geq 32$, le calcul par FFT est plus rapide que le calcul direct. Pour $M \leq 16$, il est plus rapide de calculer la convolution directement, [98].

Dans le cas d'un signal x à L échantillons non nuls avec un signal causal plus petit h à M échantillons, la convolution s'effectue par une procédure de calcul par blocs. Le signal x est décomposé en L/M blocs x_r à M échantillons non nuls :

$$x[n] = \sum_{r=0}^{L/M-1} x_r[n-rM] \text{ avec } x_r[n] = x[n+rM] \mathbf{1}_{[0, M-1]}[n]. \quad (\text{II.45})$$

Pour chaque $0 \leq r < L/M$, les $2M$ échantillons non nuls de $g_r = x_r * h$ sont calculés par l'algorithme de convolution rapide utilisant FFT. La décomposition par blocs (II.45) implique :

$$x * h[n] = \sum_{r=0}^{L/M-1} g_r[n-rM]. \quad (\text{II.46})$$

L'addition de ces L/M signaux translatés de taille $2M$ est donc effectuée en $2L$ additions élémentaires. Au total, la convolution est calculée par $O(L \log_2 M)$ opérations.

II.4 L'analyse temps–fréquence

Un signal musical est un bon exemple de signal constitué de *fréquences* sonores qui varient dans le *temps*. Ces évolutions sont mises en évidence en décomposant le signal en fonctions élémentaires bien concentrées en temps et en fréquence. La transformée de Fourier à fenêtre et la transformée en ondelettes sont deux exemples importants de décomposition temps–fréquence. La mesure de « fréquences instantanées » illustre les limites de la résolution temps–fréquence, imposées par le principe d'incertitude.

II.4.1 Atomes temps–fréquence

Une transformée temps–fréquence linéaire corrèle le signal avec une famille de fonctions bien concentrées en temps et en fréquence. Ces fonctions sont nommées *atomes temps–fréquence*. Notons la famille générale d'atomes temps–fréquence $\{\phi_\gamma(t)\}_{\gamma \in \Gamma}$, où γ peut être un indice de dimension supérieur à 1. Des exemples sont l'atome de Fourier et l'atome d'ondelette. Un atome de Fourier à fenêtre s'obtient à partir d'une fenêtre g , que l'on translate de u et l'on module à la fréquence ξ :

$$\text{avec } \phi_{u,\xi}(t) = g_{u,\xi} = e^{i\xi t} g(t-u). \quad (\text{II.47})$$

Un atome d'ondelette est obtenu par une dilatation de s et une translation de u d'une ondelette, dite mère :

$$\phi_\gamma = \psi_{u,s}(t) = \frac{1}{\sqrt{s}} \psi\left(\frac{t-u}{s}\right). \quad (\text{II.48})$$

Les ondelettes et les atomes de Fourier à fenêtre ont une énergie bien localisée en temps, tandis que leur transformé de Fourier est essentiellement concentrée sur un intervalle de fréquences relativement étroit.

La tranche d'information contenue dans $\langle x, \phi_\gamma \rangle$ est représentée dans le plan temps–fréquence (t, ω) par une région dont la position et la taille dépendent de l'étalement de ϕ_γ en temps et en fréquence. La résolution de ϕ_γ est représentée par une **boîte de Heisenberg**, centrée en (u, ξ) , d'écart type $\sigma_t(\gamma)$, et dont l'écart type fréquentiel vaut $\sigma_\omega(\gamma)$, Figure II.4. Le théorème d'incertitude de Heisenberg, [98], montre que la surface de ce rectangle est supérieure ou égale à $1/2$:

$$\sigma_t \sigma_\omega \geq 1/2. \quad (\text{II.49})$$

Ce résultat limite les résolutions conjointes en temps et en fréquence de ϕ_γ . Il faut manipuler le plan temps–fréquence avec précaution car il n'existe pas de définition du point (t_0, ω_0) . Aucune fonction n'est entièrement concentrée en un temps t_0 et à une fréquence ω_0 . Seuls les rectangles de surface supérieure à $1/2$ peuvent correspondre à des atomes temps–fréquence.

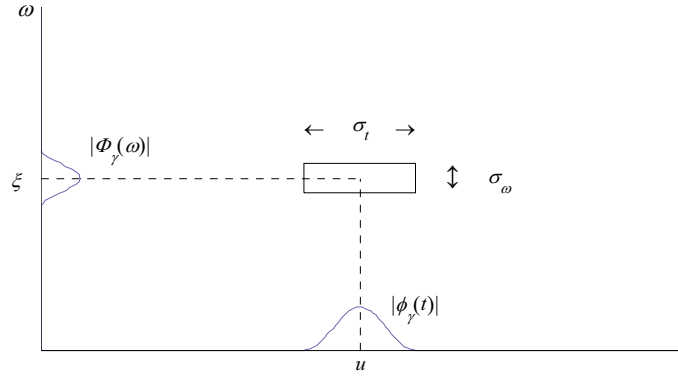


Figure II.4 : Boîte de Heisenberg représentant un atome temps–fréquence.

Supposons que, pour tout (u, ξ) , il existe un unique atome $\phi_{\gamma(u, \xi)}$ centré en (u, ξ) dans le plan temps–fréquence. La boîte de temps–fréquence de $\phi_{\gamma(u, \xi)}$ définit un voisinage de (u, ξ) sur lequel la **densité d'énergie** de x est mesurée par :

$$P_T x(u, \xi) = \left| \langle x, \phi_{\gamma(u, \xi)} \rangle \right|^2 = \left| \int_{-\infty}^{\infty} x(t) \phi_{\gamma(u, \xi)}^*(t) dt \right|^2. \quad (\text{II.50})$$

II.4.2 Transformée de Fourier à fenêtre

En 1946, Gabor a introduit les atomes de Fourier à fenêtre afin de mesurer « les variations fréquentielles » des sons. La transformée de Fourier fenêtrée est définie par :

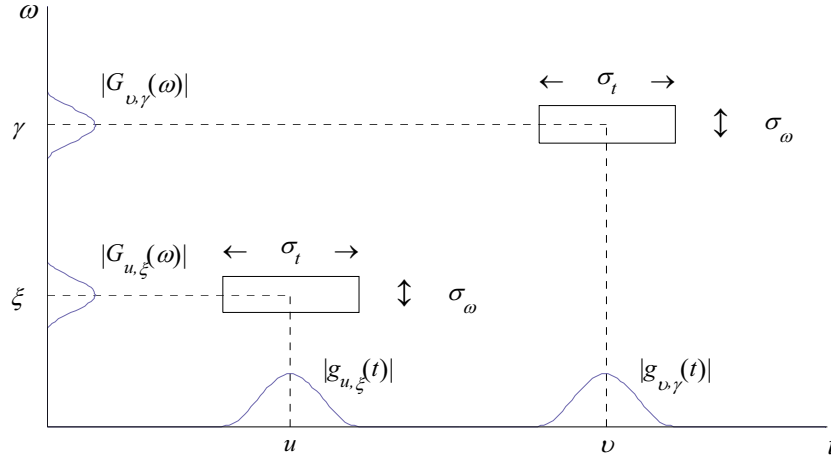
$$Sx(u, \xi) = \langle x, g_{u, \xi} \rangle = \int_{-\infty}^{\infty} x(t) g(t-u) e^{-j\xi t} dt, \quad (\text{II.51})$$

en utilisant l'atome de Fourier à fenêtre (II.47), que l'on translate de u et que l'on module à la fréquence ξ . La multiplication de $x(t)$ par $g(t-u)$ localise l'intégrale au voisinage de $t=u$. Le symbole S indique la transformée de Fourier à court terme (short-time Fourier transform, STFT).

La densité d'énergie, appelée **spectrogramme** est notée P_S :

$$P_S x(u, \xi) = |Sx(u, \xi)|^2 = \left| \int_{-\infty}^{\infty} x(t) g(t-u) e^{-j\xi t} dt \right|^2. \quad (\text{II.52})$$

Le spectrogramme mesure l'énergie de x dans le voisinage temps–fréquence de (u, ξ) défini par la boîte de Heisenberg de $g_{u, \xi}$. Comme $g_{u, \xi}$ est paire, elle est centrée fréquentiellement en ξ et symétrique par rapport à u , [98]. En plus, l'étalement en temps et en fréquence de g est indépendant de u et de ξ . En conséquence, l'atome $g_{u, \xi}$ a une boîte de Heisenberg de surface $\sigma_t \sigma_\omega$, centrée en (u, ξ) , Figure II.5. La taille de cette boîte ne dépend pas de (u, ξ) , ce qui implique que la résolution de la transformée de Fourier à fenêtre reste la même sur tout le plan temps–fréquence.

Figure II.5 : Boîte de Heisenberg de deux atomes de Fourier à fenêtre $g_{u,\xi}$ et $g_{v,\gamma}$.

Lorsque les coordonnées temps–fréquence (u, ξ) parcourent $\mathfrak{R}^2$, les boîtes de Heisenberg des atomes $g_{u,\xi}$ recouvrent tout le plan temps–fréquence. On peut donc s'attendre à pouvoir **reconstituer** x à partir de sa transformée de Fourier à fenêtre $Sx(u, \xi)$:

$$x(t) = \frac{1}{2} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} Sx(u, \xi) g(t-u) e^{-j\xi t} d\xi du. \quad (\text{II.53})$$

La transformée de Fourier fenêtrée a une résolution temps–fréquence fixe. Cette résolution peut être modifiée par un changement d'échelle S sur la fenêtre g . C'est une représentation complète, stable et redondante du signal. Elle est donc inversible. La redondance se traduit par l'existence d'un noyau « reproduisant ». Pour plus d'informations, voir [98].

II.4.2.1 Choix de la fenêtre

Les propriétés de la transformée de Fourier fenêtrée sont déterminées par sa fenêtre g , ou plutôt par sa transformée de Fourier, dont on cherche à concentrer l'énergie autour de 0. La répartition de cette énergie est caractérisée par trois paramètres :

- La largeur de bande $\Delta\omega$, définie par :

$$\frac{|\hat{g}(\Delta\omega/2)|^2}{|\hat{g}(0)|^2} = \frac{1}{2}.$$

Si $\Delta\omega$ est petit, l'énergie de la fenêtre est bien concentrée autour de 0

- L'amplitude maximale A des lobes latéraux situés en $\omega = \pm\omega_0$, mesurée en décibels :

$$A = 10 \log_{10} \frac{|\hat{g}(\omega_0)|^2}{|\hat{g}(0)|^2}.$$

- L'exposant polynomial p donnant la décroissance asymptotique de la transformée de Fourier de g $|\hat{g}(\omega)|$ dans les hautes fréquences :

$$|\hat{g}(\omega)| = O(\omega^{-p-1}). \quad (\text{II.54})$$

qui résume ce qui se passe au delà du premier lobe latéral, Figure II.6 a).

Le Tableau II.2 donne les valeurs de ces trois paramètres pour plusieurs fenêtres à support, dont les graphes sont tracés dans la Figure II.6 b).

TABLEAU II.2

Paramètres fréquentiels de cinq fenêtres g à support dans $[-1/2, 1/2]$. Ces fenêtres sont normalisées pour que $g(0)=1$, mais $\|g\| \neq 1$.

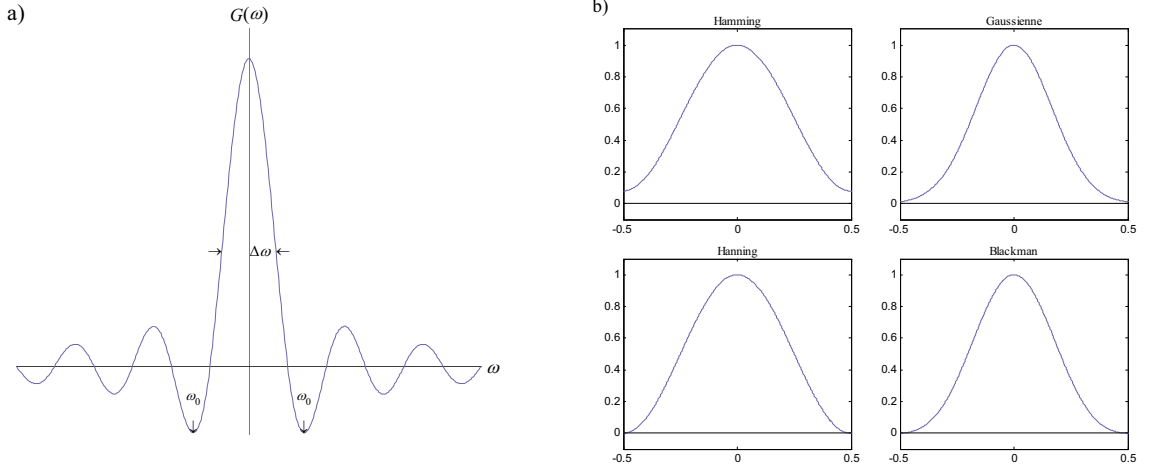
NOM	$g(t)$	$\Delta\omega$	A	P
Rectangle	1	0.89	-13 dB	0
Hamming	$0.54+0.46\cos(2\pi t)$	1.36	-43 dB	0
Gaussienne	$\exp(-18t^2)$	1.55	-55 dB	0
Hanning	$\cos^2(\pi t)$	1.44	-32 dB	2
Blackman	$0.42+0.5\cos(2\pi t)+0.08\cos(4\pi t)$	1.68	-58 dB	2

II.4.2.2 Transformée de Fourier à fenêtre rapide

La discrétisation et le calcul rapide de la transformée de Fourier à fenêtre relève des mêmes idées que la discrétisation de la FFT. La fenêtre $g[n]$ est un signal discret symétrique, de période N et de norme $\|g\|=1$. La transformée de Fourier fenêtrée discrète d'un signal x est définie par :

$$Sx[m, l] = \langle x, g_{m,l} \rangle = \sum_{n=0}^{N-1} x[n]g[n-m]\exp\left(\frac{i2\pi ln}{N}\right). \quad (\text{II.55})$$

Pour chaque $0 \leq m \leq N$, $Sx[m, l]$ se calcule pour $0 \leq l \leq N$ par transformée de Fourier discrète de $x[n]g[n-m]$. Ce calcul est effectué avec N FFT de taille N , ce qui demande un total de $O(N^2 \log_2 N)$ opérations.



II.4.3 Transformée en ondelettes

Afin d'analyser des composantes transitoires de durées différentes, il est nécessaire d'utiliser des atomes dont les supports temporels ont des tailles variables. Pour cela, la transformée en ondelettes décompose les signaux sur une famille d'atomes temps-fréquence en dilatant l'ondelette (II.48) par un facteur s et en translatant par u . La transformée en ondelettes est donc définie par :

$$Wx(u, s) = \langle x, \psi_{u,s} \rangle = \int_{-\infty}^{\infty} x(t) \frac{1}{\sqrt{s}} \psi^*\left(\frac{t-u}{s}\right) dt, \quad (\text{II.56})$$

où le symbole W indique la transformée en ondelettes (Wavelet transform).

Comme une transformée de Fourier à fenêtre, une transformée en ondelettes permet de mesurer l'évolution en temps de composantes fréquentielles. Pour cela, il faut une **ondelette analytique complexe**, afin de séparer la phase et l'amplitude. Par contre, les **ondelettes réelles** sont souvent utilisées pour la détection des transitions brutales d'un signal. Nous aborderons ce point dans la section suivante.

Pour analyser l'évolution temporelle des « fréquences » d'un signal, il est nécessaire d'utiliser une ondelette analytique, afin de séparer les informations de phase et d'amplitude. L'étalement de l'énergie de sa transformée en ondelettes analytique (II.56) correspond à une boîte de Heisenberg contrée en $(u, \eta/s)$, de largeur temporelle $s\sigma_t$ et de largeur fréquentielle σ_ω/s . La surface de ce rectangle reste égale à $\sigma_t\sigma_\omega$ pour toutes les échelles, mais sa résolution en temps et en fréquence dépend de s , Figure II.7.

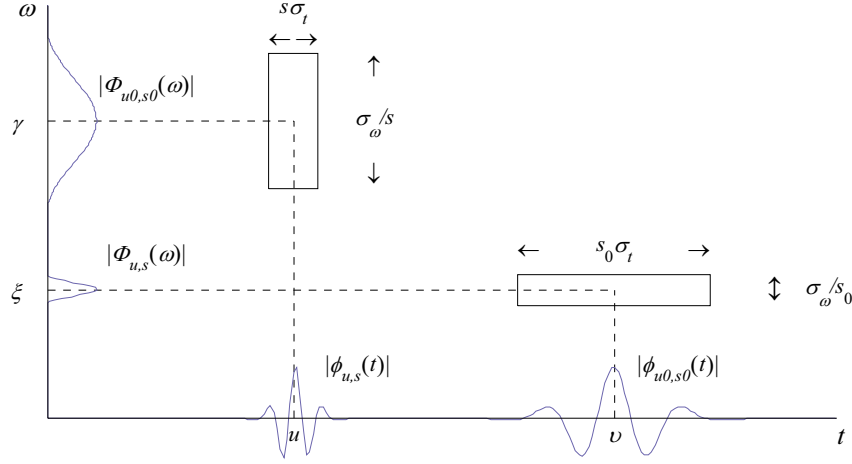


Figure II.7 : Boîte de Heisenberg de deux atomes de Fourier à fenêtre $g_{u,\xi}$ et $g_{v,\gamma}$.

Une transformée en ondelettes analytique définit une densité d'énergie en temps–fréquence $P_W x$, qui mesure l'énergie de x dans la boîte de Heisenberg de chaque ondelette $\psi_{u,s}$ centrée en $(u, \xi = \eta/s)$:

$$P_W x(u, \xi) = |Wx(u, s)|^2 = \left| Wx\left(u, \frac{\eta}{\xi}\right) \right|^2. \quad (\text{II.57})$$

Cette densité s'appelle *scalogramme*.

Dans le cas des ondelettes réelles, l'équation (II.56) mesure la variation de x dans un voisinage de u de taille proportionnelle à s . Une transformée en ondelettes réelles est complète et préserve l'énergie, tant que l'ondelette satisfait la *condition d'admissibilité* :

$$C_\psi = \int_0^{+\infty} \frac{|\hat{\psi}(\omega)|^2}{\omega} d\omega < +\infty. \quad (\text{II.58})$$

Pour que l'intégrale soit finie, il faut assurer que $\hat{\psi}(0) = 0$, ce qui explique pourquoi les ondelettes doivent être de moyenne nulle. Il s'agit d'une représentation complète, stable et redondante du signal ; en particulier, la transformée en ondelettes est inversible à gauche. La redondance se traduit par l'existence d'un noyau « reproduisant », [98].

Un exemple typique de ce type d'ondelettes est la double dérivation d'une gaussienne. Largement connue sous le nom de « chapeau mexicain » elles sont apparues en vision par ordinateur pour détecter des contours multi-échelles. L'ondelette en chapeau mexicain normalisée est :

$$\psi(\omega) = \frac{2}{\pi^{1/4} \sqrt{3}\sigma} \left(\frac{t^2}{\sigma^2} - 1 \right) \exp\left(\frac{-t^2}{2\sigma^2} \right). \quad (\text{II.59})$$

La Figure II.8 montre $-\psi$ et sa transformée de Fourier :

$$\hat{\psi}(\omega) = \frac{-\sqrt{8}\sigma^{5/2}\pi^{1/4}}{\sqrt{3}} \omega^2 \exp\left(\frac{-\sigma^2 \omega^2}{2} \right). \quad (\text{II.60})$$

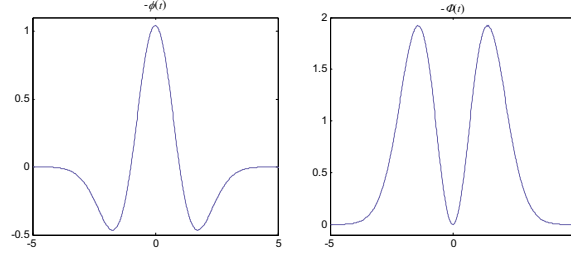


Figure II.8 : Ondelette « chapeau mexicain » pour $\sigma=1$ et sa transformée de Fourier.

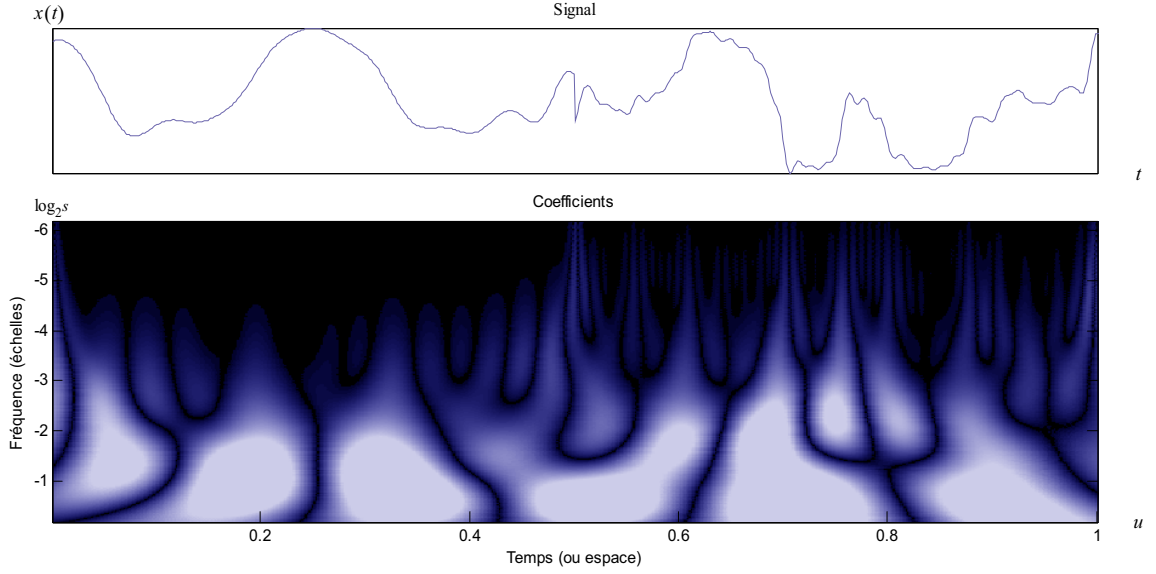


Figure II.9 : Transformée en ondelettes réelle calculée avec une ondelette en chapeau mexicain. Les axes verticaux et horizontaux représentent respectivement $\log_2 s$ et u . Les échelles les plus fines sont en haut. Les coefficients nuls correspondent à du noir, donc à des parties régulières.

II.4.3.1 Transformée en ondelettes discrète

La transformée en ondelettes discrète (discret wavelet transform, DWT) utilise une suite d'échelles discrètes a^j pour $j < 0$ avec $a=2^{1/\nu}$, ce qui fournit ν échelles intermédiaires pour chaque octave $[2^j, 2^{j+1}]$. Pour une ondelette $\psi(t)$ en temps continu, dont le support est inclus dans $[-K/2, K/2]$, pour $2 \leq a^j \leq NK^{-1}$, l'ondelette discrète dilatée par a^j est définie :

$$\psi_j[n] = \frac{1}{\sqrt{a^j}} \psi\left(\frac{n}{a^j}\right). \quad (\text{II.61})$$

Dans le cadre discret, on traite $x[n]$ et les ondelettes $\psi_j[n]$ comme des signaux de période N . La transformée en ondelettes discrète de x peut alors s'écrire comme une convolution circulaire avec :

$$Wx[n, a^j] = \sum_{m=0}^{N-1} x[m] \psi_j^*[m-n] = x \circledast \bar{\psi}_j[n]. \quad (\text{II.62})$$

La **transformée en ondelettes rapide** effectue le calcul de (II.62) par des convolutions circulaires, elles-mêmes calculées par FFT, nécessitant $O(M \log_2 N)$ opérations.

II.4.3.2 Transformée en ondelettes dyadiques

Pour construire une représentation par ondelettes invariante par translation, l'échelle s est discrétisée, mais pas le paramètre de translation u . L'échelle est échantillonnée sur une suite dyadique $\{2^j\}_{j \in \mathbb{Z}}$, pour simplifier l'implémentation. La transformée en ondelettes dyadique de $x \in \mathbf{L}^2(\mathbb{R})$ est définie par :

$$Wx(u, 2^j) = \int_{-\infty}^{+\infty} x(t) \frac{1}{\sqrt{2^j}} \psi\left(\frac{t-u}{2^j}\right) dt, \quad (\text{II.63})$$

avec l'ondelette dyadique :

$$\bar{\psi}_{2^j}(t) = \psi_{2^j}(t) = \frac{1}{\sqrt{2^j}} \psi\left(\frac{-t}{2^j}\right). \quad (\text{II.64})$$

Si l'axe des fréquences est totalement recouvert par les ondelettes dyadiques dilatées, alors cette transformée est complète et stable. La formule de reconstruction est :

$$x(t) = \sum_{j=-\infty}^{+\infty} \frac{1}{2^j} Wx(., 2^j) * \tilde{\psi}_{2^j}(t). \quad (\text{II.65})$$

Supposons que les fonctions d'échelle et les ondelettes ϕ , ψ , $\tilde{\phi}$ et $\tilde{\psi}$ soient construites avec les filtres h , g , $\tilde{h}$ et $\tilde{g}$. La **transformée en ondelettes dyadique rapide** est calculée par « l'algorithme à trous », [81]. Cet algorithme s'implémente avec un banc de filtres analogue à celui d'une transformée en ondelettes biorthogonales, rapide mais sans échantillonnage, [98].

Supposons que les échantillons $a_0[n]$ du signal d'entrée discret ne valent pas $x(t)$, mais une moyenne locale de x au voisinage de $t=n$. Alors, les échantillons s'écrivent comme des moyennes de $x(t)$ pondérées avec les noyaux d'échelle $\phi(t-n)$. Pour tout $j \geq 0$, on pose :

$$a_j[n] = \langle x(t), \phi_{2^j}(t-n) \rangle, \text{ avec } \phi_{2^j}(t) = \frac{1}{\sqrt{2^j}} \phi\left(\frac{t}{2^j}\right).$$

Les coefficients d'ondelettes sont calculées pour $j > 0$ sur la grille des entiers :

$$d_j[n] = Wx(n, 2^j) = \langle x(t), \psi_{2^j}(t-n) \rangle.$$

Etant donné tout filtre $r[n]$, on note $r_j[n]$ le filtre obtenu en insérant $2^j - 1$ zéros entre chaque coefficient de x (d'où le nom d'algorithme à trous), et on pose $\bar{r} = r[-n]$. L'algorithme à trous permet de calculer la transformée dyadique rapide de la manière suivante :

$$a_{j+1}[n] = a_j * \bar{h}_j[n], \quad d_{j+1}[n] = a_j * \bar{g}_j[n], \text{ et} \quad (\text{II.66})$$

$$a_j[n] = \frac{1}{2} (a_{j+1} * \tilde{h}_j[n] + d_{j+1} * \tilde{g}_j[n]). \quad (\text{II.67})$$

les filtres avec un tilde $\sim$ sont les filtres duaux du système biorthogonal, [98].

II.4.4 Analyse en ondelettes

Une transformée en ondelettes permet d'analyser les structures locales d'un signal avec un zoom qui réduit progressivement le paramètre d'échelle. Par exemple, les transitoires d'un électrocardiogramme ou d'un signal de radar portent des informations importantes. De même, les discontinuités de l'intensité d'une image indiquent la présence de contours dans une scène.

II.4.4.1 Régularité

Afin de caractériser les structures singulières, il faut quantifier précisément la régularité locale d'un signal $x(t)$. L'analyse de Fourier permet de caractériser la régularité globale d'une fonction, tandis que la transformée en ondelettes permet d'analyser la régularité ponctuelle d'une fonction, [98].

Les **exposants de Lipschitz** (où exposants de Hölder) fournissent des mesures de régularité uniforme sur des intervalles, mais également en un point v quelconque. Un signal est régulier (ponctuellement Lipschitz $\alpha > 0$ en v) s'il existe $K > 0$ en un polynôme p_v de degré $m = [\alpha]$ tels que :

$$\forall t \in \mathfrak{R}, \quad |x(t) - p_v(t)| \leq K |t - v|^\alpha. \quad (\text{II.68})$$

Ainsi, une fonction x est uniformément Lipschitz α sur $[a, b]$ si elle vérifie la condition (II.68) pour tout $v \in [a, b]$, avec une constante K indépendante de v . La régularité de x en v ou sur $[a, b]$ est le *sup* des α pour lesquels x est Lipschitz α .

La régularité Lipschitzienne uniforme de x sur $\mathfrak{R}$ est liée à la décroissance asymptotique de sa transformée de Fourier. Donc, une fonction x est bornée et uniformément Lipschitz α sur $\mathfrak{R}$ si elle satisfait la **condition de Fourier** (condition de régularité globale), définie par :

$$\int_{-\infty}^{+\infty} |\hat{x}(\omega)| (1 + |\omega|^\alpha) d\omega < +\infty. \quad (\text{II.69})$$

Pour mesurer la régularité locale d'un signal, l'important n'est pas d'utiliser une ondelette à support fréquentiel étroit, mais de préciser le nombre de moments nuls. Si l'ondelette a n moment nuls alors, sa transformée en ondelettes peut s'interpréter comme un opérateur différentiel multiéchelle d'ordre n (**condition sur la transformée en ondelettes**) :

$$\int_{-\infty}^{+\infty} t^k \psi(t) dt = 0, \text{ pour } 0 \leq k < n. \quad (\text{II.70})$$

Une ondelette à n moments est orthogonale aux polynômes de degré $n-1$. Pour qu'une ondelette à décroissance rapide ait ses n moments nuls, elle doit être la dérivée $n^{\text{ème}}$ d'une fonction à décroissance rapide θ :

$$\psi(t) = (-1)^n \frac{d^n \theta(t)}{dt^n}. \quad (\text{II.71})$$

Par conséquent, sa transformée en ondelettes devient :

$$Wx(u, s) = s^n \frac{d^n}{du^n} (x * \bar{\theta}_s)(u), \quad (\text{II.72})$$

avec $\bar{\theta}_s(t) = s^{-1/2} \theta(-t/s)$. De plus, ψ n'a pas plus de n moments nuls quand $\int_{-\infty}^{+\infty} \theta(t) dt \neq 0$. Dans le cas où x est une fonction un peu plus que n fois différentiable en un point v , on peut l'approximer par un polynôme de degré n . La transformée en ondelettes de ce polynôme est nulle ; autour de v , elle est donc de l'ordre de l'erreur entre le polynôme et la fonction. Si cette erreur peut être estimée uniformément sur un intervalle, on obtient un outil d'étude de la régularité sur l'intervalle.

II.4.4.2 Maxima de la transformée en ondelettes et détection de singularités

La régularité d'une fonction x en v dépend de la décroissance aux échelles fines de $|Wx(u, s)|$ au voisinage de v . La mesure de la décroissance dans le plan (u, s) n'est pas nécessaire ; elle peut être contrôlée par les valeurs de ses maxima locaux. On appelle **ligne de maxima** toute courbe connexe du plan espace-échelle (u, s) dont tous les points sont de modules maximaux.

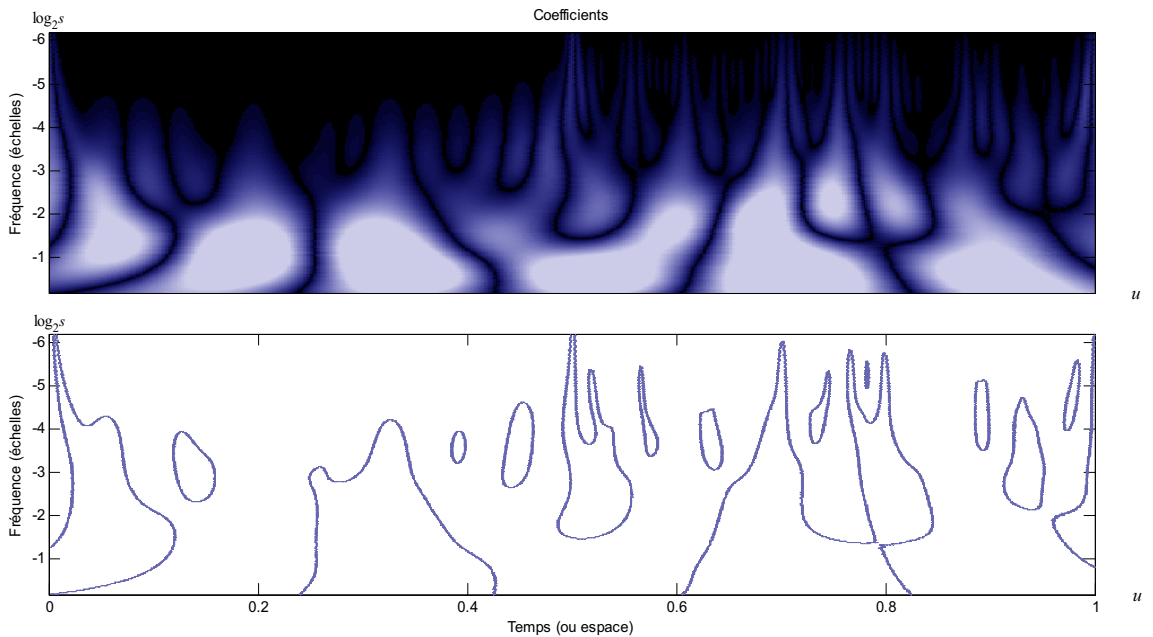


Figure II.10 : Modules maximaux de $Wx(u, s)$ de l'exemple de la Figure II.9.

La **détection de singularités** se fait en cherchant les abscisses où convergent les modules maximaux d'ondelettes aux fines échelles. Afin de mieux comprendre les propriétés de ces maxima, on écrit la transformée en ondelettes comme un opérateur différentiel multiéchelle. Si l'ondelette ψ a exactement n moment nuls et un support compact, alors θ a un support compact tel que :

$$\psi = (-1)^n \theta' \text{ avec } \int_{-\infty}^{+\infty} \theta(t) dt \neq 0. \quad (\text{II.73})$$

En [98], on a démontré que ces modules maximaux peuvent être ou ne pas être sur une même ligne de maxima, ce qui garantit que toutes les singularités soient détectées en suivant les modules maximaux d'ondelettes dans les fines échelles. La Figure II.10 montre les modules maximaux d'ondelettes du signal de l'exemple de la Figure II.9, trouvant toutes les singularités.

En général, pour des ondelettes (II.73), rien ne garantit qu'un module maximal situé en (u_0, s_0) appartienne à une ligne de **propagation des maxima** vers les fines échelles. Dans le cas où θ est une gaussienne, les modules maximaux de $Wx(u, s)$ appartiennent à des courbes connexes qui ne s'interrompent jamais quand l'échelle diminue.

Le taux de décroissance des maxima le long des courbes indique l'ordre des **singularités isolées** :

$$\log_2 |Wx(u, s)| \leq \log_2 A + \left(\alpha + \frac{1}{2} \right) \log_2 s. \quad (\text{II.74})$$

II.4.5 Bases d'ondelettes

On peut construire des ondelettes qui génèrent des bases orthonormées de $L^2(\mathbb{R})$ par translation et dilatation :

$$\left\{ \psi_{j,n}(t) = \frac{1}{\sqrt{2^j}} \psi \left(\frac{t - 2^j n}{2^j} \right) \right\}_{(j,n) \in \mathbb{Z}^2}. \quad (\text{II.75})$$

Les ondelettes orthogonales dilatées de 2^j reproduisent les variations d'un signal à la résolution 2^{-j} . La construction de ces bases est ainsi liée à l'approximation multirésolution des signaux. C'est précisément ce lien qui conduit à une surprenante équivalence entre les bases d'ondelettes et les filtres miroirs conjugués utilisés dans les bancs de filtres sous-échantillonnées.

Il y a différents types de familles d'ondelettes, dont les qualités changent selon plusieurs critères. Les critères principaux sont :

- Le support de ψ , $\hat{\psi}$ (ϕ et $\hat{\phi}$) : la vitesse de convergence vers zéro de ces fonctions ($\psi(t)$ ou $\psi(\omega)$) quand le temps t ou la fréquence ω tends vers l'infini, quantifiant les localisations en temps et en fréquence,
- La symétrie, très utile pour éviter le déphasage dans le traitement d'images,
- Le nombre de moments nuls pour ψ ou pour ϕ (s'il existe), utile pour des besoins de compression,
- La régularité, qui est utile pour obtenir des caractéristiques intéressantes, comme la continuité du signal reconstruit ou l'image, pour la fonction estimée en analyse de régression non linéaire.

Ces critères sont associés à deux propriétés qui permettent l'implémentation rapide et efficace d'un algorithme :

- L'existence d'une fonction d'échelle ϕ ,
- L'orthogonalité ou la biorthogonalité de l'analyse résultante.

Ils peuvent également être associés à ces propriétés moins importantes :

- L'existence d'une expression explicite,
- La facilité de tabulation,
- La familiarité d'utilisation.

Pour plus de renseignements sur les types d'ondelettes, voir [98].

II.5 Extraction de caractéristiques

Les fonctions aléatoires et les processus stochastiques font l'objet de nombreux ouvrages dans la littérature scientifique, [112]. Dans ce paragraphe, nous présentons la définition, de façon restrictive, d'un processus aléatoire (ou stochastique) comme une famille de fonctions réelles bidimensionnelles $x(t, \zeta)$ où ζ est une variable aléatoire. Suivant les phénomènes observés, la fonction x peut être monodimensionnelle. Il est possible, également, de classer de manière distincte les processus aléatoires en fonction de la nature de la variable aléatoire ζ : continue, discrète, monodimensionnelle, multidimensionnelle.

II.5.1 Analyse statistique

Considérons le cas général d'un fonction aléatoire $X(t)$, $X \in \mathcal{R}^n$. Pour chaque valeur de t_1 , notons X_1 l'amplitude des $x_i(t_1)$, $i=1, \dots, n$, dont le comportement peut être caractérisé par la fonction de répartition $f(x_1/t_1)$ et sa loi de densité de probabilité $p(x_1/t_1)$. Par extension, si l'on considère un ensemble de k instants $t=t_1, \dots, t_k$, on définit un variable aléatoire $\mathbf{x}$ de dimension k caractérisé par sa loi de densité de probabilité conditionnelle $p(x_1, \dots, x_k | t_1, \dots, t_k)$. La connaissance de cette loi de densité de probabilité définit les statistiques d'ordre k du processus aléatoire.

En théorie, la description d'un processus aléatoire continue nécessite la connaissance de cette statistique pour un ensemble infini de valeurs de t , mais pour des considérations pratiques, nous nous limitons à la détermination des statistiques d'ordre 1 et 2.

II.5.1.1 Statistiques du premier ordre

La fonction de répartition de X_1 est définie par $f(x_1/t_1) = \Pr(X_1 \leq x_1)$, x_1 étant une valeur donnée et sa loi de densité de probabilité $p(x_1/t_1)$ est la dérivée de la fonction de répartition :

$$p(x_1 | t_1) = \frac{df(x_1, t_1)}{dx_1}.$$

La loi de densité de probabilité conditionnelle $p(x_1/t_1)$ peut être approximée par :

$$p(x_1, t_1) = \lim_{\Delta x \rightarrow 0} \frac{\Pr(x_1 \leq X_1 < x_1 + \Delta x_1 | t = t_1)}{\Delta x_1}.$$

De façon similaire, on peut calculer les différents moments et la variance par les expressions suivantes :

$$E[X_1^n] = \int_{-\infty}^{+\infty} x_1^n p(x_1 | t_1) dx_1, \quad (\text{II.76})$$

En particulière la **moyenne** $\mu(x_{t_1})$ est définie par :

$$\bar{X}_1 = \mu(x_{t_1}) = \int_{-\infty}^{+\infty} x_1 p(x_1 | t_1) dx_1, \quad (\text{II.77})$$

En pratique, cette moyenne est estimée par :

$$\bar{X}_1 = \mu(x_1) = \frac{1}{N} \sum_{k=0}^{N-1} x_1(t_1 - k), \quad (\text{II.78})$$

Moments

L'application immédiate de l'espérance de x_1^m de la variable aléatoire X_1 est le calcul le moments. Il existe deux types de moments : par rapport à l'origine et par rapport à la moyenne, [115].

L'utilisation directe de (II.76) donne les moments par rapport à l'origine. Les moments par rapport à la moyenne, également connus comme **moments centraux**, sont définis comme :

$$\mu_m = E[(X_1 - \bar{X}_1)^m] = \int_{-\infty}^{+\infty} (x_1 - \bar{X}_1)^m p(x_1 | t_1) dx_1. \quad (\text{II.79})$$

De cette manière, $\mu_0 = 1$ et $\mu_1 = 0$.

Le deuxième moment central μ_2 est largement connu sous le nom de **variance**, notée $\sigma_{X_1}^2$. Elle est donnée par :

$$\sigma_X^2 = \mu_2 = E[(X_1 - \bar{X}_1)^2] = \int_{-\infty}^{+\infty} (x_1 - \bar{X}_1)^2 p(x_1 | t_1) dx_1. \quad (\text{II.80})$$

La racine carrée σ_{X_1} définie l'**écart type** de X_1 . La variance des signaux discrets, finis est calculée par :

$$\sigma_{X_1}^2 = \frac{1}{N} \sum_{k=0}^{N-1} (x^2(t_1 - k) - \mu^2), \quad (\text{II.81})$$

Le troisième moment $\mu_3 = E[(X_1 - \bar{X}_1)^3]$ est une mesure d'asymétrie de $p(x_1 | t_1)$, par rapport à $\bar{X}_1$. Il est connu comme le **skewness** de la fonction de densité. Si la densité est symétrique par rapport à $\bar{X}_1$, la valeur du « skewness » associé est zéro. Ce moment est calculé par :

$$\mu_3 = \frac{Ns_{t_1}^3}{(N-1)(N-2)\sigma_{X_1}^3}, \text{ avec } s_{t_1}^j = \sum_{k=0}^{N-1} (x(t_1 - k) - \mu)^j \quad (\text{II.82})$$

Une variation du quatrième moment μ_4 est le facteur **kurtosis** $= E[(X_1 - \bar{X}_1)^4] / \sigma_{X_1}^4$. Le facteur kurtosis est une valeur statistique vérifiant le caractère « gaussien » d'un signal. Ce moment est calculé par :

$$\mu_4 = \frac{N(N+1)s_{t_1}^4 - 3s_{t_1}^2 s_{t_1}^2 (N-1)}{(N-1)(N-2)(N-3)\sigma_{X_1}^4}. \quad (\text{II.83})$$

II.5.1.2 Statistiques du deuxième ordre

Considérant deux instants différents t_1 et t_2 auxquels correspondent les deux variables aléatoires X_1 et X_2 . La fonction de répartition conjointe est définie par $p(x_1, x_2 | t_1, t_2) = \Pr(X_1 \leq x_1, X_2 \leq x_2)$. La loi de densité de probabilité conjointe du couple de variables aléatoires X_1 et X_2 est donnée par :

$$p(x_1, x_2 | t_1, t_2) = \frac{\partial^2 f(x_1, x_2)}{\partial x_1 \partial x_2}.$$

La **fonction d'autocorrélation** du processus aléatoire est définie par :

$$R_{xx}(t_1, t_2) = E[X_1 X_2] = [X_1(t_1) X_2(t_2)] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} x_1 x_2 p(x_1, x_2 | t_1, t_2) dx_1 dx_2, \quad (\text{II.84})$$

La fonction d'autocovariance est définie par l'expression :

$$C_{xx}(t_1, t_2) = E[(X_1(t_1) - \mu(x_1))(X_2(t_2) - \mu(x_2))] = R_{xx}(t_1, t_2) - \mu(x_1)\mu(x_2), \quad (\text{II.85})$$

Pour une translation dans le temps τ , la loi de densité de probabilité $p(x_1, x_2 | t_1, t_2)$ ne dépend que de $\tau = t_1 - t_2$, d'où les propriétés de la fonction de corrélation et de la fonction de covariance :

$$\begin{cases} R_{xx}(t_1, t_2) = R_{xx}(\tau) \\ C_{xx}(t_1, t_2) = R_{xx}(\tau) - \mu_x^2 \end{cases}, \quad (\text{II.86})$$

Si $\mu_x = 0$, la fonction de corrélation et la fonction de covariance coïncident. Un signal est **stationnaire** au sens large si sa moyenne est indépendante de l'origine des temps et si sa fonction d'autocorrélation ne dépend que de $\tau = t_1 - t_2$.

II.5.2 Analyse spectrale

La connaissance de la densité spectrale d'un signal est très importante pour identifier la bande de fréquence où se situe l'information spectrale. La densité spectrale d'un signal $x(t)$ et sa fonction de corrélation $R_{xx}(\tau)$ sont liées par la transformée de Fourier :

$$\begin{cases} S_{xx}(\omega) = F(R_{xx}(\tau)) = \int_{-\infty}^{+\infty} R_{xx}(\tau) e^{-i2\pi\omega\tau} d\tau \\ R_{xx}(\tau) = F^{-1}(S_{xx}(\omega)) = \int_{-\infty}^{+\infty} S_{xx}(\omega) e^{+i2\pi\omega\tau} d\omega \end{cases}. \quad (\text{II.87})$$

Pour un signal réel, la fonction d'autocorrélation est une fonction paire, ce qui entraîne que la densité spectrale de puissance soit une fonction réelle paire, [172].

L'équation (II.87) fournit la densité spectrale de puissance $S_{XX}(\omega)$ dont l'intégration sur l'intervalle $[\omega_0, \omega_1]$ permet d'évaluer la **puissance spectrale** du signal sur ce même intervalle de fréquences :

$$P(\omega_0, \omega_1) = \int_{\omega_0}^{\omega_1} S_{XX}(\omega) d\omega. \quad (\text{II.88})$$

A partir de sa puissance spectrale, nous pouvons calculer les caractéristiques de la bande en fréquences du signal. Ainsi, la **fréquence centrale** :

$$\omega_c = \frac{\int_{\omega_0}^{\omega_1} \omega P(\omega_0, \omega_1) d\omega}{\int_{\omega_0}^{\omega_1} P(\omega_0, \omega_1) d\omega}, \quad (\text{II.89})$$

représente le centre de gravité de la bande en fréquences. La **largeur de la bande en fréquences** :

$$\omega_{bw} = \frac{\int_{\omega_0}^{\omega_1} (\omega - \omega_c)^2 P(\omega_0, \omega_1) d\omega}{\int_{\omega_0}^{\omega_1} P(\omega_0, \omega_1) d\omega}, \quad (\text{II.90})$$

représente la moitié de la largeur de bande de le spectre en fréquences. En autres termes, il s'agit de la variance par rapport à ω_c . La **racine carrée de la fréquence moyenne** :

$$\omega_r^2 = \frac{\int_{\omega_0}^{\omega_1} \omega^2 P(\omega_0, \omega_1) d\omega}{\int_{\omega_0}^{\omega_1} P(\omega_0, \omega_1) d\omega}, \quad (\text{II.91})$$

représente le deuxième moment du spectre de puissance du signal. Il est important de noter que la relation : $\omega_r^2 = \omega_c^2 + \omega_{bw}^2$ est satisfaite.

II.5.3 Fenêtres d'analyse

Les fenêtres glissantes fournissent une base d'analyse de signaux sur un intervalle de temps spécifique. Ces fenêtres sont définies comme des filtres discrets L , utilisant des fenêtres rectangulaires à largeur N :

$$h = \mathbf{1}_{[n-N, n]}, \quad (\text{II.92})$$

pour l'instant n , car nous travaillons avec des systèmes réels, donc causaux.

De cette manière, si l'on veut calculer la moyenne sur les N échantillons passés à l'instant n , du signal discret x , la somme (II.25) devient :

$$y[n] = Lx[n] = \frac{1}{N} \sum_{p=n-N+1}^n x[p] h[n-p] = x * h[n] \quad (\text{II.93})$$

La largeur de la fenêtre peut être définie par l'utilisateur, par la transformée de Fourier fenêtrée ou par la transformée en ondelettes, (plus précisément par la détection de singularités).

II.5 Conclusion

Dans ce chapitre, nous avons présenté les différentes méthodes d'analyse de signaux pour leur prétraitement et l'extraction d'informations pertinentes pour leur étude. Nous nous sommes limité à l'étude de signaux monodimensionnels pour l'analyse de leur dynamiques en temps et en fréquence et de leurs caractéristiques statistiques et fréquentielles.

Nous avons brièvement abordé l'analyse des signaux continus, pour les systèmes le plus simples, les systèmes LTI, afin d'introduire la transformée de Fourier et transposer les résultats au cadre des signaux

discrets. Ainsi, nous avons expliqué l'importance du calcul des convolutions ainsi que leur implémentation à complexité moindre à l'aide de la transformée de Fourier rapide.

Une fois les concepts de base expliqués, nous avons abordé l'analyse temporelle et fréquentielle des signaux, afin d'examiner leurs composantes transitoires de durées différentes à des différentes fréquences. Pour cela nous avons introduit et discuté la transformée de Fourier à fenêtre et la transformée en ondelettes.

Chapitre III

Techniques d'apprentissage et d'intelligence artificielle pour la modélisation de systèmes complexes

Sommaire

III.1	Introduction.....	51
III.2	Modélisation à partir des données.....	51
III.2.1	L'apprentissage statistique.....	51
III.2.1.1	Les machines d'apprentissage	52
III.2.1.2	Les fonctions de perte et la minimisation du risque	52
III.2.1.3	Les trois problèmes principaux d'apprentissage.....	52
III.2.1.4	Les principes d'induction.....	52
III.2.2	Les machines d'apprentissage classiques : l'approche neuronale.....	55
III.2.2.1	Les machines linéaires.....	55
III.2.2.2	Les machines non-linéaires	59
III.2.2.3	L'estimation de la densité.....	63
III.2.2.4	Remarques	66
III.2.3	Les machines à vecteurs de support SVM	67
III.2.3.1	Règles de décision non linéaires.....	67
III.2.3.2	SVM pour la reconnaissance de formes	69
III.2.3.3	SVM pour la régression	71
III.2.3.4	SVM pour l'estimation de la densité	73
III.2.3.5	Remarques	74
III.2.4	Le mécanisme d'apprentissage des SVM.....	74
III.2.4.1	Caractéristiques du problème d'optimisation quadratique SVM.....	74
III.2.4.2	Les algorithmes d'optimisation des problèmes quadratiques	75
III.2.4.3	Implémentation des SVM	78
III.2.4.4	Application des SVM sur des bases de données.....	80
III.3	Modélisation à partir de la connaissance	81
III.3.1	Généralités sur l'intelligence artificielle	81
III.3.1.1	Agents intelligents	82
III.3.1.2	Systèmes Experts.....	82
III.3.2	Les systèmes d'inférence flous	83
III.3.3.1	Eléments d'un système d'inférence flou.....	83
III.3.3.2	Généralités sur la logique floue	84
III.3.3.3	Le mécanisme de fuzzification, d'inférence et de défuzzification.....	86
III.3.3.4	Types de modèles flous.....	89
III.3.3.5	Génération de règles à partir de l'expertise	94
III.4	Modélisation à partir de la connaissance et les données.....	94
III.4.1	Modèles flous conventionnels avec apprentissage.....	94
III.4.1.1	Choix du nombre d'ensembles flous	95

III.4.1.2	Paramétrisation des règles	95
III.4.2	Modèles neuro-flous	96
III.4.2.1	Les algorithmes de groupage par données (clustering).....	96
III.4.2.3	Le groupage SVM.....	98
III.4.2.4	Initialisation des algorithmes de groupage	99
III.4.3	Identification des modèles TS pour des systèmes MIMO.....	100
III.4.3.1	Structure d'un modèle flou TS pour des systèmes MIMO	100
III.4.3.2	Méthode d'identification.....	101
III.4.3.3	Exemple d'identification.....	102
III.5	Conclusion	104

III.1 Introduction

Les techniques d'apprentissage et d'intelligence artificielle sont des disciplines définies comme la conception et la réalisation de systèmes informatiques résolvant des problèmes pour lesquels on ne dispose pas de solutions algorithmiques simples. Dans ce chapitre, nous présentons les outils méthodologiques utilisés dans cette recherche.

Dans la première partie, nous présentons la modélisation de systèmes à partir de données. Nous commençons avec une introduction à la théorie de l'apprentissage pour, ensuite, présenter des machines d'apprentissage classiques, notamment dans un cadre neuronal. Ceci nous permet d'introduire les machines à vecteurs de support, machines directement issues de la théorie de l'apprentissage statistique, pour lesquelles nous faisons une analyse plus complète, particulièrement du problème d'optimisation quadratique au cœur de leur processus d'apprentissage. Nous présentons des résultats sur quelques bases de données existantes dans la littérature.

La deuxième partie présente les systèmes d'inférence flous, en tant que branche des systèmes experts et de l'intelligence artificielle, pour modéliser la connaissance humaine (expertise) sous forme des règles. De manière générale, nous présentons les caractéristiques de base de la logique floue, les éléments d'un système d'inférence floue et les types de modèles existants.

La troisième partie aborde les aspects hybrides de la modélisation à partir des données et de l'expertise. Nous présentons les modèles flous conventionnels dotés d'un mécanisme d'optimisation aussi bien pour leur architecture que pour leurs paramètres. Ensuite, nous exposons les généralités sur les modèles neuro-flous, pour introduire une méthode hybride SVM-floue. Nous finissons avec un exemple d'identification de modèles flous Takagi–Sugeno.

III.2 Modélisation à partir des données

III.2.1 L'apprentissage statistique

L'apprentissage statistique [155] désigne l'estimation d'une fonction :

$$y=f(\mathbf{x}), \quad (\text{III.1})$$

à partir de ℓ observations indépendantes et identiquement distribuées (i.i.d.) tirées conformément à $P(\mathbf{x}, y)=P(\mathbf{x})P(y|\mathbf{x})$:

$$(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_\ell, y_\ell), \quad (\text{III.2})$$

couramment appelé *ensemble d'apprentissage*. Notons que $\mathbf{x} \in \mathcal{R}^n$ et, selon le problème que l'on traite, $y \in \mathcal{N}$, pour les problèmes de reconnaissance de formes, ou $y \in \mathcal{R}$, dans le cas des problèmes de régression et estimation de la densité.

Le mécanisme d'apprentissage doit trouver dans l'ensemble des fonctions $f(\mathbf{x}, \boldsymbol{\alpha})$, $\boldsymbol{\alpha} \in \mathbf{A}$, celle qui approxime au mieux la fonction inconnue, où $\boldsymbol{\alpha}$ peut être n'importe quel ensemble abstrait de paramètres, choisi *a priori*. Ce paradigme restreint le type de fonctions que l'on peut considérer :

1. On suppose une distribution probabiliste (statistique) des données : une valeur est tirée selon une distribution fixe mais inconnue, chaque point pouvant être contaminé par un bruit également de distribution fixe et inconnue. Ceci empêche de considérer les problèmes d'apprentissage avec cibles mobiles comme dans l'apprentissage renforcé, ou des données dépendantes, comme dans l'analyse des séries temporelles.
2. Les exemples $\mathbf{x}_1, \dots, \mathbf{x}_\ell$, sont des vecteurs de $\mathcal{R}^n$, ce qui permet de donner une interprétation géométrique à certaines méthodes.

III.2.1.1 Les machines d'apprentissage

Afin de construire une machine d'apprentissage, on a besoin de quatre éléments : un domaine (qui correspond au problème d'apprentissage avec une fonction de perte associée), un principe d'induction, un ensemble de fonctions de décision et un algorithme qui implémente les trois autres éléments.

III.2.1.2 Les fonctions de perte et la minimisation du risque

Afin de choisir la meilleure approximation du comportement du système à modéliser, une approche couramment utilisée consiste à mesurer la perte ou la divergence $L(y, f(\mathbf{x}, \alpha))$ entre la réponse y du système à une entrée $\mathbf{x}$ et la réponse $f(\mathbf{x}, \alpha)$ donnée par la machine d'apprentissage. Considérons la valeur prévue de perte, donnée par la fonction de risque :

$$R(\alpha) = \int L(y, f(\mathbf{x}, \alpha)) dP(\mathbf{x}, y). \quad (\text{III.3})$$

Le but de l'apprentissage est alors de trouver la fonction $f(\mathbf{x}, \alpha^*)$ qui minimise la fonction de risque $R(\alpha)$ dans le cas où la fonction de distribution de probabilité jointe $P(\mathbf{x}, y)$ est inconnue et la seule information disponible est l'ensemble d'apprentissage. Suivant cette définition, on peut définir différents problèmes d'apprentissage (domaines) avec les fonctions de perte associées. Dans cette thèse, nous considérons les trois problèmes principaux d'apprentissage, définis dans la section suivante.

III.2.1.3 Les trois problèmes principaux d'apprentissage

Dans [151], Vapnik définit trois principaux problèmes d'apprentissage: la reconnaissance de formes, l'estimation de la régression et l'estimation de la densité.

La reconnaissance de formes

Dans ce type de problèmes, on considère une sortie binaire $y \in \{-1, 1\}$ et $f(\mathbf{x}, \alpha)$, $\alpha \in \mathbf{A}$, un ensemble de fonctions indicateur (fonctions binaires). Considérons également la fonction de perte :

$$L(y, f(\mathbf{x}, \alpha)) = \begin{cases} 0 & \text{Si } y = f(\mathbf{x}, \alpha) \\ 1 & \text{Si } y \neq f(\mathbf{x}, \alpha) \end{cases} \quad (\text{III.4})$$

Pour cette fonction de perte, la fonction de risque (III.3) détermine la probabilité d'erreur de classification.

L'estimation de la régression

La réponse y prend, dans ce cas, une valeur réelle, et $f(\mathbf{x}, \alpha)$, $\alpha \in \mathbf{A}$, est un ensemble de fonctions réelles appelés *fonctions de régression*. La fonction de perte suivante est retenue :

$$L(y, f(\mathbf{x}, \alpha)) = (y - f(\mathbf{x}, \alpha))^2 \quad (\text{III.5})$$

Le problème est donc de trouver l'estimation de régression qui minimise la fonction de risque (III.3) avec la fonction de perte (III.5).

L'estimation de la densité

Finalement, considérons le problème d'estimation de densité $p(\mathbf{x}, \alpha)$, $\alpha \in \mathbf{A}$. La probabilité inconnue peut être trouvée en minimisant la fonction de risque :

$$R(\alpha) = \int L(y, p) dF(\mathbf{x}) \quad (\text{III.6})$$

avec la fonction de perte suivante :

$$L(p(\mathbf{x}, \alpha)) = -\log p(\mathbf{x}, \alpha). \quad (\text{III.7})$$

III.2.1.4 Les principes d'induction

Un principe d'induction permet de construire une règle de décision qui peut classer chaque point dans un espace donné à partir d'un nombre fini d'exemples (ensemble d'apprentissage). Nous allons considérer

l'un des principes d'induction les plus simples : le principe de *Minimisation du Risque Empirique* (Empirical Risk Minimisation, ERM). Nous décrirons, ensuite, le principe principal d'induction de la théorie de l'apprentissage statistique, le principe de *Minimisation Structurale du Risque* (Structural Risk Minimisation, SRM) et présenterons les similarités entre le SRM et la théorie de la Régularisation.

Le principe de minimisation du risque empirique (ERM)

Afin de minimiser la fonction de risque (III.3), il faut élire la fonction qui fournit la déviation minimale, dans le sens de notre fonction de perte, à partir d'une fonction réelle à travers tout l'espace de la fonction (pour chaque point $\mathbf{x}$). En réalité, la fonction de distribution jointe $F(\mathbf{x}, y)$ est inconnue et nous n'avons pas la valeur de y pour chaque point $\mathbf{x}$ dans l'espace de la fonction, mais seulement les paires de l'ensemble d'apprentissage (III.2).

Dans la pratique, nous pouvons approximer la fonction (III.3) en considérant la bien nommée *fonction de risque empirique* :

$$R_{emp}(\boldsymbol{\alpha}) = \frac{1}{\ell} \sum_{i=1}^{\ell} L(y_i, f(\mathbf{x}_i, \boldsymbol{\alpha})). \quad (\text{III.8})$$

Nous pouvons remarquer qu'elle ne fait pas intervenir de distribution de probabilité et que le calcul de R_{emp} est réalisable grâce au choix des paramètres $\boldsymbol{\alpha}$. Néanmoins, le minimum de (III.8) n'approxime pas nécessairement le minimum de (III.3) quand ℓ est petit. Vapnik, [151], considère que ℓ est petit si la relation ℓ/h (quotient du nombre d'échantillons d'apprentissage et de la dimension VC –définie ci-dessous– des fonctions d'une machine d'apprentissage) est petite, c'est à dire $\ell/h \leq 20$.

Dans le cas de la reconnaissance de formes, la borne suivante est valable pour le risque attendu avec une probabilité $1-\eta$, $0 \leq \eta \leq 1$:

$$R(\boldsymbol{\alpha}) \leq R_{emp}(\boldsymbol{\alpha}) + \sqrt{\frac{h(\log(\frac{2\ell}{h}) + 1) - \log(\frac{\eta}{4})}{\ell}}, \quad (\text{III.9})$$

où h est la dimension de Vapnik-Chernovenkis, VC, de l'ensemble de fonctions de décision paramétrées par $\boldsymbol{\alpha}$. La dimension VC d'un ensemble de fonctions est le nombre maximal de points qui peuvent être séparés sous toutes les formes possibles par cet ensemble [151]. Donc, si on connaît h , en choisissant η suffisamment petit, on peut utiliser cette borne pour calculer le meilleur ensemble $\boldsymbol{\alpha}$, trouvant ainsi la meilleure fonction de l'ensemble des fonctions de décision.

Le principe ERM, qui minimise le risque empirique (premier terme), donne seulement une valeur faible du risque attendu quand le ratio ℓ/h est grand. Si la dimension VC de l'ensemble de fonctions est grande, l'*intervalle de confiance* (deuxième terme) sera grand. Afin de minimiser les deux termes, la dimension VC doit être une variable contrôlable et n'est donc pas choisi *a priori*.

La capacité de généralisation

L'idée principale de l'apprentissage est d'adopter une fonction à partir d'un ensemble de règles de décision possibles (choisies *a priori*) qui s'adapte le mieux aux données. Cependant, puisque la quantité de données est souvent un petit échantillon, la fonction choisie peut décrire certaines dépendances entre les données.

La capacité de généralisation est contrôlée par le choix de l'ensemble de fonctions $f(\mathbf{x}, \boldsymbol{\alpha})$, $\boldsymbol{\alpha} \in \mathbf{A}$, le plus pertinent. La capacité de cet ensemble de fonctions (le nombre de séparations possibles des données par cet ensemble est une évaluation de la dimension VC) détermine le risque empirique atteint.

Un ensemble de fonctions, pouvant effectuer plusieurs séparations possibles avec une faible valeur du risque empirique, aura une capacité de généralisation réduite : c'est le phénomène de *sur-apprentissage* ou fausse structure. Au contraire, le choix d'un ensemble de fonctions avec une meilleure capacité de généralisation induira une capacité limitée pour décrire les dépendances entre les données : c'est le phénomène de *sous-apprentissage*. Pour illustrer ces situations, prenons le cas extrême pour $f(\mathbf{x}, \boldsymbol{\alpha})$, d'un côté, un ensemble de fonctions de décision linéaires (hyperplans) et, de l'autre côté, un ensemble de

fonctions sinus. Le premier pourra seulement décrire des dépendances linéaires, tandis que le deuxième pourra décrire n'importe quelle dépendance avec une fonction sinus haute fréquence sans aucune capacité de généralisation, Figure III.1 (a).

Typiquement, la capacité de généralisation peut se caractériser par une hyper-ellipse qui est une fonction de la capacité de l'ensemble de fonctions de décision, [163] et qui peut être contrôlée par un bon choix de la dimension VC ou par une transformation de l'ensemble des fonctions. Dans la section suivante, nous décrivons le principe inductif de minimisation structurelle du risque (SRM) qui tente de contrôler le risque empirique et la capacité de l'ensemble des fonctions à obtenir.

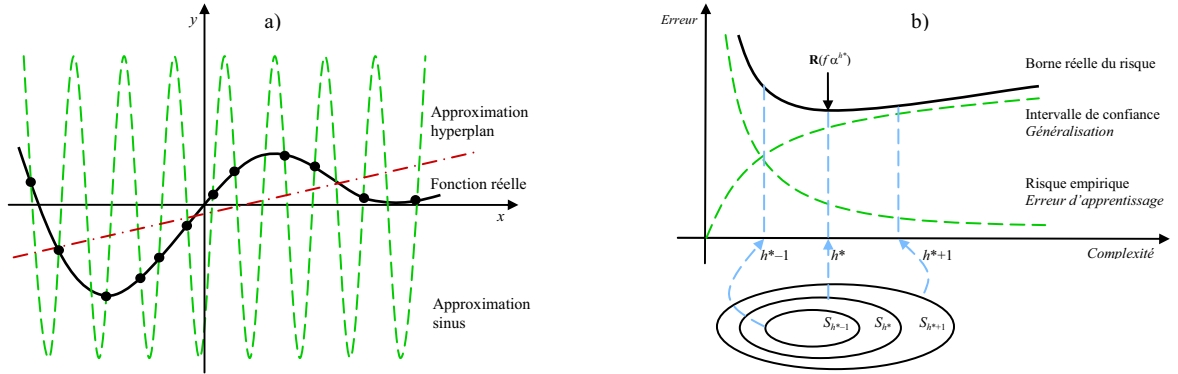


Figure III.1 : a) Estimation d'une fonction par un hyperplan, décrivant seulement les dépendances linéaires, et par une courbe sinus de haute fréquence, approximant la fonction en ses échantillons mais sans la décrire réellement. b) La borne réelle du risque est un compromis entre le risque empirique (erreur d'apprentissage) et l'intervalle de confiance (capacité de l'ensemble de fonctions S_k utilisées)

Le principe inductif de Minimisation Structurelle du Risque (SRM)

Le SRM défini par l'équation (III.9), [151], a pour objectif de minimiser le risque empirique et de maximiser l'intervalle de confiance simultanément. Pour cela, on définit la structure hiérarchisée des sous ensembles de fonctions de décision $f(\mathbf{x}, \alpha)$, $\alpha \in \mathbf{A}$, telle que :

$$S_1 \subset S_2 \subset \dots \subset S_n \dots$$

Leurs dimensions VC h_k correspondantes sont finies et satisfont:

$$h_1 \leq h_2 \leq \dots \leq h_n \dots$$

Le but est de choisir l'élément S_k le plus approprié qui minimise la borne (III.9), dans le cas de la reconnaissance de formes, Figure III.1 (b).

Théorie de la régularisation

Dans les années 1960, Tikhonov, [147], a proposé la théorie de la régularisation afin de résoudre des problèmes avec quelques propriétés, définies *a priori*, de la fonction désirée. Le principe des SRM a une forte analogie avec la théorie de régularisation. Au lieu de minimiser la fonction de perte (fonction de risque empirique) :

$$R^*(\alpha) = \frac{1}{\ell} \sum_{i=1}^{\ell} L(y_i, f(\mathbf{x}_i, \alpha)) \quad (\text{III.10})$$

on minimise, à la place, la fonction de régularisation :

$$R^*(\alpha) = \frac{1}{\ell} \sum_{i=1}^{\ell} L(y_i, f(\mathbf{x}_i, \alpha)) + \lambda Q(\alpha). \quad (\text{III.11})$$

où λ est une constante choisie de façon adéquate. Typiquement, $Q(\alpha)$ est une mesure de relaxation de $f(\mathbf{x}, \alpha)$. Si $Q(\alpha)$ est une mesure de généralisation de capacité quelconque, alors la méthode de régularisation est analogue au principe SRM. La constante λ définit le compromis entre le relâchement de l'approximation (capacité de généraliser ou intervalle de confiance) et la précision de l'approximation (risque empirique ou erreur d'apprentissage).

III.2.2 Les machines d'apprentissage classiques : l'approche neuronale

III.2.2.1 Les machines linéaires

Le perceptron monocouche

Le perceptron monocouche (single layer perceptron, SLP) est le modèle mathématique simplifié d'une unité élémentaire des systèmes nerveux des organismes vivants: le neurone, Figure III.2. Le cerveau humain est composé de milliards de ces unités interconnectées entre elles. La complexité des interconnexions entre les neurones ressemble à des réseaux. D'une manière plus modeste, les unités de traitement d'information (les neurones artificiels par exemple) ont des structures d'interconnexion souvent organisées en couches. Les unités d'une couche sont reliées à celles de la couche adjacente. C'est ainsi que le terme réseaux de neurones artificiels (artificial neural networks, ANN) est apparu.

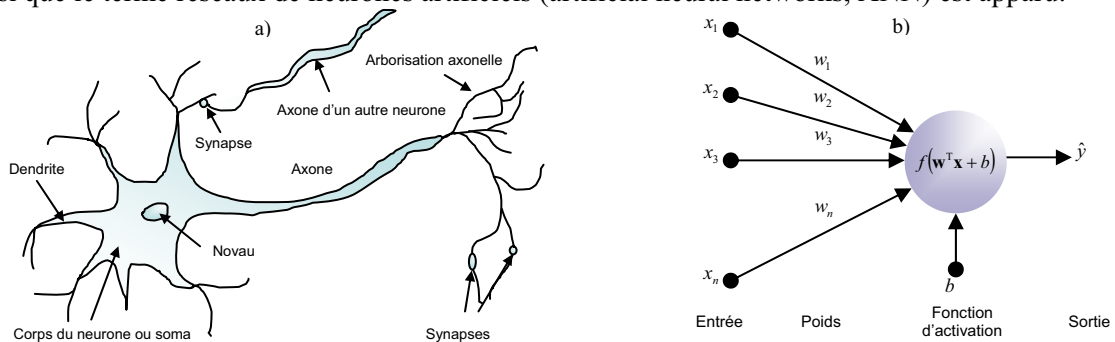


Figure III.2 : a) Le neurone biologique et b) le perceptron.

Le perceptron considère l'ensemble des fonctions linéaires :

$$\hat{y} = \mathbf{w}^T \mathbf{x} + b. \quad (\text{III.12})$$

Dans le cas de la classification, cet hyperplan découpe l'espace en deux régions, de telle sorte que :

$$\hat{y} = \begin{cases} 1 & \text{si } \mathbf{w}^T \mathbf{x} + b > 0 \\ -1 & \text{si } \mathbf{w}^T \mathbf{x} + b < 0, \end{cases} \quad (\text{III.13})$$

ce qui peut être représenté par la fonction signe. Ainsi, l'ensemble de tous les $\mathbf{x}$ qui satisfont l'égalité $\mathbf{w}^T \mathbf{x} + b = 0$ définit l'hyperplan séparateur ou la surface linéaire recherchée dans le cas de la régression.

La généralisation du perceptron revient à utiliser des fonctions, appelées *fonctions d'activation*, pour définir des sorties réelles ou pour introduire des non-linéarités (approche abordée dans la section suivante) :

$$\hat{y} = \phi(\mathbf{w}^T \mathbf{x} + b) \quad (\text{III.14})$$

En théorie, on peut utiliser n'importe quelle fonction, mais les fonctions continues et bornées ont des propriétés plus intéressantes. Le Tableau III.1 montre les fonctions d'activation les plus utilisées et leur intervalle de définition. La Figure III.3 illustre ces fonctions appliquant (III.14) avec des entrées $\mathbf{x} \in \mathbb{R}^2$.

TABLEAU III.1

Les fonctions d'activation les plus utilisées et leur intervalle de définition.

	FUNCTION D'ACTIVATION	INTERVALLE	
Identité	$\phi(\xi) = \xi$	$(-\infty, \infty)$	(III.15)
Logistique (Sigmoidale)	$\phi(\xi) = \frac{1}{1 + \exp(-\xi)}$	$(0, 1)$	(III.16)
Gaussienne	$\phi(\xi) = \exp\left(-\frac{\ \xi - \mathbf{v}\ ^2}{2\sigma^2}\right)$	$(0, 1)$	(III.17)
Hyperbolique (Tangente hyperbolique)	$\phi(\xi) = \frac{\exp(\xi) - \exp(-\xi)}{\exp(\xi) + \exp(-\xi)}$	$(-1, 1)$	(III.18)
Sinusoïdale	$\phi(\xi) = \sin(\xi)$	$[-1, 1]$	(III.19)

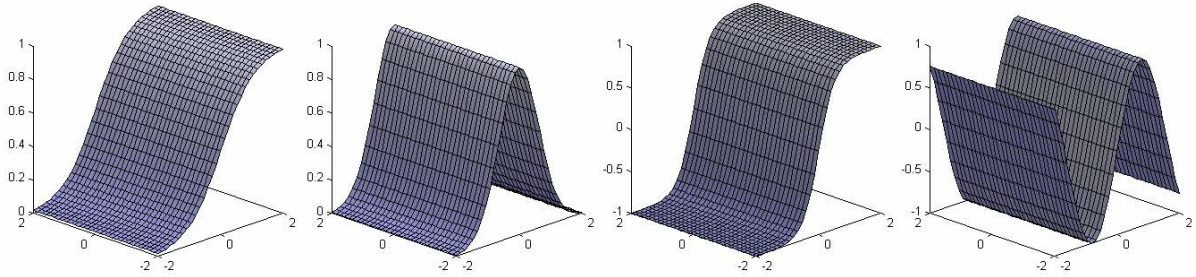


Figure III.3 : Les fonctions d'activation les plus utilisées. De gauche à droite : sigmoïdale, gaussienne, tangente hyperbolique et sinusoidale.

La recherche des paramètres $\mathbf{w}$ et b de la fonction (III.14) revient à minimiser l'erreur d'approximation ou d'apprentissage à partir de l'ensemble d'apprentissage (III.2). Une fonction objectif générale est :

$$J = \frac{1}{\ell} \sum_{i=1}^{\ell} \mathcal{G}(y_i - \hat{y}_i), \quad (\text{III.20})$$

où y_i est la sortie désirée (cible) pour l'entrée $\mathbf{x}_i$, les deux faisant partie de la $i^{\text{ème}}$ observation, $\hat{y}_i$ est la sortie estimée par (III.14), ℓ est le nombre de vecteurs d'apprentissage et $\mathcal{G}$ la fonction perte.

Les fonctions perte $\mathcal{G}$ de type $|\xi|^p$ avec $p > 1$ peuvent devenir indésirables car un incrément superlinéaire amène à une perte de la robustesse de l'estimateur [143]: dans ce cas, la dérivée de la fonction de coût peut croître sans borne. Pour $p < 1$ la fonction perte $\mathcal{G}$ devient non-convexe. Pour le cas $p=2$, nous retrouvons l'approche des **moindres carrés** (least mean squares, LMS). Le choix de la fonction perte $\mathcal{G}$ doit assurer un compromis entre sa complexité, pour éviter des problèmes d'optimisation difficiles, et le niveau d'adaptation aux données. Néanmoins, supposant que les exemples aient été générés par une fonction de dépendance $y_i = f_{\text{réelle}}(\mathbf{x}_i) + \varepsilon_i$, entachée d'un bruit de densité $p(\varepsilon)$, la fonction optimale au sens du **maximum de vraisemblance** (maximum likelihood, ML) est [143] :

$$\mathcal{G}(\mathbf{x}, y, f(\mathbf{x})) = -\log p(y - \hat{y}). \quad (\text{III.21})$$

Bien que la fonction de coût résultante puisse être non-convexe, on peut trouver une fonction convexe proche afin de réaliser l'implémentation correspondante. Le Tableau III.2 contient quelques modèles de densité correspondant aux fonctions perte, comme celles définies par (III.21), tandis que la Figure III.4 illustre ces fonctions.

TABLEAU III.2
Les fonctions perte plus communes et leurs modèles de densité correspondants.

	FONCTION DE PERTE	MODELE DE DENSITE
Gaussienne ou quadratique	$\mathcal{G}(\xi) = \frac{1}{2} \xi^2$	$p(\xi) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{\xi^2}{2}\right)$
Laplacienne	$\mathcal{G}(\xi) = \xi $	$p(\xi) = \frac{1}{2} \exp(- \xi)$
Fonction de perte de Huber	$\mathcal{G}(\xi) = \begin{cases} \frac{1}{2\sigma} \xi^2 & \text{si } \xi < \sigma \\ \xi & \text{autrement} \end{cases}$	$p(\xi) \propto \begin{cases} \exp\left(-\frac{\xi^2}{2}\right) & \text{si } \xi < \sigma \\ \exp\left(\frac{\sigma}{2} - \xi \right) & \text{autrement} \end{cases}$
ε -insensible	$\mathcal{G}(\xi) = \xi _\varepsilon$	$p(\xi) = \frac{1}{2(1+\varepsilon)} \exp\left(- \xi _\varepsilon\right)$
Polynomiale	$\mathcal{G}(\xi) = \frac{1}{p} \xi^p$	$p(\xi) = \frac{p}{2\Gamma(1/p)} \exp\left(- \xi ^2\right)$
Polynomiale par morceaux	$\mathcal{G}(\xi) = \begin{cases} \frac{1}{p\sigma^{p-1}} (\xi)^p & \text{si } \xi < \sigma \\ \xi - \sigma \frac{p-1}{p} & \text{autrement} \end{cases}$	$p(\xi) \propto \begin{cases} \exp\left(-\frac{\xi^p}{p\sigma^{p-1}}\right) & \text{si } \xi < \sigma \\ \exp\left(\sigma \frac{p-1}{p} - \xi \right) & \text{autrement} \end{cases}$

La version robuste de la fonction objectif (III.20) est :

$$J = \frac{1}{\ell} \sum_{i=1}^{\ell} \mathcal{G}(y_i - \hat{y}_i) + \lambda (\mathbf{w}^T \mathbf{w} - c^2)^2, \quad (\text{III.22})$$

où λ et c^2 sont les paramètres du terme de pénalisation (régularisation). Le premier terme est la partie de la fonction de coût qui minimise l'erreur d'approximation ou d'apprentissage, tandis que le deuxième contrôle le nombre d'erreurs ou d'exemples sortant de l'ensemble d'apprentissage que nous pouvons ignorer au moment de construire une fonction. Ensemble, ces deux termes essaient de minimiser la borne

(III.9). L'utilisation de la fonction perte gaussienne associée à ce terme de régularisation est connu comme distance de Kulback–Leibler, [76]. Une alternative au deuxième terme est le « weight decay » $+\lambda \mathbf{w}^T \mathbf{w}$ [26].

Finalement, dans les réseaux de neurones artificiels, la méthode du gradient (connu sous le nom de **règle delta** dans ce domaine) est la méthode la plus utilisée pour minimiser la fonction objectif (III.20) et trouver les poids. A l'itération k , les paramètres sont actualisés selon les équations :

$$\mathbf{w}_{(k+1)} = \mathbf{w}_{(k)} - \eta \frac{\partial J_{(k)}}{\partial \mathbf{w}}, \quad b_{(k+1)} = b_{(k)} - \eta \frac{\partial J_{(k)}}{\partial b}, \quad (\text{III.23})$$

où $\partial J_{(k)} / \partial (\cdot)$ est le gradient de la fonction objectif J par rapport à $(\cdot)$ et η est le pas d'apprentissage. On peut, évidemment, envisager l'utilisation d'autres techniques d'optimisation plus puissantes.

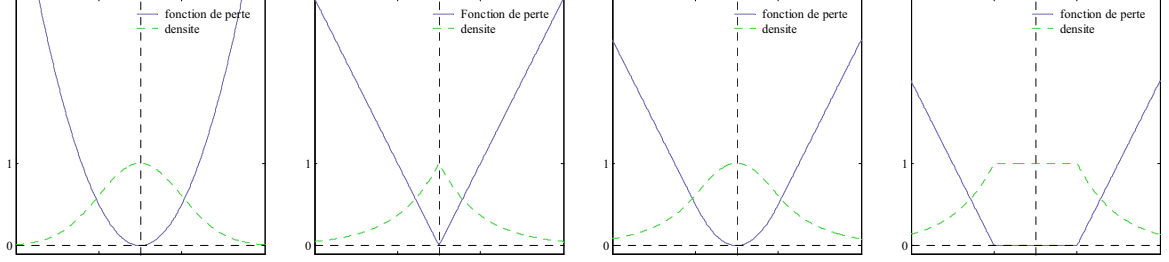


Figure III.4 : Des fonctions de perte et leurs modèles de densité correspondants. De gauche à droite : gaussienne, Laplacienne, robuste de Huber et ε -insensible.

Le perceptron monocouche comme classifieur statistique

Le perceptron monocouche a la propriété de devenir un bon classifieur statistique comparatif qui, après une seule itération, permet de retrouver d'autres classifieurs connus. Pour cela, il faut satisfaire les conditions suivantes [122] :

- C1) Le centre de l'ensemble d'apprentissage est zéro, $\mathbf{m} = \frac{1}{2}(\mathbf{m}_1 + \mathbf{m}_{-1}) = \mathbf{0}$;
- C2) Les conditions initiales sont $\mathbf{w} = \mathbf{0}$ et $b = 0$;
- C3) La cible est $y_{-1} = -y_1 \ell_1 / \ell_{-1}$;
- C4) Le gradient total d'apprentissage est utilisé.

Il faut noter que, dans le gradient total (connu dans le domaine des ANN comme époque), on actualise les poids après avoir calculé les gradients pour *tous* les vecteurs d'apprentissage. Pour le gradient stochastique, on actualise les poids après avoir calculé le gradient pour chaque vecteur d'apprentissage.

Le classifieur à distance euclidienne

Pour retrouver le classifieur à distance euclidienne (euclidean distance classifier, EDC) à partir du perceptron, considérons la fonction d'activation tangente hyperbolique (III.18), et les cibles $y_1 > 0$, $y_{-1} < 0$. Définissons $\ell = \ell_1 + \ell_{-1}$, $k_1 = \ell_1 / \ell$, $k_{-1} = \ell_{-1} / \ell$, et soit :

$$\mathbf{m} = \frac{1}{\ell} \sum_{i=1}^{\ell} \mathbf{x}_i = k_1 \mathbf{m}_1 + k_{-1} \mathbf{m}_{-1}$$

le centre de l'ensemble d'apprentissage. Durant la première itération, les poids et le terme $\mathbf{w}^T \mathbf{x} + b$ sont proches de zéro. En conséquence, pour la fonction d'activation (III.18), nous avons :

$$\hat{y}(\mathbf{w}^T \mathbf{x} + b) \cong \mathbf{w}^T \mathbf{x} + b, \quad \text{et} \quad \partial \hat{y}(b) \cong db.$$

Donc, le gradient de la fonction objectif (III.20), avec une fonction perte gaussienne est :

$$\left. \frac{\partial J}{\partial b} \right|_{b=b_{(k)}} = 2(\mathbf{w}_{(k)}^T (k_1 \mathbf{m}_1 + k_{-1} \mathbf{m}_{-1}) + b_{0(k)}) - (y_1 k_1 + y_{-1} k_{-1}), \quad (\text{III.24})$$

$$\left. \frac{\partial J}{\partial \mathbf{w}} \right|_{\mathbf{w}=\mathbf{w}_{(k)}} = 2(\mathbf{K} \mathbf{w}_{(k)} + (k_1 \mathbf{m}_1 + k_{-1} \mathbf{m}_{-1}) b_{(k)}) - (y_1 k_1 \mathbf{m}_1 + y_{-1} k_{-1} \mathbf{m}_{-1}), \quad (\text{III.25})$$

$$\text{où} \quad \mathbf{K} = \frac{1}{\ell} \sum_{i=1}^{\ell} \mathbf{x}_i \mathbf{x}_i^T. \quad (\text{III.26})$$

En considérant que $b_{(0)}=0$ et $\mathbf{w}_{(0)}=\mathbf{0}$ (condition C2), après la première itération dans le processus d'apprentissage, nous avons donc :

$$b_{(1)} = 2\eta(y_1 k_1 + y_{-1} k_{-1}), \text{ et } \mathbf{w}_{(1)} = \mathbf{w}_{(0)} - 2\eta \frac{\partial J}{\partial \mathbf{w}} \Big|_{\mathbf{w}=\mathbf{w}_{(0)}} = 2\eta(y_1 k_1 \mathbf{m}_1 + y_{-1} k_{-1} \mathbf{m}_{-1}). \quad (\text{III.27})$$

Si les cibles satisfont la condition C3, alors :

$$b_{(1)} = 0, \quad \mathbf{w}_{(1)} = \eta \Delta \mathbf{m}, \quad (\text{III.28})$$

où $\Delta \mathbf{m} = \mathbf{m}_1 - \mathbf{m}_{-1}$.

De cette manière, le vecteur (III.28) est proportionnel au vecteur de poids du EDC pour des données centrées ($\mathbf{m}=\mathbf{0}$). La classification avec le signe de la fonction discriminante aux poids (III.28) est asymptotiquement optimale ($\ell \rightarrow \infty$) quand les classes sont des gaussiennes à norme euclidienne, mais aussi dans d'autres situations. Si le nombre de vecteurs d'apprentissage de chaque classe est différent ($\ell_1 \neq \ell_{-1}$), pour obtenir le EDC après la première utilisation, il suffit d'utiliser les cibles asymétriques donnant (III.27).

L'analyse discriminante régularisée

A présent, nous pouvons analyser les dynamiques du vecteur de poids aux itérations qui suivent la première. Soit $\ell_1 = \ell_{-1} = \ell_m > n/2 + 1$, alors, la matrice de covariance espérée $\hat{\Sigma}$ n'est pas singulière. Supposons que le pas d'apprentissage η est petit et que, après k itérations, les poids restent petits. Alors, la quantité $|\mathbf{w}^T \mathbf{x}_i + b|$ est aussi petite. Dans ce cas, la fonction d'activation (III.18) a un comportement linéaire. Après la deuxième itération, l'utilisation du gradient total d'apprentissage (III.23) avec les gradients (III.24) et (III.25) résulte en :

$$b_{(2)} = 0, \text{ et } \mathbf{w}_{(2)} = \mathbf{w}_{(1)} - 2\eta \frac{\partial J}{\partial \mathbf{w}} \Big|_{\mathbf{w}=\mathbf{w}_{(1)}} = \eta \Delta \mathbf{m} - \eta(-\Delta \mathbf{m} + 2\mathbf{K}\mathbf{w}_1) = (\mathbf{I} - (\mathbf{I} - \eta\mathbf{K})^2) \mathbf{K}^{-1} \Delta \mathbf{m}.$$

après plusieurs itérations,

$$b_{(k)} = 0, \text{ et } \mathbf{w}_{(k)} = (\mathbf{I} - (\mathbf{I} - \eta\mathbf{K})^k) \mathbf{K}^{-1} \Delta \mathbf{m}. \quad (\text{III.29})$$

$$\text{où } \mathbf{K} = \frac{1}{\ell} \sum_{i=1}^{\ell} \mathbf{x}_i \mathbf{x}_i^T = \frac{\ell_m + 1}{\ell_m} \hat{\Sigma} + \frac{1}{4} \Delta \mathbf{m} \Delta \mathbf{m}^T.$$

Pour dériver (III.27), supposons la matrice $\mathbf{K}$ non singulière. En utilisant les premiers termes des développements :

$$(\mathbf{I} - \eta\mathbf{K})^k = \mathbf{I} - k\eta\mathbf{K} + \frac{1}{2}k(k-1)\eta^2\mathbf{K}^2 - \dots \text{ et } (\mathbf{I} - \beta\mathbf{K})^{-1} = \mathbf{I} + \beta\mathbf{K} + \dots$$

pour une valeur petite de η et une valeur petite de k , on peut obtenir :

$$\begin{aligned} \mathbf{w}_{(k)} &\approx (k\eta\mathbf{K} - \frac{1}{2}k(k-1)\eta^2\mathbf{K}^2) \mathbf{K}^{-1} \Delta \mathbf{m} = k\eta(\mathbf{I} - \frac{1}{2}\mathbf{K}(k-1)\eta\mathbf{K})^k \Delta \mathbf{m} \\ &\approx k\eta(\mathbf{I} + \frac{1}{2}(k-1)\eta\mathbf{K})^{-1} \Delta \mathbf{m} = \frac{2k}{(k-1)} \left(\mathbf{I} - \frac{2}{(k-1)\eta} + \mathbf{K} \right)^{-1} \Delta \mathbf{m} \\ &= \frac{2k}{(k-1)} \left(\mathbf{I} - \frac{2}{(k-1)\eta} + \left(\frac{\ell_m - 1}{\ell_m} \hat{\Sigma} + \frac{1}{4} \Delta \mathbf{m} \Delta \mathbf{m}^T \right)^{-1} \right) \Delta \mathbf{m}. \end{aligned}$$

Supposons que la matrice de covariance espérée $\hat{\Sigma}$ est non singulière. Après un peu de calcul, on trouve :

$$\mathbf{w}_{(k)} = \left(\mathbf{I} - \frac{2}{(k-1)\eta} \frac{\ell_m - 1}{\ell_m} + \hat{\Sigma} \right)^{-1} \Delta \mathbf{m} k_R, \quad b_{(k)} = 0 \quad (\text{III.30})$$

où k_R est un coefficient scalaire.

Pour classifier avec le signe de la fonction discriminante, le coefficient k_R n'est pas significatif. Ainsi, le vecteur $\mathbf{w}_{(k)}$ est équivalent à l'analyse discriminante régularisée (regularized discriminant analysis, RDA) :

$$\hat{\Sigma} + \lambda \mathbf{I} \quad (\text{III.31})$$

$$\text{avec } \lambda = \frac{2}{(k-1)\eta} \frac{\ell_m - 1}{\ell_m}.$$

Le classifieur linéaire standard de Fisher

L'expression (III.30) indique que, au moment d'augmenter le nombre d'itérations, $\left(\mathbf{I} \frac{2}{(k-1)\eta} \frac{\ell_m-1}{\ell_m} + \hat{\Sigma}\right) \rightarrow \hat{\Sigma}$ et le classifieur résultant approche la fonction de décision linéaire standard de Fisher [122]. Ce classifieur peut être également obtenu en utilisant le perceptron linéaire (fonction d'activation identité), en faisant les dérivés (III.24) et (III.25) égales à zéro et en résolvant les équations résultantes par rapport à $\mathbf{w}$ et b . Alors, pour des données centrées, nous avons $b=0$, et $\mathbf{w}=\mathbf{K}^{-1}(\mathbf{m}_1-\mathbf{m}_{-1})$.

Le classifieur pseudo Fisher

En dérivant le vecteur (III.30), on suppose que la matrice de covariance espérée est non singulière. Si le nombre d'exemples d'apprentissage $\ell < n+2$, la matrice $\hat{\Sigma}$ est singulière. Maintenant, avec un incrément du nombre d'itérations, le classifieur s'apparente à la DF de Fisher avec pseudo inversion [122]. Le vecteur (III.29) peut être écrit sous une forme telle que la matrice $\mathbf{K}$ n'a plus besoin d'être non-singulière :

$$\mathbf{w}_{(k)} = \sum_{s=1}^k C_k^s (-\eta)^s (\mathbf{K}^s)^{-1} \Delta \mathbf{m}. \quad (\text{III.32})$$

Soit la matrice orthogonale $\phi \in \mathcal{R}^{n \times n}$ qui satisfait $\phi^T \mathbf{K} \phi = \begin{bmatrix} \lambda & 0 \\ 0 & 0 \end{bmatrix}$. Alors

$$\mathbf{w}_{(k)} = \phi \left(\mathbf{I} - \left(\mathbf{I} - \eta \begin{bmatrix} \lambda & 0 \\ 0 & 0 \end{bmatrix} \right)^T \right) \phi^T \mathbf{w}^{\text{PF}}, \quad (\text{III.33})$$

$$\text{où } \mathbf{w}^{\text{PF}} = \phi \begin{bmatrix} \lambda & 0 \\ 0 & 0 \end{bmatrix} \phi^T \Delta \mathbf{w} = \mathbf{K}^+ \Delta \mathbf{w}.$$

Pour des valeurs faibles de η avec un incrément en k :

$$\left(\mathbf{I} - \left(\mathbf{I} - \eta \begin{bmatrix} \lambda & 0 \\ 0 & 0 \end{bmatrix} \right)^T \right) \rightarrow \mathbf{I}.$$

Alors $\mathbf{w}_{(k)} \rightarrow \mathbf{w}^{\text{Fpseudo}}$, conclusion qui reste valable quand les entrées $\mathbf{w}^T \mathbf{x} + b$ varient dans la partie linéaire de la fonction d'activation (III.18).

III.2.2.2 Les machines non-linéaires

Transformation des approches linéaires dans des espaces de caractéristiques

Une approche standard pour obtenir des frontières de décision non-linéaires, en utilisant des algorithmes linéaires, est de transformer les vecteurs d'entrée $\mathbf{x}$ dans un espace de dimension plus importante à l'aide d'une fonction non-linéaire, choisie *a priori* :

$$\phi(\mathbf{x}) = [x_{i_1}, \dots, x_{i_d}], \text{ avec } 1 \leq i_1 \dots i_d \leq n,$$

Dans cet espace, on construit un hyperplan qui définit une fonction non-linéaire dans l'espace original d'entrée. Une première approche correspond aux transformations polynomiales [137]. Une approche alternative est l'utilisation de fonctions *noyau* ou *kernel*. Cette astuce a connu un énorme succès ; des exemples sont les SVM [155], le perceptron [93], l'analyse discriminante régularisée [16], entre autres. Nous abordons cette approche plus en détail dans la section dédiée aux SVM.

La transformation polynomiale est le choix le plus populaire pour former « l'espace de caractéristiques » à travers une transformation non-linéaire, choisie *a priori*, par exemple :

$$\phi_{\beta}(\mathbf{x}) = (\mathbf{v}_{\beta}^T \mathbf{x})^k + c_{\beta} \quad (\text{III.34})$$

où k est le degré du polynôme, c_{β} et $\mathbf{v}_{\beta} \in \mathcal{R}^{n_{\beta}}$, $n_{\beta} > n$, sont des coefficients de la β^{me} fonction de transformation. Un deuxième choix correspond aux fonctions à base radiale comme la densité sphérique gaussienne :

$$\phi_{\beta}(\mathbf{x}) = \frac{1}{\sqrt{2\pi\sigma_{\beta}}} \exp\left(-\frac{\|\mathbf{x} - \mathbf{v}_{\beta}\|^2}{2\sigma_{\beta}^2}\right) \quad (\text{III.35})$$

où σ_β est le rayon et $\mathbf{v}_\beta \in \mathbb{R}^{n_\beta}$, $n_\beta > n$, le centre de la $\beta^{\text{ème}}$ gaussienne. En principe, on peut utiliser des fonctions à base radiale plus compliquées. Une autre transformation est la somme pondérée non-linéaire :

$$\phi_\beta(\mathbf{x}) = f(\mathbf{v}_\beta^T \mathbf{x} + c_\beta) \quad (\text{III.36})$$

où $f(\cdot)$ est une fonction non-linéaire, comme celles décrites dans le Tableau III.1, c_β et $\mathbf{v}_\beta \in \mathbb{R}^{n_\beta}$, $n_\beta > n$, sont des coefficients de la $\beta^{\text{ème}}$ fonction de transformation.

Dans cette approche, les paramètres σ_β , c_β , et $\mathbf{v}_\beta$ sont également choisis *a priori*. Donc pour régler la fonction, il reste seulement trouver les paramètres $\mathbf{w}$ et b de la fonction finale :

$$\hat{y} = \sum_{\beta=1}^{n_\beta} w_\beta \phi_\beta(\mathbf{x}) + b = \mathbf{w}^T \boldsymbol{\phi}(\mathbf{x}) + b, \quad (\text{III.37})$$

à l'aide de n'importe quel algorithme linéaire. Malheureusement, il existe plusieurs problèmes dans l'espace des caractéristiques, comme le *fléau de la dimension*. Nous aborderons ce problème dans la section dédiée aux SVM.

Le perceptron multicouche

Le perceptron multicouche (multi layer perceptron, MLP) est le réseau le plus ancien et le plus utilisé des réseaux de neurones artificiels. Un MLP typique peut avoir plusieurs couches de neurones artificiels formés de perceptrons monocouche, Figure III.5. Les nœuds d'entrée correspondent aux n composantes de l'espace d'entrée. La couche de sortie peut avoir un neurone (pour obtenir une fonction discriminante à deux classes ou une fonction de régression) ou plusieurs neurones (correspondant chacun à une classe ou à une fonction de régression : régression multiple). Habituellement appelés ANN d'injection en avant (feedforward), les signaux d'entrée se propagent de la couche d'entrée jusqu'à la couche de sortie.

On peut construire des MLP complexes en ajoutant plusieurs couches cachées Figure III.5, mais la théorie montre que le MLP standard, avec une seule couche cachée et des fonctions d'activation régulières, non-linéaires, peuvent approximer n'importe quelle fonction non-linéaire [26]. Le schéma de traitement de l'information est similaire à (III.37), avec des transformations non-linéaires (III.36), sauf que les coefficients c_β , et $\mathbf{v}_\beta$ doivent être trouvés à partir de l'ensemble d'apprentissage.

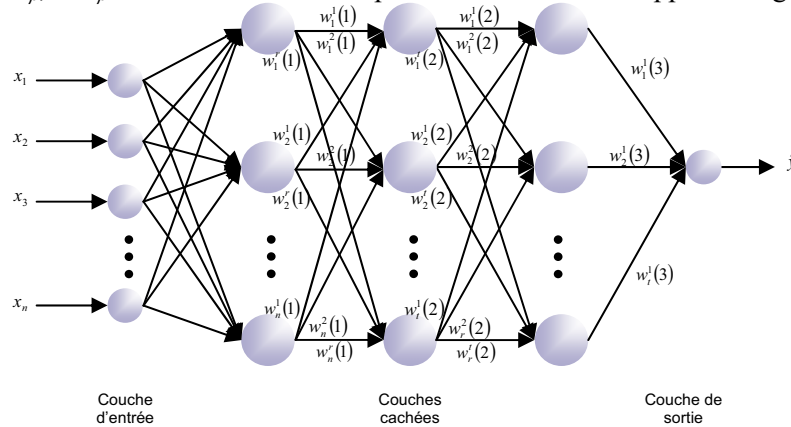


Figure III.5 : Exemple de perceptron multicouche (trois couches cachées).

L'algorithme d'apprentissage des MLP le plus connu est la méthode du gradient ou de rétropropagation de l'erreur (*backpropagation*) [103]. Des algorithmes de deuxième ordre comme la méthode de gradient conjugué, la méthode de *Levenberg-Marquardt* et la méthode de *Quasi-Newton*, sont substantiellement plus rapides (d'un ordre d'amplitude) pour plusieurs problèmes. La méthode de Quasi-Newton est recommandée pour traiter des problèmes de classification tandis que la méthode de Levenberg-Marquardt est mieux adaptée pour des problèmes de régression [26]. Néanmoins, la méthode de rétropropagation de l'erreur a encore des avantages en quelques circonstances et est l'algorithme le plus facile à comprendre.

Dans la méthode de rétropropagation, similairement au SLP, on calcule les dérivés de la fonction objectif et les gradients des neurones des couches cachées. Itérativement, à chaque pas, ou époque, l'algorithme calcule l'erreur entre la sortie estimée $\hat{y}_i$ et la cible y_i pour tous les vecteurs d'apprentissage et l'utilise pour actualiser les poids grâce aux gradients des neurones. La configuration initiale du réseau

est aléatoire et le processus d'apprentissage s'achève quand un nombre maximal d'époques est atteint ou quand l'erreur atteint un certain niveau ou quand l'erreur ne peut être plus minimisée (critère d'arrêt). Il existe des modifications heuristiques comme l'algorithme quick propagation et l'algorithme Delta-Bar-Delta. Pour plus de détails, voir [26], [103] et [139].

Les réseaux à fonctions de base radiale

Au milieu des années 1960, Aizerman [1] a proposé la *méthode des fonctions potentiel* pour estimer la dépendance fonctionnelle des exemples (III.2) en utilisant l'ensemble de fonctions :

$$\hat{y} = f\left(\sum_{i=1}^{\ell} w_i \phi(\|\mathbf{x} - \mathbf{x}_i\|)\right). \quad (\text{III.38})$$

La fonction $\phi(\xi)$, appelée *fonction potentiel* (par analogie aux potentiels physiques), est une fonction monotone qui converge vers zéro quand $|\xi|$ augmente, par exemple $\phi(\xi) = \exp\{-\gamma|\xi|\}$. Après 20 ans d'oubli, ce type de fonctions est réapparu. Dans les années 1980, Powell a proposé les réseaux à fonctions de base radiale (radial basis functions networks, RBFN) [121].

L'architecture des RBFN est similaire aux réseaux MLP, à une ou plusieurs couches cachées. Les RBFN standards ont une seule couche cachée, avec des fonctions (III.35) ou, de façon plus générale :

$$\phi_{\beta}(\mathbf{x}_i) = z_{\beta} = \frac{1}{(2\pi)^{n/2} \sqrt{\det(\mathbf{A}_{\beta})}} \exp\left(-\frac{1}{2}(\mathbf{x}_i - \mathbf{v}_{\beta})^T \mathbf{A}_{\beta}^{-1}(\mathbf{x}_i - \mathbf{v}_{\beta})\right). \quad (\text{III.39})$$

où $\mathbf{A}_{\beta}$ est la matrice de variance-covariance (si $\mathbf{A}_{\beta} = \mathbf{I}\sigma_{\beta}^2$ on retrouve (III.35)) pour le $\beta^{\text{ème}}$ neurone, mais on peut utiliser d'autres fonctions à base radiale. Comme dans les MLP, la différence est que, dans les RBFN, les coefficients des fonctions de transformation doivent être trouvés à partir de la base d'apprentissage.

Dans l'approche RBFN, un problème important est le choix des cibles, y_i . Une possibilité est une notation binaire, $\{0,1\}$, pour indiquer le niveau d'appartenance de chaque vecteur $\mathbf{x}_i$ aux classes, dans le cadre de la reconnaissance de formes. Ceci donne une interprétation très utile des sorties car ce que l'algorithme estime est la ***fonction de densité de probabilité*** (probability density function, PDF). Une représentation similaire pour le cas de la régression existe si la sortie est vue comme la valeur espérée du modèle en un point donné de l'espace d'entrée. Cette valeur est reliée à la PDF adjointe de la sortie et des entrées. Dans la section suivante, nous aborderons le problème de l'estimation de la densité.

La couche de sortie est en général une somme pondérée de gaussiennes, similaire à (III.37) avec $b=0$. Comme sortie du réseau, on peut utiliser également la relation :

$$P_c(\mathbf{x}) = \frac{\sum_{\beta=1}^{n_{\beta}} w_{c\beta} \phi(\mathbf{x}, \mathbf{v}_{c\beta}, \sigma_{c\beta})}{\sum_{i=1}^{n_c} \sum_{\beta=1}^{n_{\beta}} w_{i\beta}}. \quad (\text{III.40})$$

Si le numérateur de l'équation (III.40) décrit le comportement de la densité conditionnelle de la classe $f_c(\mathbf{x})$, l'expression représente donc la probabilité *a posteriori* de la classe c , $c \in \{1, \dots, n_c\}$. L'utilisation de (III.40) dans la fonction de coût (III.20), ou (III.22), résulte en une fonction différentiable avec un gradient facile à trouver. Cette approche peut être utilisée pour introduire formellement une option de rejet dans l'algorithme. Ceci donne lieu aux ***réseaux de neurones RBFN probabilistes***, l'un des réseaux Bayésiens, [26], [48].

Avant de lancer le processus d'apprentissage, il faut choisir le nombre d'unités et, si possible, les valeurs initiales des coefficients $\sigma_{i\beta}$, $\mathbf{v}_{i\beta}$, et $\mathbf{w}_{i\beta}$. Pour ce faire, on peut utiliser des méthodes de décomposition en mélanges statistiques ou d'analyse de groupes ou clusters standard, comme les *k*-moyennes (*k*-means) [48], dont on parlera dans la section dédiée aux techniques de groupage flou.

L'estimation des paramètres revient à la minimisation de la fonction de coût (III.20) ou, s'il faut ajouter un terme de pénalité, la fonction de coût (III.22). L'algorithme d'optimisation dépendra alors du type de problème que l'on traite, comme cité dans la section dédiée au MLP.

Les réseaux d'apprentissage à quantification vectorielle

Les réseaux d'apprentissage à quantification vectorielle (learning vector quantization, LVQ) sont une simplification du processus de calcul dans les réseaux RBFN durant la phase de reconnaissance. A la place du numérateur de l'équation (III.40), on analyse chacune des n_β composantes comme des nouvelles sous-classes. De cette manière, on peut construire une fonction linéaire discriminante par morceaux :

$$h(\mathbf{x}) = \min_j \{ (\mathbf{x} - \mathbf{v}_{-1j})^T \mathbf{A}^{-1} (\mathbf{x} - \mathbf{v}_{-1j}) \} - \min_j \{ (\mathbf{x} - \mathbf{v}_{1j})^T \mathbf{A}^{-1} (\mathbf{x} - \mathbf{v}_{1j}) \}. \quad (\text{III.41})$$

où $\mathbf{A}$ est la matrice diagonale $\mathbf{I}\sigma^2$ ou la matrice de variance covariance. A nouveau, comme les valeurs initiales des coefficients $\sigma_{i\beta}$ et $\mathbf{v}_{i\beta}$ sont inconnues, on peut utiliser des méthodes de décomposition en mélanges statistiques ou d'analyse de clusters standard [48].

Le Tableau III.3 montre une classification des principaux algorithmes pour la reconnaissance de formes et la régression, selon leur type de frontière de décision associée ; linéaire ou non-linéaire.

TABLEAU III.3
Taxonomie des méthodes pour la reconnaissance de formes et la régression.

	LINEAIRES	NON-LINEAIRES
Machines statistiques paramétriques	- Les méthodes des moindres carrés (least squares methods) : classification - La DF linéaire standard de Fisher - La DF linéaire régularisée - La DF linéaire à distance Euclidienne - La DF linéaire Anderson-Bahadur - La méthode des moindres carrés (least square estimate) : régression - Gauss et Legendre - Pondérés (weighthed), - Robuste (ridge regression) - Prédiction et corrélation - Modèles linéaires généralisés	- La DF quadratique standard - La méthode des moindres carrés non-linéaires (Least squares for nonlinear models) - Transformations pour linéariser le modèle - Empiriques - Transformation des entrées $\mathbf{x}$ - Transformations en y (la méthode Box-Cox) - Exponentielles (un et deux termes) - Gaussiennes - Séries de puissances (power series) - Rationnelles - Somme de sinus - Estimation Bayésienne - Maximum de vraisemblance - Mélanges de gaussiennes - A fenêtres de Parzen - Plus proche voisin - Interpolation linéaire - Polynomiales par morceaux : splines, Hermite cubique - Termes polynomiaux et trigonométriques, Orthogonalisation
Machines statistiques non-paramétriques	- Linéaires par morceaux (piece-wise) - A fenêtres de Parzen - Interpolation (Linéaire, plus proche voisin) - Polynomiales de degré un	- Recherche stochastique - Les réseaux de Boltzmann et les modèles graphiques - Méthodes évolutionnaires : programmation génétique - Transformation des approches linéaires dans des espaces de caractéristiques - Méthodes à base de fonctions kernel - Maximum de vraisemblance - Les SVM, le SLP, - Le perceptron multicouches (MLP) - Les réseaux à fonctions de base radiale (RBFN) - Les réseaux d'apprentissage à quantification vectorielle (LVQ)
Méthodes Stochastiques	- Approximation stochastique et l'algorithme LMS - Recherche stochastique - Les réseaux de Boltzmann et les modèles graphiques - Méthodes évolutionnaires : programmation génétique - Le perceptron monocouche (SLP) - SLP comme classifieur statistique - linéaire standard de Fischer (et pseudo) - à distance euclidienne - analyse discriminante régularisée - analyse discriminante robuste - à marge maximale SVM - SLP comme régresseur statistique - L'algorithme de Widrow-Hoff - Les SVM	
Les réseaux de neurones		

III.2.2.3 L'estimation de la densité

Le but de l'estimation de la densité est de trouver une approximation d'une *fonction de densité de probabilité* $p(\mathbf{x})$ à partir d'exemples avec la fonction de distribution correspondante :

$$F(\mathbf{x}) = \Pr(X \leq \mathbf{x}) = \int_{-\infty}^{\mathbf{x}} p(t) dt. \quad (\text{III.42})$$

Une fonction $p(\mathbf{x})$ est une fonction de densité de probabilité de la variable aléatoire et continue X si, pour n'importe quel intervalle $(\mathbf{x}_1, \mathbf{x}_2)$ de $\mathcal{R}^n$, elle satisfait :

$$p(\mathbf{x}) \geq 0 \quad (\text{III.43})$$

$$\int_{-\infty}^{\infty} p(\mathbf{x}) d\mathbf{x} = 1 \quad (\text{III.44})$$

$$\Pr(\mathbf{x}_1 \leq X \leq \mathbf{x}_2) = \int_{\mathbf{x}_1}^{\mathbf{x}_2} p(u) du \quad (\text{III.45})$$

La propriété (III.45) définit la densité comme la dérivée de la fonction de distribution cumulée $F(\mathbf{x})$.

Le problème de l'estimation de la densité revient donc à résoudre en $p(t)$ l'équation linéaire :

$$\int_{-\infty}^{\infty} \theta(\mathbf{x} - t) p(t) dt = F(\mathbf{x}) \quad (\text{III.46})$$

où, au lieu de connaître la fonction de distribution $F(\mathbf{x})$, nous avons les exemples indépendants et identiquement distribués (iid), générés par F :

$$\mathbf{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_\ell\}. \quad (\text{III.47})$$

Le problème de l'estimation de la densité est *mal posé* puisque pour trouver p satisfaisant $Ap = F$, avec A un opérateur linéaire, de grandes variations de p peuvent résulter en des petites variations de F . Il existe deux types d'estimation de densité : *paramétrique* et *non-paramétrique*.

Estimation paramétrique

Dans l'approche paramétrique, on doit supposer que les données appartiennent à une famille paramétrique de distributions connue. La densité $p(\mathbf{x})$ sous-jacente aux données peut être trouvée en cherchant les valeurs des paramètres $\theta = \{\theta^1, \dots, \theta^n\}$ à partir des données (III.47). Ainsi, dans l'estimation paramétrique, nous pouvons améliorer la qualité des modèles que l'expertise développe en incorporant tout type de connaissance qualitative ou quantitative disponible.

Maximum de vraisemblance

Dans l'approche ML, la densité résultante pour les exemples est :

$$p(\mathbf{X} | \theta) = \prod_{i=1}^{\ell} p(\mathbf{x}_i | \theta) = L(\theta | \mathbf{X}). \quad (\text{III.48})$$

La fonction $L(\theta | \mathbf{X})$ est la fonction de vraisemblance, [25]. La vraisemblance est fonction des paramètres θ où les données $\mathbf{X}$ sont fixes. Le but du problème du maximum de vraisemblance est de trouver θ qui maximise L , c'est-à-dire :

$$\hat{\theta} = \arg \max [L(\theta | \mathbf{X})], \quad (\text{III.49})$$

mais, en général, on maximise $\log(L(\theta | \mathbf{x}))$ car, analytiquement, plus facile à traiter, grâce à la monotonie de la fonction log. Alors, le calcul du maximum de vraisemblance des paramètres θ est :

$$\hat{\theta} = \arg \max \left[\log \prod_{i=1}^{\ell} p(\mathbf{x}_i | \theta) \right] = \arg \max \left[\sum_{i=1}^{\ell} \log p(\mathbf{x}_i | \theta) \right]. \quad (\text{III.50})$$

Estimation Bayésienne

Dans l'approche Bayésienne, la densité d'une nouvelle observation $\mathbf{x}$, connaissant les données $\mathbf{X}$, est calculée par la *densité prédictive a posteriori*, [41], en intégrant les paramètres :

$$p(\mathbf{x} | \mathbf{X}) = \int p(\mathbf{x} | \theta) p(\theta | \mathbf{X}) d\theta. \quad (\text{III.51})$$

Pour calculer cette intégrale, il faut déterminer la valeur de $p(\theta | \mathbf{X})$, connue comme la *densité a posteriori* de θ :

$$p(\boldsymbol{\theta} | \mathbf{X}) = \frac{p(\mathbf{X} | \boldsymbol{\theta}) p(\boldsymbol{\theta})}{p(\mathbf{X})} = \frac{\left\{ \prod_{i=1}^{\ell} p(\mathbf{x}_i | \boldsymbol{\theta}) \right\} p(\boldsymbol{\theta})}{\int_{\boldsymbol{\theta}} \left\{ \prod_{i=1}^{\ell} p(\mathbf{x}_i | \boldsymbol{\theta}) \right\} p(\boldsymbol{\theta}) d\boldsymbol{\theta}} \quad (\text{III.52})$$

où $p(\mathbf{X} | \boldsymbol{\theta})$ peut être calculé en utilisant (III.48). L'équation (III.52) constitue le théorème de Bayes qui, de façon simplifiée, en termes de $\boldsymbol{\theta}$ est :

$$p(\boldsymbol{\theta} | \mathbf{X}) \propto p(\mathbf{X} | \boldsymbol{\theta}) p(\boldsymbol{\theta}), \quad \text{Posteriori} \propto \text{Vraisemblance} \times \text{Priori} \quad (\text{III.53})$$

où $\propto$ indique que les deux termes sont proportionnels. Conséquemment, la distribution *a posteriori* est définie en combinant la distribution *a priori* pour les paramètres, $p(\boldsymbol{\theta})$, avec la probabilité d'observer les données avec ces paramètres, $p(\mathbf{X} | \boldsymbol{\theta})$, (la vraisemblance).

Evidemment, il existe une relation entre le maximum de vraisemblance et l'estimation Bayésienne. Si la vraisemblance, $p(\mathbf{X} | \boldsymbol{\theta})$, est relativement accentuée et la densité *a priori* est large, alors l'intégrale suivante va être dominée par la région autour du terme de ML :

$$p(\mathbf{x} | \mathbf{X}) = \int p(\mathbf{x} | \boldsymbol{\theta}) p(\boldsymbol{\theta} | \mathbf{X}) d\boldsymbol{\theta} \cong p(\mathbf{x} | \hat{\boldsymbol{\theta}}) \int p(\boldsymbol{\theta} | \mathbf{X}) d\boldsymbol{\theta} = p(\mathbf{x} | \hat{\boldsymbol{\theta}}).$$

En fait, quand le nombre d'exemples augmente, la densité *a posteriori* $p(\boldsymbol{\theta} | \mathbf{X})$ s'accroît. Donc, l'approche Bayésienne approxime le ML quand $\ell \rightarrow \infty$. En réalité, les deux approches donneront des résultats différents quand le nombre d'exemples est très limité.

Dans la pratique, la solution des problèmes en appliquant directement la règle de Bayes est, en général, infaisable, car son implémentation implique une somme ou une intégration sur un espace large de modèles. Les travaux menés pour éviter ce type de calcul ont fait des méthodes d'approximation Bayésienne un sujet de recherche intéressant. Quelques exemples des machines d'apprentissage Bayésienne sont : l'approximation de Laplace (Laplace's approximation), [7], les approximations de variations (variational approximations), [27], la propagation de l'espérance (expectation propagation), [102], et les chaînes de Markov Monte Carlo (Markov chain Monte Carlo), [106].

Maximisation de l'espérance

L'algorithme de maximisation de l'espérance (expectation maximization, EM), [25], [26], [77], est une méthode générale qui estime le maximum de vraisemblance des paramètres d'une distribution à partir d'un ensemble de données. Dans le cas de l'utilisation des distributions gaussiennes, l'algorithme donne lieu à des modèles de combinaison gaussienne (gaussian mixture models, GMM), donnant le modèle de densité paramétrique suivant :

$$p(\mathbf{x}) = \sum_{j=1}^{n_\beta} \alpha_j p_j(\mathbf{x} | \boldsymbol{\theta}_j), \quad \alpha_j > 0, \quad \sum_{j=1}^{n_\beta} \alpha_j = 1 \quad (\text{III.54})$$

où $p_j(\mathbf{x} | \boldsymbol{\theta}_j)$ sont des gaussiennes sous la forme (III.39), avec $\boldsymbol{\theta}_j$ composée de la moyenne $\mathbf{v}_j$ et de la matrice de covariance $\mathbf{C}_j$, et α_j est la proportion de la $j^{\text{ème}}$ gaussienne dans la combinaison et n_β est le nombre total d'unités.

L'algorithme est composé de deux étapes principales : étape *E* qui calcule la valeur de l'espérance mathématique du logarithme des données inconnues par rapport aux données connues, étape *M* qui maximise le vecteur des paramètres de l'étape *E*. L'itération k de l'algorithme est :

Algorithme III.1. L'algorithme de maximisation de l'espérance, EM.

1. Calcul des poids α_j :

$$\alpha_j^{(k)} = \frac{1}{\ell} \sum_{i=1}^{\ell} \frac{\alpha_j^{(k-1)} p_j^{(k-1)}(\mathbf{x}_i)}{p_j^{(k-1)}(\mathbf{x}_i)}, \quad j=1, \dots, n_\beta, \quad (\text{III.55})$$

2. Calcul des centres $\mathbf{v}_j$:

$$\mathbf{v}_j^{(k)} = \frac{\sum_{i=1}^{\ell} \mathbf{x}_i \frac{\alpha_j^{(k-1)} p_j^{(k-1)}(\mathbf{x}_i)}{p_j^{(k-1)}(\mathbf{x}_i)}}{\sum_{i=1}^{\ell} \frac{\alpha_j^{(k-1)} p_j^{(k-1)}(\mathbf{x}_i)}{p_j^{(k-1)}(\mathbf{x}_i)}}, \quad (\text{III.56})$$

3. Calcul des matrices de covariance

$$\mathbf{C}_j^{(k)} = \frac{\sum_{i=1}^{\ell} (\mathbf{x}_i - \mathbf{v}_j^{k-1})(\mathbf{x}_i - \mathbf{v}_j^{k-1})^T \frac{\alpha_j^{(k-1)} p_j^{(k-1)}(\mathbf{x}_i)}{p_j^{(k-1)}(\mathbf{x}_i)}}{\sum_{i=1}^{\ell} \frac{\alpha_j^{(k-1)} p_j^{(k-1)}(\mathbf{x}_i)}{p_j^{(k-1)}(\mathbf{x}_i)}}, \quad (\text{III.57})$$

Le calcul des centres $\mathbf{v}_i$, et en conséquence le nombre d'unités, peut se faire en utilisant plusieurs approches, mais la plus utilisée est l'algorithme *k-means* pour une partition optimale en k classes [94].

Estimation non-paramétrique

Le but de l'approche non-paramétrique est d'estimer directement la distribution de densité de probabilité $p(\mathbf{x})$ à partir des données (III.47), sans supposer aucune forme paramétrique de la distribution recherchée.

La formulation générale de l'estimation de densité non paramétrique est basée sur deux faits 106[26]: D'un côté, la probabilité qu'un vecteur $\mathbf{x}$, généré par la distribution $p(\mathbf{x})$, appartienne également à une région X est $P = \int_X p(\mathbf{x}) d\mathbf{x}$. Maintenant, supposons que ℓ exemples soient générés à partir de la distribution $p(\mathbf{x})$, la probabilité d'avoir k des ℓ exemples sur la région X est donnée par la distribution binomiale : $P(k) = \binom{\ell}{k} P^k (1-P)^{\ell-k}$. Alors, quand $\ell \rightarrow \infty$, une bonne estimation de P est obtenue par la fraction moyenne des points qui appartiennent à X :

$$P \cong \frac{k}{\ell}. \quad (\text{III.58})$$

D'un autre côté, supposons que la région X est si petite que $p(\mathbf{x})$ ne varie pas de façon considérable à l'intérieur de cette région, alors $\int_X p(\mathbf{x}) d\mathbf{x} = p(\mathbf{x})V$, où V est le volume capturé par la région X . En combinant avec (III.58), nous obtenons :

$$p(\mathbf{x}) \cong \frac{k}{\ell V}. \quad (\text{III.59})$$

L'application de ce résultat à la résolution de problèmes est faite selon deux approches différentes :

- soit on choisit une valeur fixe du volume V et on détermine la valeur de k à partir des données, donnant lieu à des méthodes connues sous le nom d'estimation de la densité par kernels (kernel density estimation, KDE),
- soit on choisit une valeur fixe de k et on détermine le volume V à partir des données, ce qui donne lieu à l'approche des k plus proches voisins (k nearest neighbor, kNN)

Il est démontré que ces deux méthodes convergent vers la densité de probabilité réelle quand $\ell \rightarrow \infty$, pourvu que V décroisse avec ℓ , et que k grandisse avec ℓ d'une façon adéquate [26].

La méthode de Parzen (estimation de la densité par kernels)

La méthode la plus populaire, la méthode de Parzen [157], utilise l'estimateur :

$$p_{\text{KDE}}(\mathbf{x}) = \frac{1}{\ell \gamma^n} \sum_{i=1}^{\ell} \phi\left(\frac{\mathbf{x} - \mathbf{x}_i}{\gamma}\right). \quad (\text{III.60})$$

où $\phi(\xi)$ est une fonction telle que $\phi(\xi) \geq 0$ et $\int \phi(\xi) d\xi = 1$ et γ est la largeur de la fenêtre. Habituellement, $\phi(\xi)$ est une PDF unimodale et radialement symétrique, comme la PDF gaussienne avec γ libre :

$$\phi_{\gamma}(\xi) = \frac{1}{(2\pi)^{n/2}} \exp\left(-\frac{\xi^T (\gamma^2 I)^{-1} \xi}{2}\right). \quad (\text{III.61})$$

On peut retrouver l'estimateur naïf à l'aide de la fonction kernel :

$$\phi(\xi) = \begin{cases} 1 & |\xi| < 1/2 \\ 0 & \text{autrement.} \end{cases} \quad (\text{III.62})$$

Le choix de la largeur de la fenêtre γ est un problème crucial. Une valeur grande de γ masquera la structure de la densité dans les données, tandis qu'une valeur faible donnera une densité très pointue, difficile à interpréter. On peut faire un choix subjectif en vérifiant graphiquement les densités résultantes, mais ce n'est pas recommandable si la dimension devient importante. Une deuxième possibilité est de faire référence à une distribution standard, où on suppose une densité standard pour trouver la valeur de γ qui minimise l'intégrale de l'erreur au carré [141] :

$$\gamma = \arg \min \left\{ E \left| \int (p_{\text{KDE}}(\mathbf{x}) - p(\mathbf{x}))^2 d\mathbf{x} \right| \right\}, \quad (\text{III.63})$$

qui, pour une densité gaussienne, est $\gamma = 1.06\sigma\ell^{-1/5}$, avec σ la variance de l'ensemble de données. Une troisième méthode est par validation croisée de la vraisemblance, où on maximise la « pseudo-vraisemblance » en utilisant la validation croisée leave-one-out [141] :

$$\gamma = \arg \max \left\{ \frac{1}{\ell} \sum_{i=1}^{\ell} \log p_{-i}(\mathbf{x}_i) \right\}, \text{ avec } p_{-i}(\mathbf{x}_i) = \frac{1}{(\ell-1)\gamma} \sum_{j=1, j \neq i}^{\ell} \phi \left(\frac{\mathbf{x}_i - \mathbf{x}_j}{\gamma} \right). \quad (\text{III.64})$$

La méthode de Parzen peut être vue comme une somme de protubérances, dont la fonction kernel détermine les formes, centrées sur les observations. La structure de l'estimateur de Parzen est complexe. Le nombre de termes en (III.60) est égal au nombre d'observations.

La méthode des k plus proches voisins

Dans cette méthode, on calcule le volume V entourant le point $\mathbf{x}$ jusqu'à ce qu'il englobe un total de k points. Donc, la densité devient :

$$p(\mathbf{x}) \cong \frac{k}{\ell V} = \frac{k}{\ell c_n d_k(\mathbf{x})}, \text{ avec } c_n = \frac{\pi^{n/2}}{(n/2)!} = \frac{\pi^{n/2}}{\Gamma(n/2 + 1)!}. \quad (\text{III.65})$$

où $d_k(\mathbf{x})$ est la distance entre le point d'estimation $\mathbf{x}$ et son $k^{\text{ème}}$ plus proche voisin $|\mathbf{x} - \mathbf{x}_k|$, c_n est le volume de la sphère unitaire n -dimensionnelle. La valeur de k contrôle le degré d'approximation, choisi pour être nettement plus petit que ℓ ; typiquement $k = \ell^{1/2}$.

La généralisation de la méthode est définie par l'introduction des fonctions kernel :

$$p(\mathbf{x}) = \frac{k}{\ell d_k(\mathbf{x})} \sum_{i=1}^{\ell} \phi \left(\frac{\mathbf{x} - \mathbf{x}_i}{d_k(\mathbf{x})} \right) \quad (\text{III.66})$$

Il est facile de vérifier que le contrôle du degré d'approximation réside dans le choix de k , mais la largeur de la fenêtre utilisée en un point particulier dépend de la densité des observations proches de ce point.

Le Tableau III.4 montre une classification des principaux algorithmes pour l'estimation de la densité, selon l'approche associée pour sa résolution.

TABLEAU III.4
Taxonomie des méthodes pour l'estimation de la densité.

	MONOVARIABLE	MULTIVARIABLE
Méthodes paramétriques	- Maximum de vraisemblance (Maximum Likelihood) - Maximisation de l'espérance (expectation maximization) - Estimation bayésienne (Bayesian estimation) - Approximation de Laplace - Approximations variationnelles - Propagation de l'espérance - Chaîne de Markov Monte Carlo	- Maximum de vraisemblance (Maximum Likelihood) - Maximisation de l'espérance (expectation maximization) - Estimation bayésienne (Bayesian estimation) - Approximation de Laplace - Approximations variationnelles - Propagation de l'espérance - Chaîne de Markov Monte Carlo
Méthodes non-paramétriques	- La méthode de Parzen (l'estimateur kernel) - Histogrammes - L'estimateur naïf (naive estimator) - Les produits kernel - La méthode du kernel variable - La méthode du plus proche voisin	- La méthode de Parzen (l'estimateur kernel) - L'estimateur naïf (naive estimator) - Les produits kernel - La méthode du kernel variable - La méthode du plus proche voisin - Estimateurs de séries orthogonales

III.2.2.4 Remarques

Les principaux problèmes que l'on peut trouver dans les machines d'apprentissage classiques sont :

1. La fonction de risque empirique associée à la machine d'apprentissage a plusieurs minimums locaux. Les procédures standard d'optimisation assurent d'atteindre l'un d'entre eux, mais la qualité de la solution dépend de plusieurs facteurs, en particulier l'initialisation des matrices de poids $\mathbf{w}$. Ce fait a amené à l'utilisation de plusieurs heuristiques pour initialiser ces matrices et, donc, trouver un meilleur minimum local.

2. La convergence des méthodes de gradient est, en général, lente. A nouveau, il existe plusieurs heuristiques pour augmenter le taux de convergence.
3. Certaines fonctions d'activation, comme les fonctions sigmoïde, ont un facteur d'échelle dont le choix est un compromis entre la qualité de l'approximation et le taux de convergence. Il existe, également, plusieurs recommandations empiriques pour effectuer ce choix.

Bien que les réseaux de neurones soient des machines d'apprentissage pas totalement « contrôlables », elles ont démontré leur efficacité dans plusieurs applications pratiques.

III.2.3 Les machines à vecteurs de support SVM

Vapnik et son équipe ont proposé les machines d'apprentissage à vecteurs de support (support vector machines, SVM) pour implémenter concrètement le principe inductif SRM [31], [155] et [156]. A partir de la définition de la machine d'apprentissage (section III.2.1.1), pour construire une machine d'apprentissage, il nous faut quatre composants principaux : un domaine, un principe d'induction, un ensemble de fonctions de décision et un algorithme que les implémente.

L'ensemble des fonctions de décision retenues pour les SVM est l'ensemble des hyperplans canoniques dans l'espace F , en d'autres termes, pour l'ensemble d'apprentissage $(\mathbf{z}_1, y_1), \dots, (\mathbf{z}_\ell, y_\ell)$, $\mathbf{z} \in F$, les couples $(\mathbf{w}, b)$ qui satisfont :

$$\min(\mathbf{w}^T \mathbf{z} + b) = 1. \quad (\text{III.67})$$

Pour cet ensemble explicite de fonctions, la dimension VC peut être bornée en respectant le principe SRM, ceci grâce au résultat suivant [151] :

Théorème 2.3.1. Soit R le rayon de la sphère la plus petite $B_R(\mathbf{a}) = \{\mathbf{z} \in F : \|\mathbf{z} - \mathbf{a}\| \leq R\}$, $\mathbf{a} \in F$, contenant les points $\mathbf{z}_1, \dots, \mathbf{z}_\ell$ et soit :

$$f_{\mathbf{w}, b} = \text{sign}(\mathbf{w}^T \mathbf{z}_i + b) \quad (\text{III.68})$$

les fonctions de décision hyperplans canoniques définies en ces points. Alors l'ensemble $\{f_{\mathbf{w}, b} : \|\mathbf{w}\| \leq A\}$ a une dimension VC h qui satisfait :

$$h < R^2 A^2 + 1. \quad (\text{III.69})$$

Comme nous préciserons plus tard, les SVM construisent des règles de décision non linéaires en effectuant une transformation non linéaire de l'ensemble d'apprentissage dans l'espace F . L'ensemble des fonctions de décision (hyperplans en F) correspond donc à l'ensemble de fonctions de décision :

$$\{\text{Error! No se pueden crear objetos modificando códigos de campo.}\} \quad (\text{III.70})$$

où la fonction $k(\cdot, \cdot)$ satisfait quelques propriétés particulières.

Les machines à vecteurs de support ont été utilisées pour la reconnaissance de formes, la régression et, plus récemment, l'estimation de densité. Premièrement, nous présentons la méthode pour construire des fonctions de décision non linéaires, pour ensuite présenter les algorithmes pour chaque domaine.

III.2.3.1 Règles de décision non linéaires

Les machines d'apprentissage qui peuvent seulement construire des solutions linéaires sont, en quelque sorte, limitées parce qu'elles décrivent, seulement, des dépendances linéaires. Les Machines à Vecteurs de Support utilisent la méthode suivante pour construire des surfaces de décision non linéaires. Les exemples d'apprentissage (III.2) dans l'espace $\mathcal{R}^n$ sont transformés non linéairement par la fonction ϕ (choisie *a priori*) dans l'espace F où les données doivent être linéairement séparables. Finalement, dans cet espace (appelé *espace des caractéristiques*), on construit un hyperplan. C'est ainsi que la fonction ϕ contrôle l'ensemble des fonctions de décision $\{f(\mathbf{x}, \boldsymbol{\alpha}), \boldsymbol{\alpha} \in \mathbf{A}\}$.

Le problème est que pour avoir des ensembles de surface de décision typiques (polynômes, fonctions à base radiale, splines), il faut des espaces de caractéristiques de dimension très importante pour pouvoir y construire des hyperplans. Par exemple, pour construire un polynôme de degré 4 dans un espace de dimension 256, l'espace F correspondant aura une dimension d'un milliard [151]. Ce problème est connu

comme le fléau de la dimension. La solution à ce problème constitue le principal objectif du développement des SVM.

La solution est la suivante. Pour n'importe quel algorithme trouvant des solutions non linéaires en F , si l'algorithme n'exécute que des produits internes en F , on peut utiliser la méthode du kernel afin d'éviter tout calcul explicite de la transformation. Donc, pour transformer les points de l'espace d'entrée $\mathcal{R}^n$ dans l'espace des caractéristiques F , en accord avec la théorie de Hilbert-Schmidt, on recherche le produit interne dans l'espace de Hilbert via des fonctions kernel :

$$k(\mathbf{x}_1, \mathbf{x}_2) \Leftrightarrow \phi(\mathbf{x}_1)^T \phi(\mathbf{x}_2) = \mathbf{z}_1^T \mathbf{z}_2 = \sum_{r=1}^{\infty} a_r \mathbf{z}_r(\mathbf{x}_1) \mathbf{z}_r(\mathbf{x}_2) \quad (\text{III.71})$$

où $\phi: \mathcal{R}^n \rightarrow F$. Ainsi, pour un certain kernel $k(.,.)$, on peut calculer le produit interne en F sans calculer la transformation ϕ .

Un exemple très connu est le kernel polynomial d'ordre d . Les produits internes dans l'espace des caractéristiques F , qui transforment les vecteurs $\mathbf{x}=[x^1, \dots, x^n]^T$ avec :

$$\phi(\mathbf{x})=[x_{i_1}, \dots, x_{i_d}], \text{ avec } 1 \leq i_1 \dots i_d \leq n, \quad (\text{III.72})$$

peuvent être calculés avec le kernel :

$$k(\mathbf{x}_1, \mathbf{x}_2) = \sum_{i_1=1}^n \dots \sum_{i_d=1}^n (x_{i_1} \dots x_{i_d})(x_{i_1} \dots x_{i_d}) \quad (\text{III.73})$$

$$= (\mathbf{x}_1^T \mathbf{x}_2)^d. \quad (\text{III.74})$$

Cependant, la forme de l'expression (III.73) reste encore très coûteuse à calculer. Pour des problèmes pratiques, le calcul du produit interne en utilisant l'expression (III.74) reste le meilleur choix. La Figure III.6 montre le cas pour $d=2$, la transformation ϕ du vecteur $\mathbf{x}$ étant :

$$[x^1, x^2]^T \xrightarrow{\phi} [(x^1)^2, (x^2)^2, x^1 x^2, x^2 x^1]^T,$$

ce qui ramène au produit interne dans l'espace des caractéristiques :

$$\begin{aligned} \phi(\mathbf{x}_1)^T \phi(\mathbf{x}_2) &= [(x_1^1)^2, (x_1^2)^2, x_1^1 x_1^2, x_1^2 x_1^1]^T [(x_2^1)^2, (x_2^2)^2, x_2^1 x_2^2, x_2^2 x_2^1] \\ &= (x_1^1)^2 (x_2^1)^2 + 2x_1^1 x_1^2 x_2^1 x_2^2 + (x_1^2)^2 (x_2^2)^2 \\ &= (\mathbf{x}_1^T \mathbf{x}_2)^2 = k(\mathbf{x}_1, \mathbf{x}_2) \end{aligned}$$

De plus, en utilisant les fonctions kernel, on prend en compte les statistiques de plus grand ordre, en l'absence d'une explosion combinatoire de la complexité que l'on aurait rencontrée pour des valeurs modérées de n et d [135].

Les fonctions kernel les plus utilisés sont les suivants :

1. Les polynômes de degré d

$$k(\mathbf{x}_1, \mathbf{x}_2) = ((\mathbf{x}_1^T \mathbf{x}_2) + 1)^d \quad (\text{III.75})$$

2. Les fonctions à base radiale

$$k(\mathbf{x}_1, \mathbf{x}_2) = \exp\left(-\frac{\|\mathbf{x}_1 - \mathbf{x}_2\|^2}{2\sigma^2}\right) \quad (\text{III.76})$$

3. Les réseaux de neurones à deux couches

$$k(\mathbf{x}_1, \mathbf{x}_2) = \tanh(\kappa(\mathbf{x}_1^T \mathbf{x}_2) + \delta) \quad (\text{III.77})$$

où κ et δ doivent vérifier l'inégalité $\kappa > \delta$.

On peut trouver plus de renseignements, exemples et discussion sur les fonctions kernel, en [37], [151] et [160].

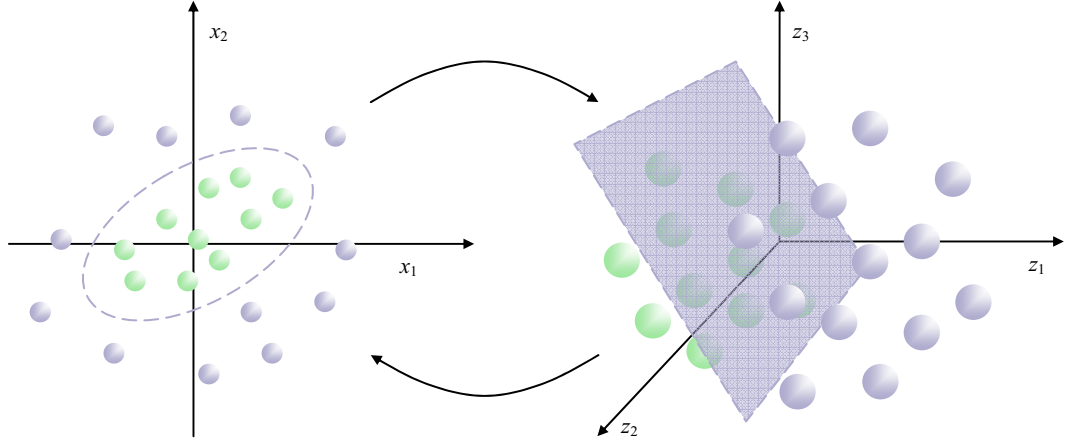


Figure III.6 : La construction d'hyperplans séparateurs dans l'espace des caractéristiques est équivalente à la construction de fonctions de décision non linéaires dans l'espace d'entrée. Donc, des données linéairement inséparables dans l'espace d'entrée peuvent être linéairement séparées dans l'espace des caractéristiques.

Ayant présenté l'ensemble des fonctions de décision à utiliser pour les SVM, il reste à décrire l'approche pour chaque domaine d'apprentissage. Chaque approche est établie de sorte que la règle de décision linéaire associée se construit dans l'espace F en n'utilisant que des produits internes. Ces produits internes peuvent être remplacés par la fonction $k(\cdot, \cdot)$ choisie.

III.2.3.2 SVM pour la reconnaissance de formes

Dans le cadre de la reconnaissance de formes, la méthode de classification à vecteurs de support (support vector classification, SVC) trouve l'hyperplan séparateur dans l'espace F qui minimise la borne (III.9), offrant la règle de décision :

$$\hat{y} = \text{sign}(\mathbf{w}^T \phi(\mathbf{x}) + b), \quad (\text{III.78})$$

où la fonction signe peut être éventuellement remplacée par une des fonctions d'activation citées dans le Tableau III.1. A partir du théorème 2.3.1, on peut minimiser le deuxième terme de la borne (III.9) avec l'ensemble des fonctions (III.78) en minimisant le terme :

$$\mathbf{w}^T \mathbf{w}. \quad (\text{III.79})$$

Vapnik a trouvé ce premier résultat en cherchant l'*hyperplan séparateur optimal* [151]. Effectivement, on peut trouver un hyperplan séparateur qui minimise (III.79). Supposons que nous n'utilisons pas de transformation ϕ pour l'instant. Afin de séparer les données (III.2), on doit séparer les contraintes :

$$\mathbf{w}^T \mathbf{x}_i + b \geq 1, \quad \text{si } y_i = 1 \quad (\text{III.80})$$

$$\mathbf{w}^T \mathbf{x}_i + b \leq -1, \quad \text{si } y_i = -1 \quad (\text{III.81})$$

qui peuvent être écrites sous une forme compacte :

$$y_i [\mathbf{w}^T \mathbf{x}_i + b] \geq 1, \quad i = 1, \dots, \ell. \quad (\text{III.82})$$

Minimiser (III.79) équivaut à trouver l'hyperplan séparateur avec la marge maximale, Figure III.7. Ceci peut se démontrer en constatant qu'à partir des contraintes (III.82) les points les plus proches de l'hyperplan satisfont $(\mathbf{w}^T \mathbf{x}^+) + b = 1$, pour $y = 1$, et $(\mathbf{w}^T \mathbf{x}^-) + b = -1$, pour $y = -1$. Maintenant, comme $(\mathbf{w}^T (\mathbf{x}^+ - \mathbf{x}^-)) = 2$, la marge est égale à $(\frac{\mathbf{w}}{\|\mathbf{w}\|}^T (\mathbf{x}^+ - \mathbf{x}^-)) = \frac{2}{\|\mathbf{w}\|}$.

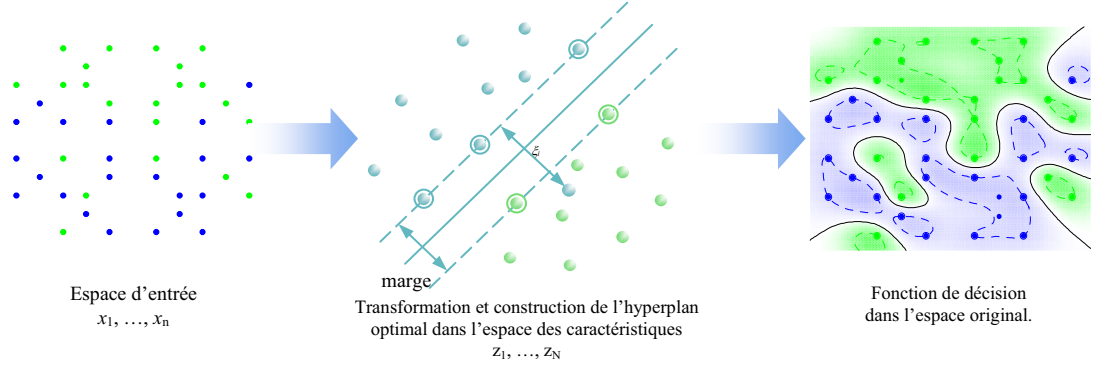


Figure III.7 : L'hyperplan séparateur optimal (avec la plus grande marge) est unique et peut être défini à l'aide des vecteurs de support (entourés). Dans le cas non linéairement séparable, le $i^{\text{ème}}$ point a une variable de relaxation associée ξ_i représentant l'amplitude de l'erreur de classification.

Pour minimiser $R_{emp}(\alpha)$, premier élément de l'équation (III.9), dans le cas général où les données ne sont pas linéairement séparables, on doit trouver la règle de décision qui donne le nombre minimal d'erreurs, en utilisant la fonction de perte (III.4). Pour cela on peut minimiser la fonction :

$$\sum_{i=1}^{\ell} \xi_i^{\sigma}, \quad (III.83)$$

avec $\sigma > 0$ suffisamment petit, sous les contraintes :

$$y_i [\mathbf{w}^T \mathbf{x}_i + b] \geq 1 - \xi_i, \quad i = 1, \dots, \ell, \quad (III.84)$$

$$\xi_i > 0. \quad (III.85)$$

Considérons le cas où $\sigma = 1$, qui est la valeur la plus petite donnant un simple problème d'optimisation (convexe)¹.

Ainsi, au lieu de minimiser la borne (III.9), les SVM, pour la reconnaissance de formes, minimisent, sous forme combinée, les termes des équations (III.79), (III.83)-(III.85), donnant le problème d'optimisation quadratique, QP :

$$\begin{aligned} &\text{Minimiser} && \frac{1}{2} \mathbf{w}^T \mathbf{w} + C \sum_{i=1}^{\ell} \xi_i \\ &\mathbf{w}, \xi, b && \\ &\text{Sous les contraintes} && y_i [\mathbf{w}^T \mathbf{x}_i + b] \geq 1 - \xi_i, \quad i = 1, \dots, \ell, \\ &&& \xi_i > 0. \end{aligned} \quad (III.86)$$

La solution à ce problème a comme arguments $\mathbf{w}$ et b et donne la règle de décision :

$$f(\mathbf{x}, \mathbf{w}, b) = \text{sign}(\mathbf{w}^T \mathbf{x} + b). \quad (III.87)$$

Le premier terme du critère maximise la marge séparatrice, tandis que le deuxième contrôle le nombre d'erreurs ou d'exemples sortants de l'ensemble d'apprentissage que nous pouvons ignorer au moment de construire une fonction de décision. Les variables de relaxation ξ_i représentent la distance du côté correct de l'hyperplan, correspondant à y_i , du $i^{\text{ème}}$ vecteur d'apprentissage. Le paramètre C détermine le compromis entre les deux termes. Il faut noter qu'un problème demeure dans le choix d'une valeur appropriée de C . Néanmoins, ce « défaut » assure, en retour, la simplicité du problème d'optimisation obtenu. Pour le cas $C = \infty$, l'algorithme ne tolère aucune erreur.

On peut trouver la solution au problème (III.86) à travers la recherche du point-col du Lagrangien :

$$L_P(\mathbf{w}, \xi, b, \alpha, \beta) = \frac{1}{2} \mathbf{w}^T \mathbf{w} + C \sum_{i=1}^{\ell} \xi_i - \sum_{i=1}^{\ell} \alpha_i (y_i [\mathbf{w}^T \mathbf{x}_i + b] - 1 + \xi_i) - \sum_{i=1}^{\ell} \beta_i \xi_i \quad (III.88)$$

où α_i et β_i sont des multiplicateurs de Lagrange. Le lagrangien doit être minimisé par rapport à $\mathbf{w}$, ξ_i et b et maximisé par rapport à $\alpha_i > 0$, $\beta_i > 0$. Sur le point-col, la solution $L_P(\mathbf{w}^*, \xi^*, b^*, \alpha^*, \beta^*)$ satisfait :

$$\frac{\partial L_P^*}{\partial \mathbf{w}} = 0, \quad \frac{\partial L_P^*}{\partial \xi_i} = 0, \quad \frac{\partial L_P^*}{\partial b} = 0 \quad (III.89)$$

ce qui mène aux conditions de stationnarité suivantes :

¹ La recherche du plus petit nombre d'erreurs d'apprentissage est un problème d'optimisation combinatoire. Nous sommes intéressé par des problèmes d'optimisation convexes garantissant un minimum global.

$$\mathbf{w}^* = \sum_{i=1}^{\ell} \alpha_i^* y_i \mathbf{x}_i, \quad (\text{III.90})$$

$$\sum_{i=1}^{\ell} \alpha_i^* y_i = 0, \quad (\text{III.91})$$

$$\alpha_i^* + \beta_i^* = C. \quad (\text{III.92})$$

Par substitution de (III.90)-(III.92) dans le Lagrangien (III.88), on obtient le problème d'optimisation :

$$\begin{array}{ll} \text{Maximiser} & L_D(\boldsymbol{\alpha}) = \sum_{i=1}^{\ell} \alpha_i - \frac{1}{2} \sum_{i,j=1}^{\ell} \alpha_i \alpha_j y_i y_j (\mathbf{x}_i^T \mathbf{x}_j) \\ \boldsymbol{\alpha} & \end{array} \quad (\text{III.93})$$

$$\begin{array}{ll} \text{Sous les contraintes} & \sum_{i=1}^{\ell} y_i \alpha_i = 0, \\ & 0 \leq \alpha_i \leq C \quad i = 1, \dots, \ell. \end{array} \quad (\text{III.94})$$

$$0 \leq \alpha_i \leq C \quad i = 1, \dots, \ell. \quad (\text{III.95})$$

En accord avec le théorème de Karush–Kuhn–Tucker, KKT, au point-col, les conditions :

$$\alpha_i^* \{y_i [\mathbf{w}^{*T} \mathbf{x}_i + b] - 1 + \xi_i^*\} = 0, \quad i = 1, \dots, \ell, \text{ et} \quad (\text{III.96})$$

$$\beta_i \xi_i = 0 \quad (\text{III.97})$$

sont satisfaites. De l'équation (III.92), on peut déduire que si $\alpha_i^* < C$ alors $\xi_i^* = 0$. Donc, de l'équation (III.96), pour n'importe quel vecteur $\mathbf{x}_i$, si $0 < \alpha_i^* < C$ (strictement entre les bornes), alors $[y_i(\mathbf{w}^{*T} \mathbf{x}) + b] - 1 = 0$ reste vrai. Ainsi, la valeur de b peut être calculée de la façon suivante :

$$b^* = \frac{1}{2} [(\mathbf{w}^{*T} \mathbf{x}^+) + (\mathbf{w}^{*T} \mathbf{x}^-)] \quad (\text{III.98})$$

où $\mathbf{x}^+$ est n'importe quel vecteur appartenant à la première classe, avec $0 < \alpha_i^+ < C$ et $\mathbf{x}^-$ est n'importe quel vecteur de la deuxième classe, avec $0 < \alpha_i^- < C$.

La fonction de décision correspondante peut être dérivée en transformant l'équation (III.87) avec (III.90) en :

$$\hat{y} = \text{sign} \left(\sum_{i=1}^{\ell} \alpha_i^* y_i (\mathbf{x}_i^T \mathbf{x}) + b^* \right). \quad (\text{III.99})$$

Les multiplicateurs de Lagrange peuvent être vues comme des « poids » ou « influences » de chaque vecteur. Pour chaque poids différent de zéro (et de par les conditions KKT, seulement quelques poids sont différents de zéro), on appelle le vecteur $\mathbf{x}_i$ *vecteur de support*, SV. Ainsi, la fonction de décision est un développement des vecteurs de support, diminuant ainsi son temps de calcul :

$$\hat{y} = \text{sign} \left(\sum_{i=1}^{nsv} \alpha_i^* y_i k(\mathbf{x}_i, \mathbf{x}) + b^* \right). \quad (\text{III.100})$$

Comme indiqué précédemment, on peut remplacer les produits internes du problème QP (III.93)-(III.95) et (III.99) par des fonctions kernel.

III.2.3.3 SVM pour la régression

Afin de résoudre des problèmes de régression, on peut suivre une procédure analogue à celle utilisée dans le cas de la reconnaissance de formes [151]. Le principe SRM est à nouveau implémenté en minimisant le terme (III.79) pour trouver la fonction de régression à partir de l'ensemble de fonctions :

$$\hat{y} = \mathbf{w}^T \phi(\mathbf{x}) + b, \quad (\text{III.101})$$

qui est choisie à partir du plus petit élément de la structure dans l'espace des caractéristiques. La construction d'une régression linéaire dans l'espace F , en utilisant que des produits internes, conditionne, à nouveau, la capacité de générer des règles de décision non linéaires à travers la transformation ϕ .

Pour le problème de régression, on peut considérer plusieurs fonctions de perte, comme indiqué dans le Tableau III.2. Si on dispose d'informations *a priori* concernant le modèle de bruit des données, on peut construire une fonction de perte spécifique, optimale pour le modèle du bruit (par exemple bruit blanc gaussien). Huber a montré que si la densité décrivant le comportement du bruit est une fonction symétrique et régulière, alors la méthode du *least modulo* fournit la meilleure fonction perte pour le modèle du pire bruit possible [156] :

$$L(y, \eta) = |y - \eta|. \quad (\text{III.102})$$

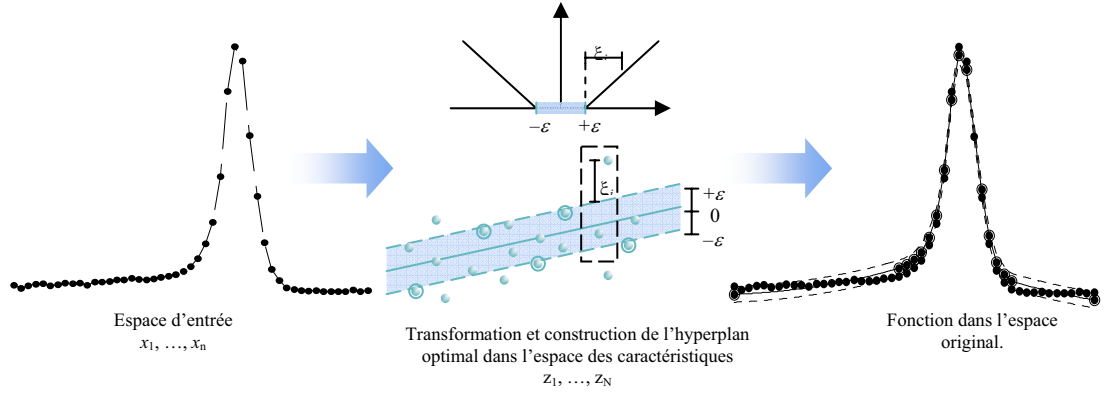


Figure III.8 : L'estimation de la régression à vecteurs de support utilisant la fonction de perte ε -insensible. Les vecteurs de support (avec $0 < \alpha_i < C$), entourés, sont dans $\pm \varepsilon$ à partir de l'hyperplan. Les autres vecteurs sont dans la zone (avec $\alpha_i = 0$) où hors zone (avec $\alpha_i = C$).

La méthode de régression à vecteurs de support (support vector regression, SVR) considère un type de fonction perte un peu plus général que (III.102), appelée fonction de perte ε -insensible :

$$|y - \eta|_{\varepsilon} = \begin{cases} 0, & \text{si } |y - \eta| \leq \varepsilon \\ |y - \eta| - \varepsilon, & \text{dans le cas contraire.} \end{cases} \quad (\text{III.103})$$

La méthode du *least modulo* est un cas particulier de la fonction de perte (III.87) quand $\varepsilon = 0$. Cette fonction de perte introduit le concept de marge dans le cas de la régression, Figure III.8.

Utilisant donc la méthode des vecteurs de support pour la régression en combinant (III.79) et (III.103), on doit résoudre le problème d'optimisation quadratique suivant :

$$\begin{aligned} &\text{Minimiser} && \frac{1}{2} \mathbf{w}^T \mathbf{w} + C \sum_{i=1}^{\ell} (\xi_i + \xi_i^*) \\ &\mathbf{w}, \xi, b && \end{aligned} \quad (\text{III.104})$$

$$\text{Sous les contraintes} \quad y_i - \mathbf{w}^T \mathbf{x}_i - b \leq \varepsilon - \xi_i^*, \quad i = 1, \dots, \ell, \quad (\text{III.105})$$

$$\mathbf{w}^T \mathbf{x}_i + b - y_i \leq \varepsilon - \xi_i, \quad i = 1, \dots, \ell, \quad (\text{III.106})$$

$$\xi > 0, \xi^* > 0. \quad (\text{III.107})$$

La solution donne la fonction de régression (III.101), sans la transformation ϕ . Pour cette dernière, dans l'espace des caractéristiques, nous utilisons la solution duale et on remplace les produits internes par la fonction $k(\cdot, \cdot)$.

Comme précédemment, à l'aide de la méthode de Lagrange, en substituant les conditions au point-col, on doit trouver les multiplicateurs de Lagrange α et α^* qui maximisent la forme quadratique :

$$\begin{aligned} &\text{Maximiser} && L_D(\alpha, \alpha^*) = -\varepsilon \sum_{i=1}^{\ell} (\alpha_i^* + \alpha_i) + \sum_{i=1}^{\ell} y_i (\alpha_i^* - \alpha_i) - \frac{1}{2} \sum_{i,j=1}^{\ell} (\alpha_i^* - \alpha_i) (\alpha_j^* - \alpha_j) (\mathbf{x}_i^T \mathbf{x}_j) \\ &\alpha, \alpha^* && \end{aligned} \quad (\text{III.108})$$

$$\text{Sous les contraintes} \quad \sum_{i=1}^{\ell} \alpha_i^* = \sum_{i=1}^{\ell} \alpha_i, \quad (\text{III.109})$$

$$0 \leq \alpha_i^* \leq C \quad i = 1, \dots, \ell, \quad (\text{III.110})$$

$$0 \leq \alpha_i \leq C \quad i = 1, \dots, \ell. \quad (\text{III.111})$$

A partir des conditions de Karush–Kuhn–Tucker, on peut observer que les points avec $\alpha_i = \alpha_i^* = 0$ demeurent à l'intérieur de la zone ε -insensible (qui est analogue à l'intérieur de la marge dans le cas de la reconnaissance de formes), et les points avec $\alpha_i = C$ ou $\alpha_i^* = C$ sont hors de la zone ε -insensible. Il faut noter que les coefficients α_i et α_i^* peuvent être actifs pour un point quelconque, selon le côté de l'hyperplan où il se trouve. Le calcul de b est établi comme suit :

$$b = y_i - (\mathbf{w}^T \mathbf{x}_i) - \varepsilon, \quad \text{si } 0 < \alpha_i < C \quad (\text{III.112})$$

$$b = y_i - (\mathbf{w}^T \mathbf{x}_i) + \varepsilon, \quad \text{si } 0 < \alpha_i^* < C \quad (\text{III.113})$$

La fonction de décision correspondante est, en conséquence :

$$f(\mathbf{x}, \boldsymbol{\beta}, b) = \sum_{i=1}^{MSV} \beta_i \mathbf{x}_i^T \mathbf{x} + b. \quad (\text{III.114})$$

où $\beta = \alpha_i^* - \alpha_i$. Une fois de plus, on substitue au produit interne la fonction kernel $k(\cdot, \cdot)$. Pour plus de détails voir [143].

III.2.3.4 SVM pour l'estimation de la densité

La méthode d'estimation de densité à vecteurs de support (support vector density estimation, SVDE) consiste à résoudre les équations linéaires :

$$Ap(\mathbf{x}) = F(\mathbf{x}), \quad (\text{III.115})$$

où A est une transformation linéaire de l'espace de Hilbert des fonctions $p(\mathbf{x})$ dans l'espace de Hilbert des fonctions $F(\mathbf{x})$. Dans l'espace des caractéristiques F , nous n'avons qu'une seule classe que nous pouvons séparer de l'origine avec une marge maximale [136], Figure III.9, permettant de paramétrer la distribution recherchée à partir de l'ensemble de fonctions :

$$f(\mathbf{x}) = \sum_{r=0}^{\infty} w_r \phi_r(\mathbf{x}) + b = \mathbf{w}^T \boldsymbol{\phi}(\mathbf{x}) + b. \quad (\text{III.116})$$

Cette fonction $f(\mathbf{x})$ procure une valeur plus grande que zéro dans une région qui capture la plupart des points d'apprentissage et une valeur négative ou égale à zéro dans le cas contraire. Ceci rompt avec le concept classique de densité de probabilité où $f(\mathbf{x}) \geq 0$, équation (III.43), mais permet d'implémenter le principe SRM en minimisant le terme (III.79) pour trouver le classifieur à partir de l'ensemble des fonctions (III.116).

Ainsi, en combinant (III.79) et (III.83), comme dans le cas de la reconnaissance de formes, on résout le problème d'optimisation quadratique suivant :

$$\begin{array}{ll} \text{Minimiser} & \frac{1}{2} \mathbf{w}^T \mathbf{w} + \frac{1}{v\ell} \sum_{i=1}^{\ell} \xi_i - b \\ \mathbf{w}, \xi, b & \end{array} \quad (\text{III.117})$$

$$\text{sous les contraintes} \quad \mathbf{w}^T \boldsymbol{\phi}(\mathbf{x}) \geq b - \xi_i, \quad (\text{III.118})$$

$$\xi_i \geq 0 \quad i = 1, \dots, \ell \quad (\text{III.119})$$

où le paramètre $v \in (0, 1)$ correspond à une borne supérieure de la fraction de vecteurs qui restent hors de la distribution apprise (outliers), ainsi qu'à une borne inférieure de la fraction de vecteurs de support ; si les données réelles sont générées de façon indépendante à partir d'une distribution $P(\mathbf{x})$, v est alors égal à ces deux fractions.

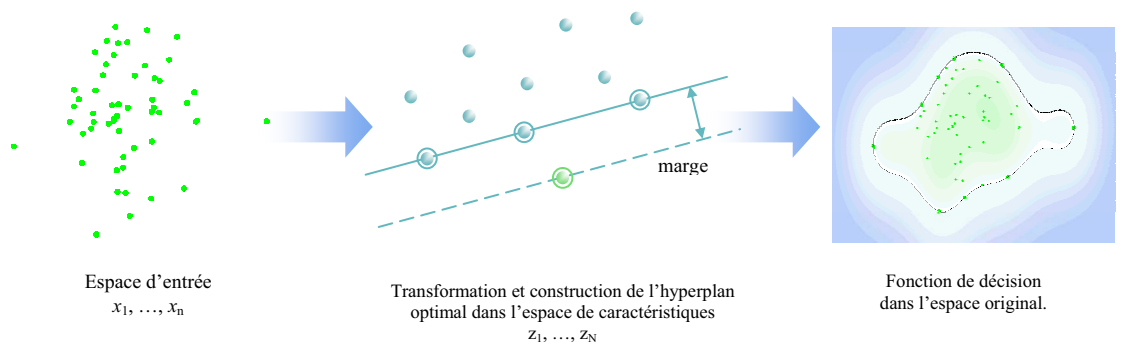


Figure III.9 : L'estimation de densité construite à partir de l'hyperplan séparateur optimal, défini à l'aide des vecteurs de support (entourés).

En utilisant, encore une fois, la méthode de Lagrange et en substituant les conditions au point-col, on doit trouver les multiplicateurs de Lagrange $\boldsymbol{\alpha}$ solution du problème quadratique :

$$\begin{array}{ll} \text{Maximiser} & L_D(\boldsymbol{\alpha}) = -\frac{1}{2} \sum_{i,j=1}^{\ell} \alpha_i \alpha_j \mathbf{x}_i^T \mathbf{x}_j \\ \boldsymbol{\alpha} & \end{array} \quad (\text{III.120})$$

$$\text{Sous les contraintes} \quad \sum_{i=1}^{\ell} \alpha_i = 1, \quad (\text{III.121})$$

$$0 \leq \alpha_i \leq \frac{1}{\nu \ell} \quad i = 1, \dots, \ell. \quad (\text{III.122})$$

La fonction de décision associée est définie comme :

$$f(\mathbf{x}) = \sum_{i=1}^{\ell} \alpha_i k(\mathbf{x}_i, \mathbf{x}) - b, \quad (\text{III.123})$$

qui sera positive pour la plupart des exemples $\mathbf{x}_i$ dans l'ensemble d'apprentissage. Les vecteurs de support correspondent, une fois de plus, aux exemples $\mathbf{x}_i$ associés aux multiplicateurs de Lagrange α_i plus grands que zéro. Géométriquement, ces vecteurs se trouvent sur la frontière de la fonction d'évaluation, Figure III.9. En analysant les conditions KKT associées au problème d'optimisation (III.120)-(III.122), on peut calculer la valeur de b en exploitant le fait que, pour chaque α_i , son correspondant $\mathbf{x}_i$ satisfait :

$$b = \sum_{j=1}^{\ell} \alpha_j k(\mathbf{x}_j, \mathbf{x}_i). \quad (\text{III.124})$$

Pour conclure cette section, il faut noter que l'on peut également construire des hypersphères pour décrire les données dans l'espace des caractéristiques [136]. Cependant, pour des kernels $k(\mathbf{u}, \mathbf{v})$ qui dépendent de $\mathbf{u}-\mathbf{v}$, la valeur de $k(\mathbf{u}, \mathbf{u})$ est constante et le problème quadratique devient équivalent à utiliser un hyperplan séparateur. On peut trouver des approches alternatives pour l'estimation de la densité dans [157] et [163].

III.2.3.5 Remarques

Dans nos recherches, nous avons constaté que la contrainte égalité associée à chacun des trois problèmes peut être enlevée quand on utilise des fonctions kernel, [65]. Pour illustrer ceci, prenons comme exemple le cas de la reconnaissance de formes, mais un raisonnement analogue peut être appliqué pour le cas de la régression et de l'estimation de la densité. Dans le passage entre le Lagrangian primal (III.88) et le Lagrangian dual (III.93) et en prenant en compte les conditions de stationnarité (III.90)-(III.92) :

$$L_D(\boldsymbol{\alpha}) = -\frac{1}{2} \sum_{i=1}^{\ell} \sum_{j=1}^{\ell} \alpha_i \alpha_j y_i y_j \mathbf{x}_i^T \mathbf{x}_j - \sum_{i=1}^{\ell} \alpha_i y_i b + \sum_{i=1}^{\ell} \alpha_i, \quad (\text{III.125})$$

nous constatons qu'il faut garder la condition (III.91) pour que le deuxième terme soit nul, car nous ne pouvons pas toujours assurer que la valeur de b soit égale à zéro.

Néanmoins, l'utilisation des fonctions kernel assure que la solution passe par l'origine de l'espace des caractéristiques de plus grande dimension où le problème est résolu, [156], faisant la valeur de b égale à zéro. Ainsi, nous pouvons enlever la contrainte égalité (III.94) du QP défini dans le processus d'apprentissage de la classification (III.93)-(III.95), ce qui permet d'utiliser des méthodes plus performantes du fait de la seule existence des contraintes de borne.

III.2.4 Le mécanisme d'apprentissage des SVM

III.2.4.1 Caractéristiques du problème d'optimisation quadratique SVM

Le processus d'apprentissage des SVM revient à résoudre le problème d'optimisation quadratique (III.93)-(III.95), dans le cas de la reconnaissance de formes, le QP définie par (III.108)-(III.111) dans le cas de la régression et le QP (III.120)-(III.122) dans le cas de l'estimation de la densité. Par des raisons de simplicité, nous n'aborderons que le cas de la reconnaissance de formes, tout en remarquant que le développement est similaire pour les trois paradigmes et que l'utilisation des fonctions kernel permet d'enlever la contrainte égalité associée. Ce résultat nous permet d'utiliser des méthodes mieux adaptées, [72].

Définissons la matrice symétrique $\mathbf{Q}$ comme $(Q)_{ij} = y_i y_j k(\mathbf{x}_i, \mathbf{x}_j)$, $i, j = 1, \dots, \ell$, en considérant le concept de kernel, et les vecteurs $\boldsymbol{\alpha} = [\alpha^1, \dots, \alpha^\ell]^T$, $\mathbf{1} = [1^1, \dots, 1^\ell]^T$, $\mathbf{y} = [y^1, \dots, y^\ell]^T$ et $\mathbf{C} = [C^1, \dots, C^\ell]^T$, le QP (III.93)-(III.95) écrit sous forme matricielle est :

$$\underset{\boldsymbol{\alpha}}{\text{Minimiser}} \quad q(\boldsymbol{\alpha}) = \frac{1}{2} \boldsymbol{\alpha}^T \mathbf{Q} \boldsymbol{\alpha} - \mathbf{1}^T \boldsymbol{\alpha} \quad (\text{III.126})$$

$$\text{sous les contraintes : } \mathbf{y}^T \mathbf{a} = 0, \quad (\text{III.127})$$

$$\mathbf{0} \leq \mathbf{a} \leq \mathbf{C} \quad (\text{III.128})$$

La matrice symétrique semi-définie positive (symmetric semi-definite positive, SSDP, matrix) $\mathbf{Q}$ assure un minimum global $\mathbf{a}^*$ du problème. Cela grâce à la (stricte) convexité de $q(\mathbf{a})$ faisant du problème (III.126)-(III.128) un problème de programmation convexe [56].

Les conditions d'optimalité

Les conditions d'optimalité sont très importantes, car elles permettent de reconnaître la solution et guident le développement des algorithmes. En général, les QP doivent satisfaire les conditions du premier et deuxième ordre mais, grâce à la convexité du problème (III.126)-(III.128), tout point solution de s conditions du premier ordre est une solution globale.

Pour trouver les conditions de stationnarité du premier ordre, on fait appel à la théorie classique de la dualité lagrangienne. En associant les multiplicateurs de Lagrange $\gamma \in \mathbb{R}$ à (III.127), $\boldsymbol{\beta} = [\beta^1, \dots, \beta^\ell]^T$ et $\boldsymbol{\chi} = [\chi^1, \dots, \chi^\ell]^T$ à (III.128), le lagrangian primal est :

$$L_p(\mathbf{a}, \boldsymbol{\beta}, \boldsymbol{\chi}, \gamma) = \frac{1}{2} \mathbf{a}^T \mathbf{Q} \mathbf{a} - \mathbf{1}^T \mathbf{a} + \gamma \mathbf{y}^T \mathbf{a} - \boldsymbol{\beta}^T \mathbf{a} + \boldsymbol{\chi}^T (\mathbf{a} - \mathbf{C}), \quad (\text{III.129})$$

qu'il faut minimiser par rapport à $\mathbf{a}$ et maximiser par rapport à γ , $\boldsymbol{\beta}$ et $\boldsymbol{\chi}$. Donc l'ensemble des conditions de stationnarité est :

$$\mathbf{y}^T \mathbf{a} = 0, \mathbf{a} \geq \mathbf{0} \text{ et } \mathbf{C} - \mathbf{a} \geq \mathbf{0} \quad (\text{faisabilité primale}) \quad (\text{III.130})$$

$$\mathbf{Q} \mathbf{a} - \mathbf{1} + \gamma \mathbf{y} - \boldsymbol{\beta} + \boldsymbol{\chi} = \mathbf{0}, \boldsymbol{\beta} \geq \mathbf{0} \text{ et } \boldsymbol{\chi} \geq \mathbf{0} \quad (\text{faisabilité duale}) \quad (\text{III.131})$$

$$\boldsymbol{\beta}^T \mathbf{a} = 0 \text{ et } \boldsymbol{\chi}^T (\mathbf{a} - \mathbf{C}) = 0 \quad (\text{souplesse complémentaire}) \quad (\text{III.132})$$

qui sont, à exception des conditions de KKT (III.132), des fonctions linéaires de $\mathbf{a}$, γ , $\boldsymbol{\beta}$ et $\boldsymbol{\chi}$. En conséquence, on peut obtenir la solution du QP (III.126)-(III.128) en trouvant une solution non négative pour l'ensemble des ℓ équations (III.130) qui satisfont également les ℓ équations (III.130)-(III.132). Nous pouvons également retrouver le problème dual du QP (III.126)-(III.128) :

$$\begin{array}{ll} \text{Maximiser} & L_D(\boldsymbol{\beta}, \boldsymbol{\chi}, \gamma) = \frac{1}{2} (\mathbf{1} - \gamma \mathbf{y} + \boldsymbol{\beta} - \boldsymbol{\chi})^T \mathbf{Q}^{-1} (\mathbf{1} - \gamma \mathbf{y} + \boldsymbol{\beta} - \boldsymbol{\chi}) - \boldsymbol{\chi}^T \mathbf{C} \end{array} \quad (\text{III.133})$$

$$\text{sous les contraintes : } \boldsymbol{\beta} \geq \mathbf{0}, \quad (\text{III.134})$$

$$\boldsymbol{\chi} \geq \mathbf{0}, \quad (\text{III.135})$$

qui a les mêmes conditions de KKT et, donc, la même solution (point col).

III.2.4.2 Les algorithmes d'optimisation des problèmes quadratiques

Il existe, essentiellement, deux types d'algorithmes de résolution des QPs :

- **Méthodes à ensemble actif** sont à la fois divisés en méthodes *primales*, qui visent la faisabilité duale en gardant l'admissibilité primale et la souplesse complémentaire, et les méthodes *duales*, qui pointent vers l'admissibilité primale en gardant l'admissibilité duale et la souplesse complémentaire. Ces dernières sont seulement applicables quand $\mathbf{Q}$ est une matrice définie positive. Elles regroupent les méthodes de programmation linéaire du simplexe dans le cas des QP, [56].
- **Méthodes de point intérieur**, qui visent la souplesse complémentaire en gardant les admissibilités primale et duale, en même temps.

Méthodes à ensemble actif

Les méthodes d'ensemble actif cherchent, après avoir trouvé un point admissible durant une phase initiale, une solution autour des bords et les faces de l'ensemble admissible en résolvant une séquence de problèmes QP avec des contraintes égalité. La seule différence entre les méthodes d'ensemble actif et la méthode du simplexe de la programmation linéaire est que ni les itérations ni la solution doivent être des sommets de l'ensemble admissible [104].

Méthodes pour des problèmes avec des contraintes égalité

La solution de base dans toutes les méthodes d'ensemble actif est celle d'un problème d'optimisation quadratique avec des contraintes égalité (equally constrained quadratic problem, EQP) :

$$\begin{aligned} &\text{Minimiser} && \frac{1}{2} \mathbf{a}^T \mathbf{Q} \mathbf{a} - \mathbf{a}^T \mathbf{g} \\ &\mathbf{a} && \end{aligned} \quad (III.136)$$

$$\text{sous les contraintes :} \quad \mathbf{A}_i^T \mathbf{a} = 0, \quad \mathbf{a} \in \mathcal{R}$$

qui suppose $\mathbf{A}$ matrice de $m \times \ell$ ($m < \ell$) de rang plein. Les conditions d'optimalité du premier ordre, avec $\boldsymbol{\mu}$ les multiplicateurs de Lagrange associés aux contraintes :

$$\begin{pmatrix} \mathbf{Q} & \mathbf{A}^T \\ \mathbf{A} & \mathbf{0} \end{pmatrix} \begin{pmatrix} \mathbf{a} \\ \boldsymbol{\mu} \end{pmatrix} = \begin{pmatrix} \mathbf{g} \\ \mathbf{0} \end{pmatrix}, \quad (III.137)$$

sont des conditions suffisantes pour résoudre (III.136), problème convexe, [104]. Pour traiter (III.137) il y a trois approches possibles : 1) l'approche de l'espace plein (full-space), 2) l'approche de l'espace borné (range-space) et 3) l'approche de l'espace nul (null-space). Pour chacune de ces trois approches, on peut utiliser des méthodes directes (factorisation de Cholesky par exemple) ou itératives (gradient conjugué). Pour plus de détails sur ces méthodes, voir [56].

Méthodes primales à ensemble actif

Dans les méthodes primales à ensemble actif, certaines contraintes inégalité indexées par un ensemble actif A , sont vues comme des contraintes égalité tandis que le reste est écarté dans l'ensemble inactif N (complément de A). La méthode ajuste cet ensemble afin d'identifier les contraintes actives correctes en la solution de (III.126)-(III.128). Chaque itération k tente de trouver la solution d'un problème EQP (formé des contraintes actives). Pour cela, l'origine change à $\boldsymbol{\alpha}^{(k)}$ et cherche une correction $\boldsymbol{\delta}^{(k)}$ qui résout :

$$\begin{aligned} &\text{Minimiser} && \frac{1}{2} \boldsymbol{\delta}^T \mathbf{Q} \boldsymbol{\delta} - \boldsymbol{\delta}^T \mathbf{g}^{(k)} \\ &\boldsymbol{\delta} && \end{aligned} \quad \text{avec } i \in A \quad (III.138)$$

$$\text{sous les contraintes :} \quad \mathbf{a}_i^T \boldsymbol{\delta} = 0,$$

avec $\mathbf{g}^{(k)} = \mathbf{1} + \mathbf{Q} \boldsymbol{\alpha}^{(k)}$, définie comme $\nabla q(\boldsymbol{\alpha}^{(k)})$ de la fonction (III.126). Si $\boldsymbol{\delta}^{(k)}$ satisfait les contraintes inactives, alors l'itération suivante est $\boldsymbol{\alpha}^{(k+1)} = \boldsymbol{\alpha}^{(k)} + \boldsymbol{\delta}^{(k)}$. Dans le cas contraire, pour trouver le meilleur point admissible, la solution de (III.138) est faite par une recherche en direction de $\boldsymbol{\delta}^{(k)}$, $\mathbf{s}^{(k)}$, avec un pas $\eta^{(k)}$, donnant le problème :

$$\eta^{(k)} = \min \left(1, \min_{\substack{i: i \notin A \\ \mathbf{a}_i^T \mathbf{s}^{(k)} < 0}} \frac{b_i - \mathbf{a}_i^T \boldsymbol{\alpha}^{(k)}}{\mathbf{a}_i^T \mathbf{s}^{(k)}} \right), \quad (III.139)$$

pour retrouver $\boldsymbol{\alpha}^{(k+1)} = \boldsymbol{\alpha}^{(k)} + \eta^{(k)} \mathbf{s}^{(k)}$. Si $\eta^{(k)} < 1$, alors la nouvelle contrainte devient active, définie par l'ensemble p des indices qui minimise (III.139), et en l'ajoutant à l'ensemble $A^{(k)}$.

Si $\boldsymbol{\alpha}^{(k)}$ (c'est-à-dire $\boldsymbol{\delta}^{(k)} = \mathbf{0}$) résout l'actuel EQP, alors il est possible de calculer les multiplicateurs de Lagrange, $\boldsymbol{\beta}$, pour les contraintes inégalité actives en utilisant une méthode de résolution des EQP. Le vecteur $\boldsymbol{\beta}^{(k)}$ doit satisfaire l'ensemble des conditions de stationnarité (III.130)-(III.132). Si les conditions de stationnarité duales, $\boldsymbol{\beta} \geq \mathbf{0}$, ne sont pas satisfaites, il faut trouver l'ensemble q des indices pour lesquels $\beta_q^{(k)} < 0$ et ensuite sélectionner q qui résout :

$$\min_{i \in A} \beta_i^{(k)}. \quad (III.140)$$

L'Algorithme III.2 synthétise la méthode primale à ensemble actif.

Algorithme III.2. Méthode primale à ensemble actif.

1. Soit le point initial $\mathbf{x}^{(1)}$ et l'ensemble actif A et $k=1$.
2. Si $\boldsymbol{\delta}^{(k)} = \mathbf{0}$ ne résout pas (III.138), alors aller au pas 4.
3. Calcul des multiplicateurs de Lagrange $\boldsymbol{\beta}^{(k)}$ et résoudre (III.140) ; si $\boldsymbol{\beta} \geq \mathbf{0}$ alors terminer avec $\boldsymbol{\alpha}^* = \boldsymbol{\alpha}^{(k)}$, autrement enlever q de l'ensemble A .
4. Résoudre (III.138) pour $\mathbf{s}^{(k)}$.
5. Trouver $\eta^{(k)} < 1$ pour résoudre (III.139) et faire $\boldsymbol{\alpha}^{(k+1)} = \boldsymbol{\alpha}^{(k)} + \eta^{(k)} \mathbf{s}^{(k)}$.
6. Si $\eta^{(k)} < 1$, ajouter p à l'ensemble A .

7. Faire $k=k+1$ et aller au pas 2.

Méthodes duales à ensemble actif

Pour notre recherche, nous utilisons un algorithme d'ensemble actif dual, proposé par Goldfarb et Idnani, [64], avec des modifications de Powell, [120]. La méthode calcule une solution sans prendre en compte les contraintes à travers une factorisation de Cholesky, pour obtenir un système triangulaire. La séquence $\alpha^{(k)}$, $\beta^{(k)}$, satisfait les conditions de KKT (III.132) à l'exception de l'admissibilité primale (III.130). Initialement, $\alpha = -Q^{-1}1$ est la minimisation du QP dual sans contraintes (III.133), $A = \emptyset$, et $\beta = 0$ est le vertex dans l'espace dual. L'itération k de la méthode consiste en les pas suivants :

Algorithme III.3. Méthode duale à ensemble actif.

1. Prendre quelques q tels que la contrainte q soit insatisfaite dans (III.126) et les ajouter à A .
2. Si $a_q^{(k)} \in \text{span}\{a_i, i \in A\}$ alors diminuer l'indice d'une contrainte inégalité qui deviendra positive dans le pas 3.
3. Aller à la solution du EQP.
4. Si $\beta_i > 0$, $i \in A$, enlever i de A et aller au pas 3.

L'itération est complétée quand la solution du EQP est trouvée au pas 3. Des itérations plus importantes continuent jusqu'à avoir une admissibilité primale au pas 1 (α est une solution optimale) ou jusqu'à l'inexistence des indices dans A (à diminuer au pas 2). Comme le QP (III.133)-(III.135) est le dual du QP (III.126)-(III.128), alors la méthode de Goldfarb et Idnani est équivalente à la méthode à ensemble actif primale, appliquée à ce problème dual, [56].

Méthodes de point intérieur

Les méthodes primales-duales, couramment connues comme méthodes de point intérieur, résolvent le problème, comme leur nom le dit, dans l'espace primal (espace des variables α) en même temps que dans l'espace dual (espace des multiplicateurs de Lagrange γ , β et χ , associés aux contraintes). En général, les algorithmes de point intérieur introduisent des variables de relaxation μ_1 , $\mu_2 \in \mathbb{R}^+$ dans les conditions de stationnarité (III.130)-(III.132) en formant l'ensemble des conditions perturbées :

$$y^T \alpha = 0, \alpha + t = C, \alpha \geq 0 \text{ et } t \geq 0 \quad (\text{admissibilité primale}) \quad (\text{III.141})$$

$$Q\alpha - 1 + \gamma y - \beta + \chi = 0, \beta \geq 0 \text{ et } \chi \geq 0 \quad (\text{admissibilité duale}) \quad (\text{III.142})$$

$$B\alpha = \mu_1 I \text{ et } X\chi = \mu_2 I \quad (\text{souplesse complémentaire}) \quad (\text{III.143})$$

où $A = \text{diag}(\alpha)$, $B = \text{diag}(\beta)$, $X = \text{diag}(\chi)$ et $T = \text{diag}(t)$. Ces algorithmes manipulent les perturbations μ_1 et μ_2 , pour qu'elles tendent vers zéro et donnent lieu à une solution approchée des conditions de stationnarité (III.130)-(III.132) avec la précision désirée. La solution est unique et réside à l'intérieur de l'espace primal-dual :

$$\alpha \geq 0, t \geq 0, \gamma, \beta \geq 0, \chi \geq 0. \quad (\text{III.144})$$

Il existe plusieurs types d'algorithmes de point intérieur, tous ayant une base mathématique commune avec la fonction barrière logarithmique, [58], mais ils peuvent se résumer en trois principaux : méthodes d'échelle affine (affine scaling), méthodes de réduction de potentiel (potential reduction) et méthodes de trajectoire centrale (central trajectory). Ces dernières sont les plus utilisées en pratique et le seul type d'algorithme de point intérieur abordé dans cette thèse.

Un algorithme de suivi de trajectoire primal-dual est composé d'un processus itératif qui commence avec un point strictement à l'intérieur de (III.144) et, à chaque itération, il estime une valeur des perturbations μ_1 et μ_2 , représentant un point de la trajectoire centrale, plus proche de la solution optimale que le point actuel, pour ensuite tenter de faire un pas vers ce point de la trajectoire centrale, [151].

Soit $(\alpha, t, \gamma, \beta, \chi)$ le point actuel et $(\alpha + \Delta\alpha, t + \Delta t, \gamma + \Delta\gamma, \beta + \Delta\beta, \chi + \Delta\chi)$ le point de la trajectoire centrale correspondant aux valeurs cibles μ_1 et μ_2 . Les équations définies pour le point de la trajectoire centrale sont :

$$y^T \Delta\alpha = -y^T \alpha \quad (\text{III.145})$$

$$\Delta\alpha + \Delta t = C - \alpha - t \quad (\text{III.146})$$

$$Q\Delta\alpha + \gamma\Delta y - \Delta\beta + \Delta\chi = 1 - Q\alpha + \gamma y + \beta - \chi \quad (\text{III.147})$$

$$\mathbf{B}^{-1}\Delta\mathbf{B} + \Delta\mathbf{A} = \boldsymbol{\mu}_1\mathbf{B}^{-1} - \mathbf{A} - \mathbf{B}^{-1}\Delta\mathbf{B}\Delta\mathbf{A} \quad (\text{III.148})$$

$$\mathbf{X}^{-1}\Delta\mathbf{X} + \Delta\mathbf{T} = \boldsymbol{\mu}_2\mathbf{X}^{-1} - \mathbf{T} - \mathbf{X}^{-1}\Delta\mathbf{X}\Delta\mathbf{T} \quad (\text{III.149})$$

Cet ensemble d'équations est presque un système linéaire pour les vecteurs de direction $(\Delta\boldsymbol{\alpha}, \Delta\mathbf{t}, \Delta\gamma, \Delta\boldsymbol{\beta}, \Delta\boldsymbol{\chi})$. Les seules non-linéarités apparaissent dans les conditions (III.148)-(III.149).

La méthode prédicteur-correcteur

Les méthodes prédicteur-correcteur sont les méthodes les plus efficaces de la programmation linéaire, [60] et [151], pour trouver une solution aux équations (III.145)-(III.149). Le pas *prédicteur* consiste à diminuer les termes $\boldsymbol{\mu}$ et les termes « delta » qui apparaissent du côté droit pour ensuite résoudre le système linéaire résultant pour les variables « delta ». Le pas *correcteur* estime une valeur cible appropriée des $\boldsymbol{\mu}$ et rétablit les termes « delta » et $\boldsymbol{\mu}$ du côté droit en utilisant les estimations actuelles et le système résultant est à nouveau résolu par rapport aux variables « delta ». Les directions résultantes sont utilisées pour aller au nouveau point dans l'espace primal-dual.

Le pas de base à l'itération k calcule la direction de Newton associée aux conditions (III.145)-(III.149) :

$$\begin{pmatrix} \mathbf{Q} & \mathbf{0} & \mathbf{y} & -\mathbf{1} & \mathbf{1} \\ \mathbf{y}^T & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{1} & \mathbf{1} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{1} & \mathbf{0} & \mathbf{0} & \mathbf{B}^{-1(k)} & \mathbf{0} \\ \mathbf{0} & \mathbf{1} & \mathbf{0} & \mathbf{0} & \mathbf{X}^{-1(k)} \end{pmatrix} \begin{pmatrix} \Delta\boldsymbol{\alpha}^{(k)} \\ \Delta\mathbf{t}^{(k)} \\ \Delta\gamma^{(k)} \\ \Delta\boldsymbol{\beta}^{(k)} \\ \Delta\boldsymbol{\chi}^{(k)} \end{pmatrix} = \begin{pmatrix} \mathbf{1} - \mathbf{Q}\boldsymbol{\alpha}^{(k)} + \boldsymbol{\beta}^{(k)} - \boldsymbol{\chi}^{(k)} \\ -\mathbf{y}^T\boldsymbol{\alpha}^{(k)} \\ \mathbf{C} - \boldsymbol{\alpha}^{(k)} - \mathbf{t}^{(k)} \\ \boldsymbol{\mu}_1\mathbf{B}^{-1(k)} - \mathbf{A}^{(k)} - \mathbf{B}^{-1(k)}\Delta\mathbf{B}^{(k)}\Delta\mathbf{A}^{(k)} \\ \boldsymbol{\mu}_2\mathbf{T}^{-1(k)} - \mathbf{T}^{(k)} - \mathbf{X}^{-1(k)}\Delta\mathbf{X}^{(k)}\Delta\mathbf{T}^{(k)} \end{pmatrix}, \quad (\text{III.150})$$

où $\boldsymbol{\alpha}^{(k)}$, et $\mathbf{t}^{(k)}$, indiquent le produit scalaire entre deux vecteurs (élément par élément). Ce système est réduit à un système quasi-défini et il est résolu en utilisant des factorisations de Cholesky. Il n'est pas garanti que $(\boldsymbol{\alpha}^{(k)} + \Delta\boldsymbol{\alpha}^{(k+1)}, \mathbf{t}^{(k)} + \Delta\mathbf{t}^{(k+1)}, \gamma^{(k)} + \Delta\gamma^{(k+1)}, \boldsymbol{\beta}^{(k)} + \Delta\boldsymbol{\beta}^{(k+1)}, \boldsymbol{\chi}^{(k)} + \Delta\boldsymbol{\chi}^{(k+1)}) > 0$. Donc, on établit à la place :

$$(\boldsymbol{\alpha}, \mathbf{t}, \gamma, \boldsymbol{\beta}, \boldsymbol{\chi})^{(k+1)} = (\boldsymbol{\alpha}, \mathbf{t}, \gamma, \boldsymbol{\beta}, \boldsymbol{\chi})^{(k)} + \nu((\Delta\boldsymbol{\alpha}, \Delta\mathbf{t}, \Delta\gamma, \Delta\boldsymbol{\beta}, \Delta\boldsymbol{\chi})^{(k+1)}) \quad (\text{III.151})$$

où $\nu \leq 1$ assure $(\boldsymbol{\alpha}, \mathbf{t}, \gamma, \boldsymbol{\beta}, \boldsymbol{\chi})^{(k)} + \nu(\Delta\boldsymbol{\alpha}, \Delta\mathbf{t}, \Delta\gamma, \Delta\boldsymbol{\beta}, \Delta\boldsymbol{\chi})^{(k+1)} > 0$. L'algorithme peut être paraphrasé comme suit :

Algorithme III.4. Méthode de point intérieur prédicteur-correcteur.

1. Vérifier l'optimalité avec les tolérances spécifiées.
 2. Etablir les nouvelles valeurs cible de $\boldsymbol{\mu}$.
 3. Résoudre le système de Newton (III.150) avec ces cibles.
 4. Réaliser un pas de Newton (III.151) dans la direction de Newton avec, éventuellement, des pas de largeur différente dans les espaces primal et dual.
-

Pour plus de renseignements sur les méthodes de point intérieur, voir [5], [6], [142] et [152].

III.2.4.3 Implémentation des SVM

Les algorithmes d'optimisation pour résoudre le problème d'optimisation quadratique (III.126)-(III.128) décrits dans la section précédente, sont opérationnels pour résoudre des problèmes avec des bases de données de moins de 2000 exemples. Au-delà de cette borne, dépendant des capacités de calcul et de mémoire, nous ne pouvons utiliser n'importe quelle technique d'optimisation quadratique sans quelques modifications. Dans des bases de données de grande taille, il est difficile de calculer et de stocker la matrice $\mathbf{Q}$ des produits vectoriels $k(\mathbf{x}_i, \mathbf{x}_j)$, à cause des limitations informatiques. En conséquence, on doit trouver des méthodes plus efficaces pour ce type de bases de données et pouvoir arriver à la solution optimale en un temps minimal, avec une demande de ressources informatiques modérée.

Décomposition du QP

Plusieurs auteurs, [111], [118], [155], ont introduit l'idée de décomposer le problème en des petits sous problèmes, plus faciles à traiter. Les stratégies considèrent deux points clés :

- les conditions d'optimalité, qui permettent de vérifier si l'algorithme a optimalement résolu le problème, et

- *la stratégie d'implémentation*, qui définit la façon d'implémenter la fonction objectif, associée aux variables qui violent les conditions d'optimalité, si une solution particulière n'est pas la solution globale.

La solution du problème (III.126)-(III.128) est optimale si et seulement si les conditions de KKT (III.132) sont satisfaites, sachant que la matrice $\mathbf{Q}$ est semi définie positive. Les conditions de KKT ont une forme simple à vérifier ; le QP peut être résolu quand, $\forall i, i=1, \dots, \ell$:

$$\alpha_i=0 \rightarrow y_i g(\mathbf{x}_i) > 1 \quad (\text{III.152})$$

$$0 < \alpha_i < C \rightarrow y_i g(\mathbf{x}_i) = 1 \quad (\text{III.153})$$

$$\alpha_i = C \rightarrow y_i g(\mathbf{x}_i) < 1 \quad (\text{III.154})$$

avec $g(\mathbf{x}_i)$ fonction de décision (III.99) à fonction d'activation identité, définie par, [65], [111] :

$$g(\mathbf{x}_i) = \sum_{j=1}^{\ell} \alpha_j y_j k(\mathbf{x}_j, \mathbf{x}_i) + b \quad (\text{III.155})$$

Afin d'incorporer les conditions d'optimalité, la stratégie d'implémentation doit prendre en compte le fait qu'une partie importante des multiplicateurs de Lagrange α_i sont égaux à zéro à la solution. D'une façon analogue aux méthodes d'optimisation à ensemble actif, une solution est de découper l'ensemble d'apprentissage en un ensemble actif A , également connu sur le nom d'ensemble de travail (working set), est son complément N . Alors, on peut récrire le QP (III.126)-(III.128) comme suit :

$$\begin{aligned} \text{Minimiser} \quad & q(\mathbf{a}_A, \mathbf{a}_N) = \frac{1}{2} \begin{bmatrix} \mathbf{a}_A \\ \mathbf{a}_N \end{bmatrix}^T \begin{bmatrix} \mathbf{Q}_{AA} & \mathbf{Q}_{AN} \\ \mathbf{Q}_{NA} & \mathbf{Q}_{NN} \end{bmatrix} \begin{bmatrix} \mathbf{a}_A \\ \mathbf{a}_N \end{bmatrix} - \begin{bmatrix} \mathbf{1}_A \\ \mathbf{1}_N \end{bmatrix}^T \begin{bmatrix} \mathbf{a}_A \\ \mathbf{a}_N \end{bmatrix} \end{aligned} \quad (\text{III.156})$$

$$\text{sous les contraintes :} \quad \begin{bmatrix} \mathbf{y}_A \\ \mathbf{y}_N \end{bmatrix}^T \begin{bmatrix} \mathbf{a}_A \\ \mathbf{a}_N \end{bmatrix} = 0, \quad (\text{III.157})$$

$$\begin{bmatrix} \mathbf{0}_A \\ \mathbf{0}_N \end{bmatrix} \leq \begin{bmatrix} \mathbf{a}_A \\ \mathbf{a}_N \end{bmatrix} \leq \begin{bmatrix} \mathbf{C}_A \\ \mathbf{C}_N \end{bmatrix}, \quad (\text{III.158})$$

dans lequel nous pouvons remplacer n'importe quel $i \in A$, par n'importe quel $j \in N$, sans modifier la fonction de coût. L'idée est d'avoir dans l'ensemble actif A tous les vecteurs de support. De cette manière, l'ensemble inactif N n'est formé que par des multiplicateurs α_N nuls, ce qui nous mène au QP :

$$\begin{aligned} \text{Minimiser} \quad & q(\mathbf{a}_A) = \frac{1}{2} \mathbf{a}_A^T \mathbf{Q}_A \mathbf{a}_A - \mathbf{1}^T \mathbf{a}_A \end{aligned} \quad (\text{III.159})$$

$$\text{sous les contraintes :} \quad \mathbf{y}_A^T \mathbf{a}_A = 0, \quad (\text{III.160})$$

$$\mathbf{0}_A \leq \mathbf{a}_A \leq \mathbf{C}_A. \quad (\text{III.161})$$

La solution de ce problème sera la solution du QP (III.126)-(III.128), si elle vérifie les conditions d'optimalité (III.152)-(III.154), en particulier $y_j g(\mathbf{x}_j) > 1$, $j \in N$ (condition pour laquelle $\alpha_j = 0$). Si ce n'est pas le cas, alors la valeur de α_j , correspondante au vecteur $\mathbf{x}_j$, doit être différente de zéro et, par conséquence, il est nécessaire de le déplacer vers l'ensemble actif A . Evidemment, nous pouvons faire de même pour les vecteurs $\mathbf{x}_i$, associés aux $\alpha_i = 0$, vers l'ensemble N . L'Algorithme III.5 définit la méthode générale de décomposition du problème quadratique SVM.

Algorithme III.5. Algorithme de décomposition du problème quadratique SVM.

-
1. Election d'un ensemble actif initial A de taille n_A .
 2. Résoudre le QP (II.183)-(III.185), défini par l'ensemble actif A .
 3. Tant qu'il existe des $j \in N$ sans satisfaire $y_j g(\mathbf{x}_j) > 1$,
 - a. déplacer les n_A vecteurs $\mathbf{x}_j$ les plus erronées vers l'ensemble A ,
 - b. déplacer tous les vecteurs $\mathbf{x}_i$ avec $\alpha_i = 0$, $i \in A$, vers l'ensemble N , et retourner au pas 2
-

Les variations autour de l'algorithme de décomposition

Vapnik, [155], a été le premier à proposer une méthode de décomposition pour résoudre le QP, connue depuis sous le nom de « chunking ». A chaque itération, l'ensemble actif A est défini par les α_A différents de zéro, et les n_A exemples les plus erronées, de l'ensemble d'apprentissage entier. Osuna, [110], suggère une taille constante de l'ensemble actif A , ce qui implique enlever et ajouter le même nombre d'exemples à chaque itération. Joachims, [91], a introduit une série de conseils pour traiter des

problèmes de grande échelle, dont la définition du meilleur ensemble d'apprentissage, basée dans la méthode de Zoutendijk, ce qui revient à prendre les exemples les plus erronés.

Une forme de décomposition extrême est l'optimisation minimale séquentielle (sequential minimal optimization, SMO), proposée par Platt, [118]. Cette approche résout une série des QP de taille deux, le QP le plus petit possible dû à la contrainte égalité, de façon analytique. A chaque pas, l'algorithme choisit deux multiplicateurs de Lagrange, trouve leurs valeurs optimales et actualise la SVM. L'algorithme a trois composantes principales : une méthode analytique pour résoudre les deux multiplicateurs de Lagrange, une heuristique pour choisir les multiplicateurs à optimiser et une méthode pour calculer la valeur de b .

III.2.4.4 Application des SVM sur des bases de données

Dans cette section, nous présentons l'application de trois SVM-QP : l'algorithme à ensemble actif dual (QP_1), un algorithme de point intérieur (QP_2) et un deuxième algorithme de point intérieur appelé proloquo (QP_3), très utilisé dans la littérature SVM, [142]. Pour chaque cas, nous avons également implémenté l'algorithme de décomposition (QP_i+Dec).

Pour montrer la performance des implémentations des SVM, nous avons restreint l'étude à des problèmes de classification ayant des bases de données à deux classes. Pour cela, nous avons utilisé huit bases de données largement connues dans la littérature, donc bien adaptées à notre propos de comparaison. Le Tableau III.5 résume les caractéristiques principales des bases de données utilisées. Toutes les bases sont disponibles dans [80].

Les bases de données du cancer de sein (par la suite nommées par wbc) ont été obtenues dans le centre hospitalier universitaire de Wisconsin. La fondation clinique de Cleveland a récolté les données traitant le diagnostic des maladies du cœur (hea). Le test sonar (son) correspond aux données issues d'une étude de classification des signaux pour discriminer les signaux réfléchis par un cylindre métallique et ceux réfléchis par un cylindre rocheux. Le bureau de recensement aux Etats Unis a collecté la base de données « adult » (adu) en 1994 afin de déterminer si une personne obtient des revenus supérieurs à 50kUSD. L'institut médical BUPA, à Nottingham, a collecté des analyses sanguines pour former la base de données « désordres du foie » (bld), suite à la consommation excessive d'alcool. Les iris (iri), la base de données la plus populaire dans la littérature de la reconnaissance de formes, contient trois types de plantes iris. Enfin, une analyse chimique de plusieurs vins italiens provenant de trois régions différentes (win) a été considérée.

TABLEAU III.5

Caractéristiques des bases de données UCI. N_{ler} est le nombre d'individus de l'ensemble d'apprentissage, N_{test} le nombre d'observations de l'ensemble de test et N est la taille de la base de données. n_{num} et n_{cat} dénotent le nombre d'attributs numériques et catégoriels respectivement, n est la dimension d'entrée.

	N_{ler}	N_{test}	N	n_{num}	n_{cat}	n
wbc	455	228	683	9	0	9
bld	230	115	345	6	0	5
iri	100	50	150	4	0	4
hea	1000	541	1541	7	6	13
adu	5000	40222	45222	6	8	14
wine	125	53	178	13	0	13
glass	150	64	214	10	0	10
sonar	140	68	208	60	0	60

La performance des algorithmes en termes de précision de séparation des classes est la même pour toutes les implémentations (solution globale du problème quadratique). Nous avons réalisé tous les tests sur des stations de travail SunBlade₁₀₀. Concernant les fonctions kernel utilisées, nous avons employé le

kernel polynomial et le kernel gaussien. Nous avons lancé une optimisation du paramètre de régularisation C et du paramètre du kernel σ , en utilisant une validation croisée. Comme processus de prétraitement, nous avons enlevé tout exemple contenant des valeurs inconnues et normalisé chaque attribut sur le domaine $[-1,1]$. Le Tableau III.6 reporte les paramètres pour chaque base de données et le taux de bonne classification pour la base d'apprentissage et la base de test.

TABLEAU III.6

Paramètres SVM optimisés pour les bases de données UCI. $kernel$ est le type de kernel, σ son paramètre correspondant, le C paramètre de régularisation, pour la performance d'apprentissage P_{ler} et de test P_{test} . n_{sv} (%) indique le nombre de vecteurs de support trouvés et le pourcentage par rapport à l'ensemble d'apprentissage

	$kernel$	σ	C	P_L	P_{Test}	n_{sv} (%)
wbc	Gaussian	1.4	100	1.0	0.9780	64 (14.065 %)
bld	Gaussian	0.22	840	0.9085	-	192 (55.65%)
iri	Polynomial	2	1120	0.9876	-	9 (6%)
hea	Gaussian	0.34	150	0.9086	-	660 (42.83%)
adu	Polynomial	3	70	0.9632	0.9213	936 (23.4%)
wine	linear	-	100	1.0	0.9732	11 (8.8%)
glass	Polynomial	5	100	1.0	0.9532	23 (15.33%)
sonar	Polynomial	2	0.01	0.9863	0.8592	84 (57.53%)

Le Tableau III.7 affiche les résultats en terme de temps pour les trois implémentations SVM-QP sur les bases de données. D'un côté, l'implémentation de la méthode à ensemble actif dual, QP_1 , se montre plus performante que les autres implémentations. Les implémentations QP_1 et QP_2 ont de meilleurs résultats que le QP_3 , même en utilisant l'algorithme de décomposition. De l'autre côté, l'algorithme de décomposition diminue substantiellement le temps d'apprentissage des trois implémentations.

TABLEAU III.7

Temps de calcul pour les implémentations SVM, utilisant les bases des données UCI : QP_i dénotent les implémentations SVM-QP et QP_i+Dec utilisent l'algorithme décomposition.

	QP_1	QP_2	QP_3	QP_1+Dec	QP_2+Dec	QP_3+Dec
wbc	16.656 sec	71.747 sec	184.816 sec	1.02 sec	1.102 sec	2.625 sec
bld	3.097 sec	5.657 sec	49.095 sec	1.558 sec	9.183 sec	38.843 sec
iri	1.065 sec	19.649 sec	4.259 sec	0.018 sec	0.038 sec	0.082 sec
hea	-	-	-	495.95 sec	1284.21 sec	4646.65 sec
adu	-	-	-	1960.59 sec	5830.2 sec	24951.9 sec
wine	0.077 sec	0.482 sec	3.987 sec	0.024 sec	0.027 sec	0.196 sec
glass	0.421 sec	3.887 sec	8.349 sec	0.038 sec	0.114 sec	0.528 sec
sonar	0.375 sec	1.701 sec	5.717 sec	0.379 sec	1.95 sec	3.237 sec

III.3 Modélisation à partir de la connaissance

III.3.1 Généralités sur l'intelligence artificielle

Pendant des milliers d'années, l'homme a essayé de *comprendre* comment nous pensons, c'est-à-dire, comment un simple individu peut percevoir, comprendre, prédire et agir sur un monde beaucoup plus vaste et compliqué que lui même. L'intelligence artificielle, AI, tente non seulement de comprendre mais également de *construire* des entités « intelligentes ».

Dans la littérature, il existe plusieurs définitions de l'intelligence artificielle. Le Tableau III.8 collecte huit définitions [132]. Ces définitions varient selon deux grandes dimensions. Les définitions de gauche mesurent le succès en termes de fidélité à la *performance humaine* tandis que celles de droite les mesurent vis-à-vis d'un concept *idéal* d'intelligence, connu sous le nom de *rationalité*. Un système est rationnel s'il exécute « l'action correcte », sur la base de ce qu'il « sait ».

TABLEAU III.8
Quelques définitions de l'intelligence artificielle, organisées en quatre catégories.

SYSTEMES QUI ESSAIENT DE PENSER COMME LES HUMAINS	SYSTEMES QUI PENSENT RATIONNELLEMENT
« Le nouvel effort excitant pour faire penser les ordinateurs... <i>machines avec du cerveau</i> , dans le sens littéral et complet »	« L'étude des facultés mentales à travers des modèles implémentés sur ordinateur »
« L'automatisation des activités que l'on associe à la pensée humaine, activités comme la prise de décision, la résolution de problèmes, l'apprentissage, etc. »	« L'étude des manipulations par ordinateur qui rendent possible la perception, le raisonnement et l'action »
SYSTEMES QUI AGISSENT COMME LES HUMAINS	SYSTEMES QUI AGISSENT RATIONNELLEMENT
« L'art de créer des machines qui exécutent des fonctions qui requièrent de l'intelligence quand les personnes le font »	« L'art de concevoir des agents intelligents »
« L'étude du « comment faire » pour que les ordinateurs fassent des choses que les hommes, en ce moment, font mieux »	« AI ... est concernée par le comportement intelligent des artefacts »

Dans notre étude, nous adoptons le point de vue où « l'intelligence » est assimilée principalement à des *actions rationnelles*. Idéalement, un agent intelligent prend la meilleure action dans une situation particulière. Nous allons aborder le problème de la construction d'*agents* qui sont « intelligents » dans ce sens.

III.3.1.1 Agents intelligents

Un agent est une entité qui peut percevoir son environnement à travers des capteurs ou d'autres agents et réagir sur celui-ci à travers des actionneurs. Nous utilisons le terme *perception* pour noter les entrées perçues par l'agent à un instant donné. Ainsi, une *séquence de perception* est tout ce que l'agent a reçu dans un intervalle de temps. En général, l'action d'un agent à un instant donné peut dépendre de toute la séquence perçue jusqu'à cet instant. Mathématiquement parlant, le comportement d'un agent est décrit par une fonction $f(x)$ qui transforme une séquence de perception en une action.

Le premier pas dans le processus de conception d'un agent rationnel est de spécifier le plus complètement possible la tâche qu'il doit accomplir. Pour cela, il faut identifier la nature de la tâche, c'est-à-dire vérifier si elle peut être complètement ou partiellement observable, déterministe ou stochastique, épisodique ou séquentielle, statique ou dynamique, discrète ou continue, et mono-agent ou multi-agent.

III.3.1.2 Systèmes Experts

Les systèmes experts sont l'application la plus commerciale de l'intelligence artificielle. Ce type de système utilise une base de connaissances correspondant à une « expertise humaine » pour résoudre des problèmes difficiles. Le degré de résolution du problème est basé sur la qualité des données et des règles obtenues de l'expert humain. Les systèmes experts sont conçus pour prendre des décisions (ou proposer des décisions) au même niveau que l'expert humain. Dans la pratique, leur performance est bien au-dessous de celle d'un individu expert, sauf en ce qui concerne la rapidité du raisonnement des machines artificielles.

Le système expert élabore ses réponses en consultant la base de connaissances à travers un moteur d'inférence, logiciel qui interagit avec l'utilisateur et calcule le résultat à partir des règles et données de la base de connaissances.

Les systèmes experts sont utilisés dans des applications telles que le diagnostic médical, la maintenance de matériel, l'analyse de l'investissement, la gestion des assurances, la planification de mouvements des véhicules de livraison, la vérification des contrats, etc. [132].

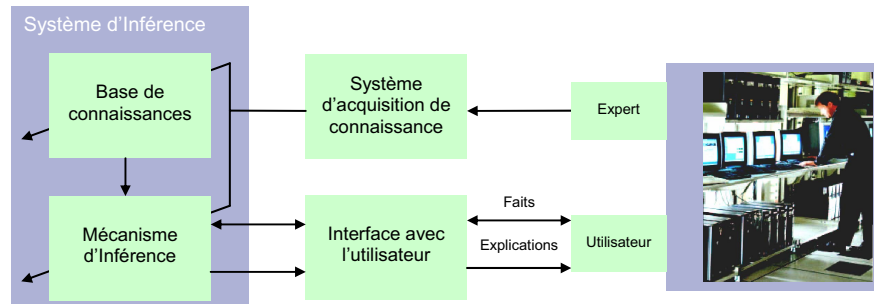


Figure III.10 : Schéma général de l'architecture d'un système expert.

Un système expert est composé de trois éléments principaux, Figure III.10: Un moteur d'inférence, un système d'acquisition de connaissance et une interface avec l'utilisateur. Si le premier composant peut être bien maîtrisé techniquement parlant, ceci est moins vrai pour les deux autres composants.

Plusieurs travaux conduisent à prendre en compte quatre recommandations lors de la conception d'un système expert [32] :

1. La planification des expériences pour acquérir la connaissance doit être sérieuse. La participation de l'expertise humaine ne peut pas être assurée *a priori*. Les experts peuvent être très occupés et facilement indisponibles.
2. Les experts doivent être rapidement impliqués dans le processus de construction d'un système expert. Il faut leur montrer des résultats le plus tôt possible. Une première version du système maintient l'intérêt des experts, d'où la nécessité d'un prototype rapide.
3. Les experts doivent énoncer, dans la mesure du possible, tous les pas de leur raisonnement. Le fait d'écrire toutes ces étapes est un point essentiel pour comprendre la façon dont les experts manipulent leur connaissance.
4. Le problème doit être suffisamment (mais pas trop) difficile pour maintenir un certain degré d'intérêt de la part des experts. Cependant, il doit être limité à un nombre d'heures de travail raisonnable.

Dans cette thèse, nous nous centrons principalement sur le développement du système d'inférence, puisqu'il constitue le principal support de la technologie des systèmes experts, tout particulièrement les systèmes d'inférence flou.

III.3.2 Les systèmes d'inférence flous

Les systèmes d'inférence flous (fuzzy inference systems, FIS) sont une des applications les plus célèbres de la théorie de la logique floue et des ensembles flous, proposés par L., Zadeh dès 1965 [170]. Le succès des FIS réside dans leur deux caractéristiques principales suivantes. D'un côté, ils sont capables de manipuler des concepts linguistiques. De l'autre, ce sont des approximateurs universels aptes à reproduire des transformations non linéaires entre les entrées et les sorties.

III.3.3.1 Éléments d'un système d'inférence flou

Dans les systèmes d'inférence flous, la base de connaissance est composée d'un ensemble de règles d'inférence *si-alors* de la forme :

si *prémisse* **alors** *conséquence*.

La prémisse (**si**) définit de manière imprécise les états du système, tandis que la partie conséquente (**alors**) représente les actions qui peuvent être prises dans cette situation. Un système d'inférence flou est composé de quatre éléments :

1. Une *base de connaissance* qui contient les définitions des ensembles flous et les opérateurs flous ;
2. Un *mécanisme d'inférence* qui exécute tous les calculs ;

3. Un *fuzzificateur* qui transforme les entrées réelles en variables qualitatives définies par des ensembles flous ;
4. Un *défuzzificateur* qui transforme l'ensemble flou de sortie en une valeur réelle.

La Figure III.11 montre la structure de base d'un système d'inférence flou, dans lequel les entrées et les sorties sont habituellement des variables réelles. Le point fort de la logique floue est la notion d'ensemble flou, car elle donne à l'expert une méthode pour définir une représentation précise des termes vagues du langage naturel comme chaud, froid, tiède, etc. La logique floue procure les éléments nécessaires pour manipuler et inférer des conclusions basées sur cette information vague, [170].

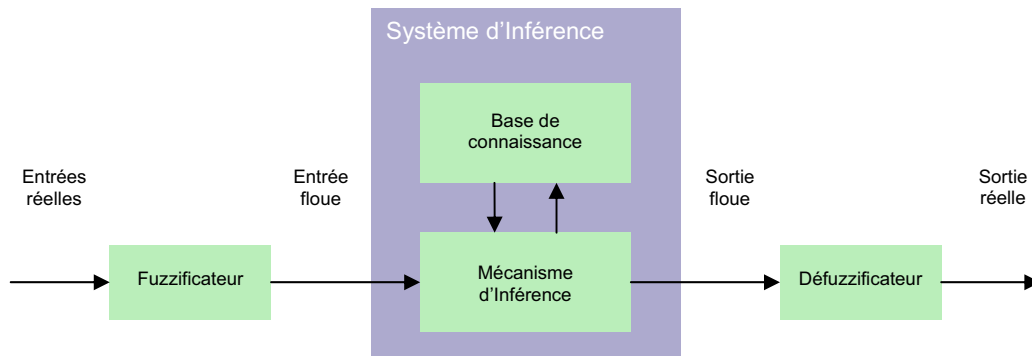


Figure III.11 : Composants basiques d'un système d'inférence flou.

III.3.3.2 Généralités sur la logique floue

Les ensembles flous et les fonctions d'appartenance

De même qu'un ensemble classique est défini par sa fonction caractéristique, un ensemble flou A , définie sur l'univers du discours (ou domaine) X , est représenté par sa *fonction d'appartenance* $\mu(x) \in [0,1]$, qui indique le degré avec lequel un événement peut être classifié comme quelque chose. Les fonctions d'appartenance fournissent l'interface entre l'espace (d'entrée) à valeurs réelles et les ensembles linguistiques de l'expert. C'est dans la précision de la représentation de ces concepts vagues que la connaissance spécifique de l'expert sur une situation particulière se trouve. Le Tableau III.9 montre les fonctions d'appartenance les plus utilisées : triangulaire, trapézoïdal, gaussienne, mais un nombre réel est un cas spécial d'ensemble flou (ensemble singleton). La Figure III.12 illustre ces fonctions. Les fonctions d'appartenance gaussiennes sont très populaires, car elles constituent la base de la connexion entre les systèmes d'inférence flou et les réseaux à fonctions de base radiale (RBFN).

TABLEAU III.9

Les fonctions d'appartenance les plus utilisées.

	FONCTION D'APPARTENANCE	RESTRICTIONS
Singleton	$\mu(x) = \begin{cases} 1 & \text{si } x = p_1 \\ 0 & \text{autrement} \end{cases}$	—
Triangulaire	$\mu(x) = \begin{cases} (x - p_1)/(p_2 - p_1) & \text{si } p_1 \leq x < p_2 \\ (p_3 - x)/(p_3 - p_2) & \text{si } p_2 \leq x < p_3 \\ 0 & \text{autrement} \end{cases}$	$p_1 < p_2 < p_3$
Trapézoïdale	$\mu(x) = \begin{cases} (x - p_1)/(p_2 - p_1) & \text{si } p_1 \leq x < p_2 \\ 1 & \text{si } p_2 \leq x < p_3 \\ (p_4 - x)/(p_4 - p_3) & \text{si } p_3 \leq x < p_4 \\ 0 & \text{autrement} \end{cases}$	$p_1 < p_2 < p_3 < p_4$
Gaussienne	$\mu(x) = \exp\left(-\frac{(x - p_1)^2}{p_2^2}\right)$	—

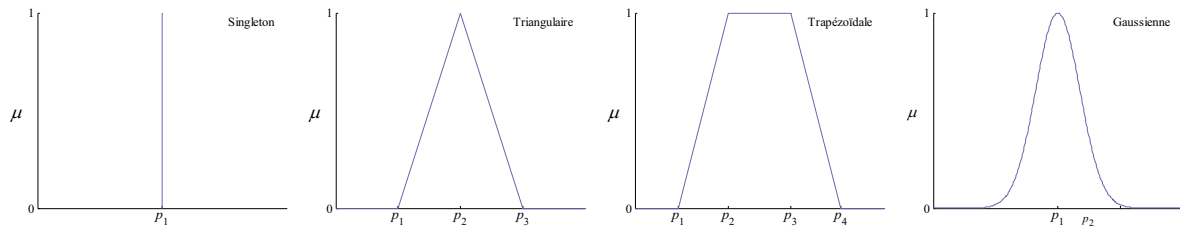


Figure III.12 : Fonctions d'appartenance : singleton, triangulaire, trapézoïdale, gaussienne.

Les variables et les valeurs linguistiques

Les propositions floues sont des expressions comme « la température est grande », où « grande » est une *valeur linguistique*, définie par un ensemble flou A sur l'univers de discours X de la *variable linguistique* « température ». Les valeurs linguistiques sont aussi désignées constantes floues, termes flous ou notions floues. Les modificateurs linguistiques peuvent modifier la signification des valeurs linguistiques [130]. Par exemple, le modificateur linguistique *plutôt* peut changer « la température est grande » en « la température est *plutôt* grande ». Le Tableau III.10 montre quelques exemples de modificateurs linguistiques.

La Figure III.13 expose un exemple de variable linguistique « température » avec trois termes linguistiques : « petite », « moyenne » et « grande » représentés par des trapézoïdes et trois exemples de modificateurs linguistiques.

TABLEAU III.10
Quelques exemples de modificateurs linguistiques.

OPERATION		OPERATION	
très A (concentration)	μ_A^2	plus ou moins A (dilatation)	$\sqrt{\mu_A}$
pas très A	$1 - \mu_A^2$	plutôt A (intensification)	$\text{int}(\mu_A) = \begin{cases} 2\mu_A^2 & \text{si } \mu_A \leq 0.5 \\ 1 - 2(1 - \mu_A)^2 & \text{autrement} \end{cases}$

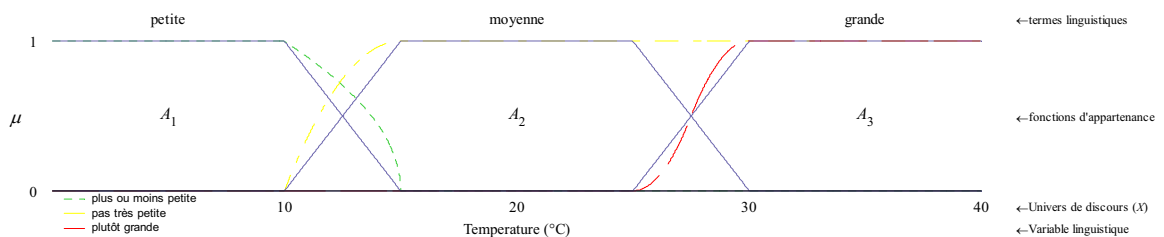


Figure III.13 : Modificateurs linguistiques : plus ou moins petit, pas très petite, plutôt grande.

Opérations entre ensembles flous

Les opérations définies dans la théorie classique des ensembles comme le complément, l'union et l'intersection ont une extension vers la théorie des ensembles flous. Comme les degrés d'appartenance ne sont plus restreints à $\{0,1\}$, mais à n'importe quelle valeur dans l'intervalle $[0,1]$, il existe plusieurs généralisations pour les opérateurs logiques flous, [75].

Le complément d'un ensemble flou

Le complément d'un ensemble flou A , dans le domaine X , est un ensemble flou, dénoté par A' , de telle sorte que pour chaque $x \in \mathfrak{R}$:

$$\mu_{A'}(x) = 1 - \mu_A(x). \quad (\text{III.162})$$

L'intersection floue : ET

L'intersection floue de deux ensembles flous A et B concerne à l'expression :

$$(x \text{ est } A) \text{ ET } (y \text{ est } B),$$

où x et y peuvent faire référence à la même variable. Cet opérateur génère un nouvel ensemble flou défini sur l'espace $X \times Y$ et est dénoté par $\mu_{A \cap B}(x, y)$. L'intersection floue peut être réalisée par une famille d'opérateurs connue comme l'ensemble des normes triangulaire ou T-normes et la nouvelle fonction d'appartenance est générée par l'expression :

$$\mu_{A \cap B}(x, y) = \mu_A(x, y) \hat{*} \mu_B(x, y). \quad (\text{III.163})$$

où $\hat{*}$ est l'opérateur T-norme. L'ensemble des T-normes doit satisfaire les conditions citées dans le Tableau III.11 et plusieurs opérateurs possibles le font, [148], mais les deux plus populaires sont l'opérateur min et le produit algébrique :

$$\mu_{A \cap B}(x, y) = \min(\mu_A(x, y), \mu_B(x, y)), \quad (\text{III.164})$$

$$\mu_{A \cap B}(x, y) = \mu_A(x, y) * \mu_B(x, y). \quad (\text{III.165})$$

L'opérateur min forme une borne supérieure dans l'espace des opérateurs d'intersection floue :

$$\mu_A(x, y) \hat{*} \mu_B(x, y) \leq \min(\mu_A(x, y), \mu_B(x, y)). \quad (\text{III.166})$$

L'union floue : OU

L'union de deux ensembles flous A et B concerne la forme :

$$(x \text{ est } A) \text{ OU } (y \text{ est } B),$$

où x et y peuvent faire référence à la même variable. Cet opérateur génère un nouvel ensemble flou défini sur l'espace $X \times Y$, et est dénoté par $\mu_{A \cup B}(x, y)$. Il existe plusieurs façons de réaliser l'union en logique floue. Cette famille d'opérateurs est connue comme w -normes, S-normes ou T-conormes et la nouvelle fonction d'appartenance est générée par l'expression :

$$\mu_{A \cup B}(x, y) = \mu_A(x, y) \hat{+} \mu_B(x, y). \quad (\text{III.167})$$

où $\hat{+}$ est l'opérateur S-norme. L'ensemble des S-normes doit satisfaire les conditions citées dans le Tableau III.11. Les opérateurs flous OU les plus utilisés sont l'opérateur max et la somme probabiliste :

$$\mu_{A \cup B}(x, y) = \max(\mu_A(x, y), \mu_B(x, y)), \quad (\text{III.168})$$

$$\mu_{A \cup B}(x, y) = \mu_A(x, y) + \mu_B(x, y) - \mu_A(x, y) \mu_B(x, y). \quad (\text{III.169})$$

L'opérateur max est la S-norme la plus pessimiste :

$$\max(\mu_A(x, y), \mu_B(x, y)) \leq \mu_A(x, y) \hat{+} \mu_B(x, y). \quad (\text{III.170})$$

TABLEAU III.11

Conditions vérifiées par les T-normes et les S-normes, avec les fonctions d'appartenance $a, b, c, d \in [0, 1]$.

T-NORMES	S-NORMES
$a \hat{*} b = b \hat{*} a$	$a \hat{+} b = b \hat{+} a$
$(a \hat{*} b) \hat{*} c = a \hat{*} (b \hat{*} c)$	$(a \hat{+} b) \hat{+} c = a \hat{+} (b \hat{+} c)$
Si $a \leq c$ et $b \leq d$ alors $a \hat{*} b = c \hat{*} d$	Si $a \leq c$ et $b \leq d$ alors $a \hat{+} b = c \hat{+} d$
$a \hat{*} 1 = a$	$a \hat{+} 0 = a$
Si une T-norme est Archimédienne, alors $a \hat{*} a < a, \forall a(0, 1)$	

III.3.3.3 Le mécanisme de fuzzification, d'inférence et de défuzzification

L'implication floue si-alors

Dans les systèmes d'inférence flous, l'implication représente une relation causale entre les ensembles d'entrée et de sortie où les idées de la représentation locale de la connaissance sont très importantes. En général, un FIS est composé d'un ensemble de règles de production correspondant à un partitionnement de l'espace des variables (multidimensionnel) qui peut se faire selon deux grandes approches :

- *Partitionnement partagé*, qui définit un certain nombre d'ensembles flous pour chacune des dimensions. Ces ensembles introduisent, en conséquence, des valeurs linguistiques communes à l'ensemble des règles.

- *Partitionnement par groupage (clustering)*, qui conduit à une partition de l'espace en groupes homogènes suivant un certain critère. Ce partitionnement conduit à une règle par groupe. Ainsi, les ensembles flous sont propres à une règle et ne sont pas partagés par l'ensemble de règles.

Partitionnement partagé (forme conjonctive)

Dans le partitionnement partagé, l'ensemble des règles de production est de la forme :

$$R_i : \text{si } (x_1 \text{ est } A_i^1 \text{ et } \dots \text{ et } x_n \text{ est } A_i^n) \text{ alors } c_i(y \text{ est } B_i), \quad i=1, \dots, k_R. \quad (\text{III.171})$$

où la variable antécédente x_j est divisé en k_j ensembles flous, l'ensemble de ses termes linguistiques associés est $\mathcal{A}_j = \{A_j^l \mid l=1, \dots, k_j\}$, avec $\mu_{A_j^l}(x_j) : X_j \rightarrow [0,1]$. Similairement, l'ensemble des termes définis pour la variable conséquente y sont dénotés par $\mathcal{B} = \{B_l \mid l=1, \dots, k_y\}$, où $\mu_{B_l}(y) : Y \rightarrow [0,1]$. La variable $c_i \in [0,1]$ représente le niveau de confiance de la règle i , pondérant la connaissance des experts par rapport à une situation particulière, mais également à partir d'un ensemble de paramètres pour ajuster le FIS à partir des données, [33].

L'ensemble des règles de production (III.186) peut être écrit comme suit :

$$R_i : \text{si } (\mathbf{x} \text{ est } A_i) \text{ alors } c_i(y \text{ est } B_i), \quad i=1, \dots, k_R. \quad (\text{III.172})$$

La fonction d'appartenance associée à chaque règle R_i est de dimension $n+1$, $\mu_{R_i}(\mathbf{x}, y)$. Elle est définie par :

$$\mu_{R_i}(\mathbf{x}, y) = c_i * \hat{\mu}_{A_i}(\mathbf{x}) * \hat{\mu}_{B_i}(y). \quad (\text{III.173})$$

dans l'espace cartésien : $R : \mathcal{A}_1 \times \dots \times \mathcal{A}_n \times \mathcal{B} \rightarrow [0,1]$, qui, en version simplifiée, est : $R : \mathcal{A} \times \mathcal{B} \rightarrow [0,1]$, si toutes les règles sont définies pour toutes les combinaisons des termes antécédents, $k_R = \text{card}(\mathcal{A})$.

La Figure III.14 illustre le partitionnement à l'aide des fonctions trapézoïdales sur $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . Les aires grisées de la partie inférieure de la figure montrent les régions de recouvrement des ensembles flous.

Prémisse		Conséquence
si		alors
$x_1 \text{ et } x_2$		y
A_1^1	A_1^2	Conséquence ₁
A_1^1	A_2^2	Conséquence ₂
A_1^1	A_3^2	Conséquence ₃
A_2^1	A_1^2	Conséquence ₄
A_2^1	A_2^2	Conséquence ₅
A_2^1	A_3^2	Conséquence ₆
A_3^1	A_1^2	Conséquence ₇
A_3^1	A_2^2	Conséquence ₈
A_3^1	A_3^2	Conséquence ₉

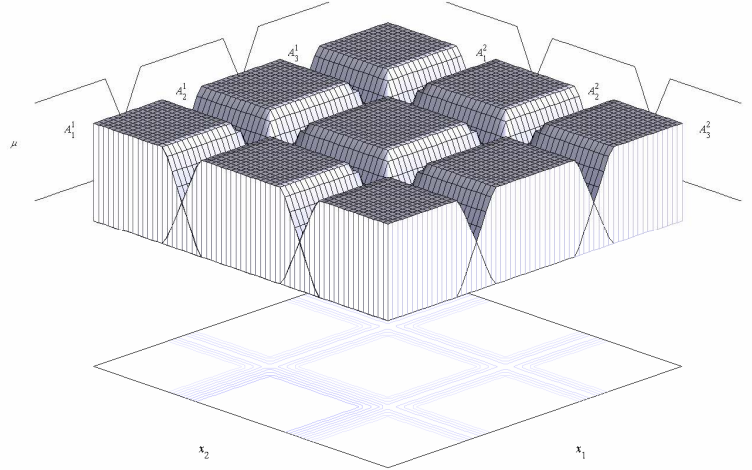


Figure III.14 : Partitionnement à l'aide des fonctions trapézoïdales sur $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . Le partitionnement de l'espace d'entrée est composé de neuf règles. Les aires grisées de la partie basse de la figure montrent les régions de recouvrement des ensembles flous.

Partitionnement par groupage (clustering)

Dans le partitionnement par groupage (clustering), l'ensemble des règles de production a la même forme que (III.172), sauf que la fonction d'appartenance $\mu_{A_i}(\mathbf{x})$ est directement définie dans l'espace multidimensionnel n . Ce partitionnement conduit à une règle par groupe, dans le cas général. Ainsi, les ensembles flous sont propres à une règle et ne sont pas partagés par l'ensemble des règles. Leur interprétation en termes linguistiques est souvent difficile. La Figure III.15 illustre un exemple de partition par groupage à l'aide de trois ensembles flous multidimensionnels A_1 , A_2 , et A_3 , et leur projection respective sur chaque axe pour leur interprétation linguistique.

Il existe d'autres techniques de construction de bases de règles, néanmoins le partitionnement par groupage établit le lien mathématique entre les systèmes flous et neuronaux, [75]. En effet, une phase d'optimisation de la structure des FIS peut avoir lieu si l'on dispose d'observations (bases de données). Nous aborderons ce sujet dans la partie dédiée à la modélisation à partir de la connaissance et des données.

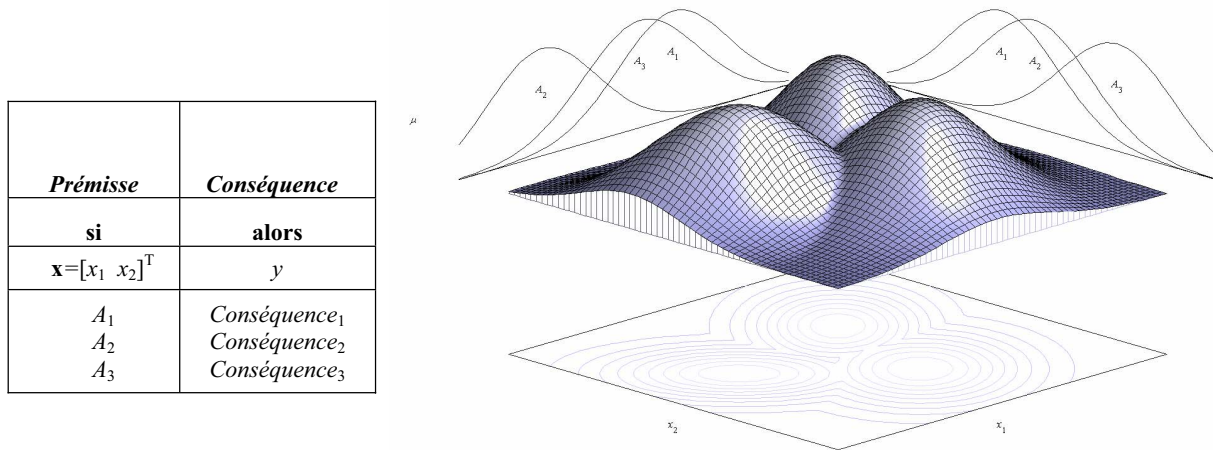


Figure III.15 : Partitionnement par groupage (clustering) à l'aide de trois ensembles flous multidimensionnels A_1 , A_2 , et A_3 , et leur projection respective sur chaque axe pour leur interprétation linguistique. Seule trois règles existent.

L'inférence floue

Les relations d'implication floue sont une représentation de l'expertise qui explique le rapport entre l'ensemble d'entrée linguistique et l'ensemble de sortie. Le mécanisme d'inférence flou établit une correspondance entre l'entrée floue courante de fonction d'appartenance $\mu_A(\mathbf{x})$ avec toutes les prémisses de toutes les règles floues pour combiner ses réponses et produire un seul ensemble flou de sortie $\mu_B(y)$. Le mécanisme d'inférence flou est défini comme suit :

$$\mu_B(y) = \hat{\Sigma}_{\mathbf{x}} (\mu_A(\mathbf{x}) \hat{*} \mu_R(\mathbf{x}, y)) \quad (\text{III.174})$$

où la S-norme $\hat{\Sigma}_{\mathbf{x}}$ est prise sur toutes les valeurs possibles de $\mathbf{x}$ et la T-norme $\hat{*}$ calcule un consensus entre les deux fonctions d'appartenance pour une valeur particulière de $\mathbf{x}$. Quand on choisit $\hat{\Sigma}$ et $\hat{\Pi}$ comme des opérateurs d'intégration (somme) et de produit, respectivement, alors :

$$\mu_B(y) = \int_D \mu_A(\mathbf{x}) \mu_R(\mathbf{x}, y) d\mathbf{x}, \quad (\text{III.175})$$

qui, pour un ensemble flou d'entrée, nécessite le calcul d'une intégrale de dimension n sur le domaine D . En résumé, la sortie floue calculée dépend de l'ensemble flou d'entrée $\mu_A(\mathbf{x})$, du modèle de la règle $\mu_R(\mathbf{x}, y)$, ainsi que des opérateurs d'inférence.

La fuzzification et la défuzzification

La fuzzification

La fuzzification transforme un signal réel en un ensemble flou et est nécessaire quand un FIS considère ce type de variables réelles comme entrée. Pour des entrées entachées par bruit, la forme de l'ensemble flou peut être le reflet de l'incertitude associée au processus d'acquisition. Par exemple, un ensemble triangulaire où le vertex correspond à la moyenne du signal de mesure et la longueur de la base est fonction de l'écart type.

Défuzzification

Quand un ensemble flou de sortie $\mu_B(y)$ est résultat d'un mécanisme d'inférence, il est nécessaire parfois de le comprimer pour produire une seule valeur réelle. Le processus de défuzzification permet d'obtenir une valeur numérique représentative de cet ensemble flou. Il existe plusieurs méthodes de

défuzzification, [148], mais les plus couramment utilisées sont le *centre de gravité* (center of gravity, COG) :

$$y'_g = \frac{\int \mu_B(y) y dy}{\int \mu_B(y) dy} \quad (\text{III.176})$$

et la *moyenne des maximums* (mean of maxima, MOM) :

$$y'_m = \frac{\int \mu_{B''}(y) y dy}{\int \mu_{B''}(y) dy} \quad (\text{III.177})$$

où $\mu_{B''}(y)$ est l'ensemble flou obtenu en prenant l' α -coupe à la hauteur H_B de l'ensemble B . La Figure III.16 illustre les points y_g et y_m correspondant à ces deux méthodes. Le domaine continu des Y est discrétisé pour calculer le centre de gravité. La méthode COG réalise une interpolation entre les conséquents et la méthode MOM sélectionne la sortie la plus « possible ».

Pour éviter une intégration numérique dans la méthode COG, on peut faire appel à la défuzzification moyenne-floue (fuzzy-mean defuzzification). Les ensembles flous conséquents sont, d'abord, défuzzifiés, afin d'obtenir des valeurs numériques représentant les ensembles flous, par exemple par la méthode MOM : $b_j = \text{mom}(B_j)$. La sortie numérique finale est calculée par la somme pondérée :

$$y' = \frac{\sum_{j=1}^{k_y} \omega_j b_j}{\sum_{j=1}^{k_y} \omega_j} \quad (\text{III.178})$$

où k_y est le nombre d'ensembles flous B_j et ω_j est le maximum des degrés de recouvrement sur toutes les règles ayant B_j pour conséquent. Cette méthode assure une interpolation linéaire entre les b_j avec des fonctions d'appartenance antécédentes linéairement distribuées. Puisque la défuzzification individuelle est pré-calculée, la forme et le recouvrement des ensembles flous conséquents n'ont aucune influence et peuvent être remplacés par des valeurs défuzzifiées (singletons). La défuzzification moyenne-floue pondérée (weighted fuzzy-mean defuzzification) prend partiellement en compte les différences entre les ensembles flous conséquents :

$$y' = \frac{\sum_{j=1}^m \gamma_j S_j b_j}{\sum_{j=1}^m \gamma_j S_j} \quad (\text{III.179})$$

où S_j est l'aire sous l'ensemble d'appartenance B_j . L'avantage des méthodes de moyenne-floue est que les paramètres b_j peuvent être estimés par des techniques d'estimation linéaires [86].

III.3.3.4 Types de modèles flous

Selon la structure de la partie conséquente de la règle, on peut distinguer trois principaux types de modèle :

- *Modèle flou linguistique*, où les parties prémisse et conséquence sont toutes les deux des propositions floues.
- *Modèle relationnel flou*, qui est, en quelque sorte, une généralisation du modèle linguistique, permettant l'association d'une prémisse à différentes propositions conséquentes, par une relation floue.
- *Modèle flou de Takagi-Sugeno (TS)*, où la partie conséquence est une fonction mathématique des variables antécédentes plutôt qu'une proposition floue.

Modèle flou linguistique

Le modèle flou linguistique, [98] et [170], est conçu pour capturer de la connaissance qualitative sous forme de règles si-alors :

$$R_i : \text{si } (\mathbf{x} \text{ est } A_i) \text{ alors } (\mathbf{y} \text{ est } B_i), \quad i=1, \dots, k_R. \quad (\text{III.180})$$

Ici, $\mathbf{x}$ est la variable linguistique d'entrée et A_i sont les valeurs linguistiques de cette variable. Similairement, $\mathbf{y}$ est la variable linguistique de sortie et B_i sont les termes linguistiques conséquents

Mécanisme d'inférence

Le mécanisme d'inférence du modèle flou linguistique est basé sur la *règle compositionnelle d'inférence* [130]. Chaque règle (III.180) est une relation floue $R : (X \times Y) \rightarrow [0,1]$:

$$\mu_R(\mathbf{x}, \mathbf{y}) = \mu_A(\mathbf{x}) \hat{*} \mu_B(\mathbf{y}), \quad (\text{III.181})$$

calculée dans l'espace cartésien $X \times Y$. Quand la règle (III.180) est vue comme une implication $A_i \rightarrow B_i$ (A_i implique B_i), en logique classique, cela veut dire que « si A est satisfaite, B doit être également satisfaite pour que l'implication soit vraie ». Néanmoins, on ne peut rien dire de B quand A n'est pas satisfaite et la relation ne peut être inversée. Quand on utilise une conjonction, $A \wedge B$, l'interprétation des règles si-alors est « c'est vrai que A et B sont satisfaites simultanément ». Cette relation est symétrique (non directionnelle) et peut être inversée. Le mécanisme d'inférence flou est basé sur la règle généralisée du *modus ponens* :

$$\frac{\text{si } x \text{ est } A \text{ alors } y \text{ est } B}{x \text{ est } A'} \quad y \text{ est } B'$$

Pour une règle si-alors activée par le fait « x est A' », l'ensemble flou de sortie B' est dérivé de la composition relationnelle max- t ,

$$B' = A' \bullet R. \quad (\text{II.182})$$

Pour la T-norme min, la composition max-min est :

$$\mu_{B'}(y) = \max_x \min(\mu_{A'}(x), \mu_R(x, y)). \quad (\text{II.183})$$

L'agrégation des relations R_i des règles individuelles en une seule relation floue R représente la base entière de règles (III.180). Si les relations R_i représentent des implications, on obtient R par l'opérateur intersection (avec la T-norme min) :

$$R = \bigcap_{i=1}^{k_R} R_i, \text{ c'est à dire, } \mu_R(x, y) = \min_{1 \leq i \leq k_R} \mu_{R_i}(x, y). \quad (\text{III.184})$$

Si l'implication est une T-norme, la relation d'agrégation est l'union des relations individuelles R_i (avec la S-norme max) :

$$R = \bigcup_{i=1}^{k_R} R_i, \text{ c'est à dire, } \mu_R(x, y) = \max_{1 \leq i \leq k_R} \mu_{R_i}(x, y). \quad (\text{III.185})$$

Ensuite l'ensemble flou de sortie B' est inféré de la même manière que dans le cas d'une seule règle en utilisant la règle compositionnelle (II.182). La Figure III.16 illustre un exemple, avec $k_R=3$ et $k_y=3$.

Mécanisme d'inférence max-min (Mamdani)

Une base de règles représente une relation floue. La sortie d'un modèle flou à base de règles est alors extraite d'une composition relationnelle max-min. Jager, [86], a montré que pour des implications floues avec entrées numériques, et pour les t -normes avec entrées numériques et floues, le schéma de raisonnement peut être simplifié. Ceci est très avantageux, car on peut éviter la discrétisation des domaines et le stockage de la relation R . Mamdani, [98], a été le premier à suggérer un algorithme implémentant cette idée :

Algorithme III.6. Inférence min-max (Mamdani).

1. Calcul du degré de recouvrement pour chaque règle par :

$$\beta_i = \max_x (\mu_{A'}(x), \mu_{A_i}(x)), \quad 1 \leq i \leq k. \quad (\text{III.186})$$

Notons que pour un ensemble singleton $(\mu_{A'}(x)=1 \text{ pour } x=x_o \text{ et } \mu_{A'}(x)=0 \text{ autrement})$ $\beta_i = \mu_{A_i}(x_o)$.

2. Calcul des ensembles flous B_i' :

$$\mu_{\beta_i}(y) = \beta_i \wedge \mu_{B_i}(y), \quad y \in Y, \quad 1 \leq i \leq k_R. \quad (\text{III.187})$$

3. Agrégation des ensembles flous B_i' :

$$\mu_{\beta_i}(y) = \max_{1 \leq i \leq k} \mu_{\beta_i}(y), \quad y \in Y. \quad (\text{III.188})$$

Modèle singleton

Il s'agit d'un cas spécial du modèle flou linguistique où les ensembles flous conséquents B_i sont des ensembles singleton. Ces ensembles sont des nombres réels, produisant les règles :

$$R_i: \text{ si } x \text{ est } A_i \text{ alors } y \text{ est } b_i, \quad i=1, \dots, k_R. \quad (\text{III.189})$$

Le nombre des singletons dans la base de règles n'est pas limité et chaque règle peut avoir son propre singleton conséquent. La défuzzification COG résulte en la méthode de la moyenne floue :

$$y = \frac{\sum_{i=1}^{k_R} \beta_i b_i}{\sum_{i=1}^{k_R} \beta_i}. \quad (\text{III.190})$$

Contrairement à (III.179), les k_{k_R} règles ont une contribution dans la défuzzification. Le modèle singleton peut être vu, aussi, comme un cas spécial (d'ordre zéro) du modèle Takagi–Sugeno, présenté dans cette section.

Données	Modèle	
	Prémisse	Conséquence
	si	alors
	x est A_1 x est A_2 x est A_3	y est B_1 y est B_2 y est B_3
	x est A'	y est B'

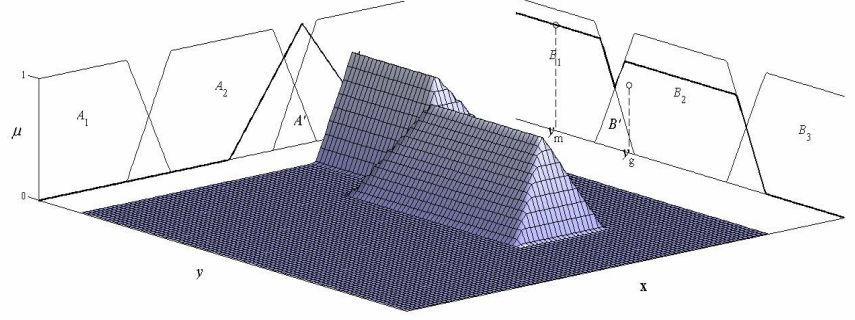


Figure III.16 : Représentation du mécanisme d'inférence max–min (Mamdani). A' est l'ensemble flou d'entrée et B' est l'ensemble flou de sortie. La défuzzification de B' par COG est y_g et par MOM est y_m .

L'avantage du modèle singleton sur le modèle linguistique est qu'il permet d'estimer les paramètres conséquents b_i à partir des données en utilisant la technique des moindres carrés [13]. Le modèle singleton flou appartient aux techniques générales d'approximation de fonctions :

$$y = \sum_{i=1}^k b_i \phi_i(\mathbf{x}) = \mathbf{w}^T \boldsymbol{\phi}(\mathbf{x}), \quad (\text{III.191})$$

où $\boldsymbol{\phi}(\mathbf{x})$ est donné par les degrés de recouvrement des antécédents, et $\mathbf{w}$ est formé par les conséquents b_i . Une interpolation linéaire entre les conséquents a lieu si les fonctions d'appartenance antécédentes sont distribuées selon une partition floue forte et si l'opérateur produit représente le connecteur logique **et** dans les parties antécédentes. Cette propriété est très intéressante, car un modèle ou un contrôleur flou peut être initialisé d'une façon grossière et être optimisé ultérieurement.

Modèle relationnel flou

Les modèles relationnels flous, [170], encodent des associations entre les termes linguistiques définis dans les domaines entrée–sortie à travers des relations floues. Considérons le modèle flou linguistique composé de l'ensemble de règles :

$$R_i : \text{si } x_1 \text{ est } A_i^1 \text{ et } \dots \text{ et } x_n \text{ est } A_i^n \text{ alors } y \text{ est } B_i, \quad i=1, \dots, k_R. \quad (\text{III.192})$$

L'objectif est de faire que la base de règles (III.192) soit représentée comme une relation flou $R=[r_{ij}]$ entre l'ensemble des termes antécédents $\mathcal{A}_j$ et l'ensemble des termes conséquents $\mathcal{B}$. Ainsi, R est représenté par une matrice $k_R \times k_y$, contrainte à avoir au moins un élément différent de zéro dans chaque ligne. Chaque règle contient tous les termes conséquents possibles, chacune avec ses propres facteurs de poids, correspondant à l'élément r_{ij} de la relation flou. Ces poids permettent au modèle d'être facilement réglé pour s'adapter si l'on dispose de données. La Figure III.17 montre un exemple dans $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 .

Mécanisme d'inférence

Le mécanisme d'inférence du modèle relationnel flou est basé sur la *composition relationnelle*, consistant en une phase de combinaison et une phase de projection. Zadeh a proposé la composition sup–min, [171] :

$$\mu_B(y) = \sup_x \min(\mu_A(x), \mu_R(x, y)), \quad (\text{III.193})$$

de l'ensemble flou représentant les degrés de recouvrement β_i , (III.186), et la relation R . L'algorithme est :

Algorithme III.7. Mécanisme d'inférence du modèle relationnel flou.

1. Calcul du degré de recouvrement par :

$$\beta_i = \mu_{A_1^i}(x_1) \wedge \dots \wedge \mu_{A_n^i}(x_n), \quad 1 \leq i \leq k_R. \quad (\text{III.194})$$

2. Application de la composition relationnelle $\omega = \beta \bullet R$, donnée par :

$$\omega_j = \max_{1 \leq i \leq k_R} (\beta_i \wedge r_{ij}), \quad 1 \leq j \leq k_y. \quad (\text{III.195})$$

3. Défuzzification de l'ensemble flou conséquent par:

$$y = \frac{\sum_{l=1}^{k_y} \omega_l b_l}{\sum_{l=1}^{k_y} \omega_l} \quad (\text{III.196})$$

où b_l sont les centres de gravité des ensembles conséquents des ensembles flous B_l , calculés en appliquant une méthode de défuzzification comme le COG ou le MOM des ensembles flous individuels B_l .

Le principal avantage du modèle relationnel flou est que la partition entrée–sortie peut être réglée sans changer les ensembles flous conséquents. Dans le modèle flou linguistique, les sorties des règles individuelles sont restreintes à la grille donnée par les centres de gravité des ensembles flous de sortie [13].

Cette liberté engage un nombre de paramètres plus élevé. S'il n'y a pas de contraintes sur ces paramètres, plusieurs éléments dans une ligne peuvent devenir zéro, ce qui peut gêner l'interprétation du modèle. De plus, la forme des ensembles flous n'a aucune influence sur la valeur défuzzifiée résultante, car le modèle ne considère que les centres de gravité des ensembles dans la défuzzification.

Si les ensembles flous antécédents forment une partition floue forte et si l'on utilise la composition somme–produit, le modèle relationnel a, au niveau implémentation, un modèle équivalent à conséquents singleton. Si les ensembles d'appartenance conséquents forment une partition floue, un modèle singleton a un modèle relationnel équivalent en calculant les degrés d'appartenance des singletons dans les ensembles flous conséquents B_j . Ces degrés d'appartenance deviennent donc des éléments de la relation floue :

$$R = \begin{bmatrix} \mu_{B_1}(b_1) & \mu_{B_2}(b_1) & \dots & \mu_{B_{k_y}}(b_1) \\ \mu_{B_1}(b_2) & \mu_{B_2}(b_2) & \dots & \mu_{B_{k_y}}(b_2) \\ \vdots & \vdots & \ddots & \vdots \\ \mu_{B_1}(b_{k_R}) & \mu_{B_2}(b_{k_R}) & \dots & \mu_{B_{k_y}}(b_{k_R}) \end{bmatrix} \quad (\text{III.197})$$

Ainsi, le modèle flou linguistique est un cas spécial du modèle relationnel flou, avec R une relation numérique contrainte telle que seul un élément soit différent de zéro dans chaque ligne de R (chaque règle a un seul conséquent).

Prémisse		Conséquence		
si		alors		
		y		
x_1 et x_2		B_1	B_2	B_3
A_1^1	A_1^2	$r_{1,1}$	$r_{1,2}$	$r_{1,3}$
A_1^1	A_2^2	$r_{2,1}$	$r_{2,2}$	$r_{2,3}$
A_1^1	A_3^2	$r_{3,1}$	$r_{3,2}$	$r_{3,3}$
A_2^1	A_1^2	$r_{4,1}$	$r_{4,2}$	$r_{4,3}$
A_2^1	A_2^2	$r_{5,1}$	$r_{5,2}$	$r_{5,3}$
A_2^1	A_3^2	$r_{6,1}$	$r_{6,2}$	$r_{6,3}$
A_3^1	A_1^2	$r_{7,1}$	$r_{7,2}$	$r_{7,3}$
A_3^1	A_2^2	$r_{8,1}$	$r_{8,2}$	$r_{8,3}$
A_3^1	A_3^2	$r_{9,1}$	$r_{9,2}$	$r_{9,3}$

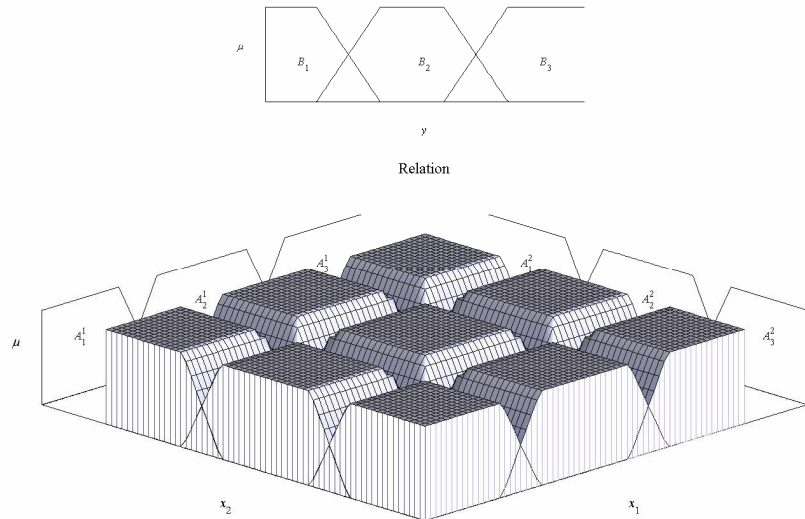


Figure III.17 : Modèle flou relationnel dans $\mathbb{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . La relation R du partitionnement de l'espace d'entrée est composée d'un total de neuf règles.

Modèle flou de Takagi–Sugeno (TS)

A la différence des deux autres types de modèles, le modèle flou Takagi–Sugeno (TS) utilise des fonctions numériques dans les conclusions, [148]. Cette technique peut être vue comme un compromis entre modélisation linguistique et régression mathématique au sens où les antécédents décrivent des

régions floues dans l'espace d'entrée dans lequel les fonctions conséquentes sont valides. Les règles TS ont la forme suivante :

$$R_i : \text{si } (x_1 \text{ est } A_i^1 \text{ et } \dots \text{ et } x_n \text{ est } A_i^n) \text{ alors } (y_i = f_i(\mathbf{x})), \quad i=1, \dots, k_R. \quad (\text{III.198})$$

Contrairement au modèle linguistique, l'entrée $\mathbf{x}$ est une variable numérique (en principe, l'utilisation d'entrées linguistiques est possible, mais requiert le principe d'extension pour calculer la valeur de y_i). En général, les fonctions f_i sont typiquement de la même structure, seuls les paramètres changent entre chaque règle. Typiquement, f_i est une fonction vectorielle, mais pour des raisons de simplicité nous considérons des fonctions scalaires f_i . Evidemment, les fonctions linéaires sont les fonctions les plus simples et les plus pratiques à paramétrer :

$$R_i : \text{si } (x_1 \text{ est } A_i^1 \text{ et } \dots \text{ et } x_n \text{ est } A_i^n) \text{ alors } (y_i = \mathbf{w}_i^T \mathbf{x} + b_i), \quad i=1, \dots, k_R. \quad (\text{III.199})$$

Ce modèle est couramment connu comme modèle TS linéaire (affine TS model). Nous retrouvons le modèle singleton (III.189) si $\mathbf{w}_i = \mathbf{0}$ (modèle TS d'ordre zéro).

Mécanisme d'inférence

La formule d'inférence est une extension du mécanisme d'inférence avec singleton :

$$y = \frac{\sum_{i=1}^{k_R} \beta_i y_i}{\sum_{i=1}^{k_R} \beta_i} = \frac{\sum_{i=1}^{k_R} \beta_i (\mathbf{w}_i^T \mathbf{x} + b_i)}{\sum_{i=1}^{k_R} \beta_i}. \quad (\text{III.200})$$

Quand les ensembles flous antécédents définissent des régions distinctes mais avec recouvrement dans l'espace antécédent, les paramètres $\mathbf{w}_i$ et b_i correspondent à une linéarisation locale d'une fonction non linéaire : le modèle TS peut être vu comme une approximation linéaire par morceaux.

Le modèle TS comme un système quasi-linéaire

Le modèle TS peut se représenter comme un système quasi-linéaire (système linéaire à paramètres dépendant des entrées). Dénotons le degré normalisé de recouvrement :

$$\gamma_i(\mathbf{x}) = \beta_i(\mathbf{x}) / \sum_{j=1}^{k_R} \beta_j(\mathbf{x}). \quad (\text{III.201})$$

Nous écrivons explicitement β_i en fonction de $\mathbf{x}$ pour souligner que le modèle TS est un modèle quasi-linéaire de la forme suivante :

$$y = \sum_{i=1}^{k_R} \gamma_i(\mathbf{x}) \mathbf{w}_i^T \mathbf{x} + \sum_{i=1}^{k_R} \gamma_i(\mathbf{x}) b_i = \mathbf{w}^T(\mathbf{x}) \mathbf{x} + b(\mathbf{x}). \quad (\text{III.202})$$

Conséquemment, les paramètres $\mathbf{w}(\mathbf{x})$ et $b(\mathbf{x})$ sont des combinaisons linéaires, convexes, des paramètres conséquents $\mathbf{w}_i$ et b_i :

$$\mathbf{w}^T(\mathbf{x}) = \sum_{i=1}^{k_R} \gamma_i(\mathbf{x}) \mathbf{w}_i^T, \quad b(\mathbf{x}) = \sum_{i=1}^{k_R} \gamma_i(\mathbf{x}) b_i. \quad (\text{III.203})$$

En ce sens, le modèle TS est une transformation de l'espace antécédent (entrée) en une région convexe.

Prémisse		Conséquence		
si		alors		
x_1 et x_2		y $y_i = \mathbf{w}_i^T \mathbf{x} + b_i$		
A_1^1	A_1^2	$w_{1,1}$	$w_{1,2}$	b_1
A_1^1	A_2^2	$w_{2,1}$	$w_{2,2}$	b_2
A_1^1	A_3^2	$w_{3,1}$	$w_{3,2}$	b_3
A_2^1	A_1^2	$w_{4,1}$	$w_{4,2}$	b_4
A_2^1	A_2^2	$w_{5,1}$	$w_{5,2}$	b_5
A_2^1	A_3^2	$w_{6,1}$	$w_{6,2}$	b_6
A_3^1	A_1^2	$w_{7,1}$	$w_{7,2}$	b_7
A_3^1	A_2^2	$w_{8,1}$	$w_{8,2}$	b_8
A_3^1	A_3^2	$w_{9,1}$	$w_{9,2}$	b_9

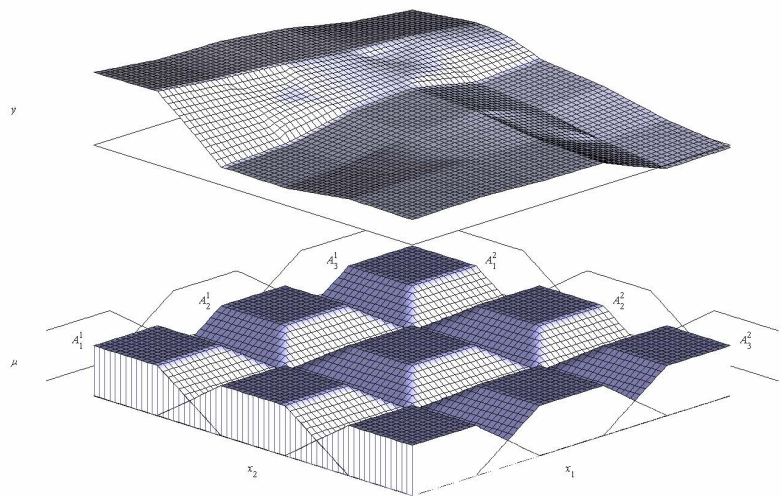


Figure III.18 : Modèle flou Takagi-Sugeno dans $\mathcal{R}^2$, avec $k_1=3$ pour la variable x_1 et $k_2=3$ pour la variable x_2 . La partition de l'espace d'entrée est composée de neuf règles (bas). La sortie y est une approximation linéaire par morceaux d'une fonction non linéaire (haut).

III.3.3.5 Génération de règles à partir de l'expertise

La génération de règles floues à partir de l'expertise se fait selon les pas suivants :

- Détermination de la structure des systèmes flous. Cette partie inclut la sélection des entrées et des sorties. Typiquement, les systèmes multi-entrée multi-sortie (multiple-input multiple-output, MIMO) sont décomposés en plusieurs systèmes multi-entrée mono-sortie (multiple-input single-output, MISO).
- Détermination des termes linguistiques et de leur fonction d'appartenance pour chaque variable floue. Les univers de discours de l'entrée comme celui de la sortie (cette dernière dépendant du type de modèle flou) doivent être couverts par des termes linguistiques (petit, moyen et grand, par exemple).
- Détermination de la méthode d'inférence floue. Cette partie est dépendante du type de modèle flou (linguistique, relationnel ou TS).
- Défuzzification. Nécessaire si le mécanisme d'inférence a comme sortie un ensemble flou. Si le modèle est de type TS, la défuzzification est incluse dans le mécanisme d'inférence.

La qualité et la performance des FIS basés sur l'expertise dépendent énormément du problème traité ainsi que de la qualité et de la quantité de connaissance disponible. Pour certains problèmes, cette situation peut amener à l'obtention rapide de modèles utiles, tandis que pour d'autres, elle peut amener à des modèles très coûteux en termes de temps et de calcul.

III.4 Modélisation à partir de la connaissance et les données

Les règles floues sont capables de représenter une connaissance compréhensible par l'être humain. La construction traditionnelle de règles floues est basée sur l'expertise et des heuristiques, ce qui induit deux désavantages. D'un côté, les règles floues sont très simples et la performance du FIS est faible ; souvent, les fonctions d'appartenance floues sont définies de façon grossière et, en conséquence, la connaissance représentée par ces règles peut être superficielle. De l'autre côté, l'extraction efficace de règles floues décrivant des systèmes de dimension importante devient très difficile à cause des limites de l'abstraction humaine.

L'hybridation des techniques d'apprentissage avec les systèmes d'inférence flous est un bon exemple de combinaison de techniques. L'idée de base est qu'un système d'inférence flou peut tirer profit d'observations mesurées pour améliorer sa qualité d'approximation, sans perdre son niveau d'interprétabilité. Il existe deux types principaux de modèles flous avec apprentissage :

- *Modèles flous conventionnels avec apprentissage*, où un algorithme d'apprentissage est directement utilisé pour ajuster les paramètres des prémisses des règles floues, sans représenter le FIS sous une structure de réseau de neurones. Ce type de modèle est directement applicable au cas du partitionnement partagé.
- *Modèles neuro-flous*, où le FIS est d'abord interprété avec une architecture de réseau de neurones générale ou spécifique, pour ensuite utiliser des techniques d'apprentissage pour le paramétrer. Ce type de modèles est le plus largement abordé dans le cadre des techniques neuro-floues. Le partitionnement par groupage est directement abordé par ce type de modèles.

III.4.1 Modèles flous conventionnels avec apprentissage

Les modèles flous avec apprentissage utilisent des algorithmes d'apprentissage afin de profiter, en même temps, de l'expertise et de l'information provenant des bases des données, sans changer l'architecture de base du système flou. Dans ces conditions, seul le partitionnement partagé est concerné car l'apprentissage se déroule par l'optimisation du nombre d'ensembles flous pour chaque dimension, de leur respective paramétrisation et du nombre de règles à implémenter.

Pour préserver l'interprétabilité des règles, ce type de techniques doit retenir les points suivants, [90] :

- Le nombre d'ensembles flous dans une partition floue ne doit pas être trop grand. Un nombre élevé de sous ensembles flous rend leur signification en termes linguistiques difficile à interpréter.

- Le nombre de règles floues dans la base de règles peut être grand, notamment si la quantité de données disponibles est importante.
- Les règles générées à partir des données peuvent entrer en conflit entre elles, ce qui peut être résolu par l'assignation d'un degré de vérité à chaque règle, [161], ou en vérifiant la consistance des règles floues, [90].

III.4.1.1 Choix du nombre d'ensembles flous

Le choix du nombre d'ensembles flous par dimension ($k_1, k_2, \dots, k_n$) est délicat et lourd de conséquences. D'une part, s'il est trop faible, le système aura une déficience à représenter les comportements non-linéaires, d'autre part, si le niveau de granularité augmente, les ensembles flous deviennent trop spécifiques, ce qui entraîne une réduction des capacités de généralisation du FIS. Pour éviter de fixer le nombre d'ensembles flous arbitrairement, des auteurs proposent d'adapter l'espace d'entrée aux données, soit par un *affinement de partition*, soit par l'utilisation d'un *algorithme génétique* [73], par exemple, d'autres méthodes étant possibles.

D'un côté, l'affinement de partition consiste en un processus itératif, commençant par une partition initiale grossière, généralement deux sous-ensembles flous par entrée et, à chaque étape, l'algorithme ajoute un sous-ensemble flou sur l'une des entrées du système, là où l'erreur de modélisation est maximale. Plusieurs travaux proposent des critères pour choisir la portion de l'espace nécessitant un découpage plus fin ainsi que pour arrêter le processus de partition lui-même. De l'autre côté, les algorithmes génétiques sont utilisés pour choisir le niveau de résolution, de granularité, le plus adapté à chacune des régions de l'espace. Pour cela, la fonction de coût vise à minimiser le nombre de règles tout en maximisant le nombre d'exemples bien classés à partir d'un ensemble de partitions avec différents valeurs de k_i pour une variable i donnée. Jeng a utilisé les SVM pour trouver le nombre optimal de règles et initialiser les fonctions d'appartenance gaussiennes des variables [86].

On peut envisager, éventuellement, l'implémentation de toutes les règles correspondant à l'ensemble des combinaisons des sous-ensembles flous des variables d'entrées [85], donnant pour un système à n entrées un total de $k_1 \times k_2 \times \dots \times k_n$ règles. Des améliorations permettent de gérer les inconvénients inhérents à cette méthode. Par exemple, il se peut que certaines règles ne soient pas initialisées du fait d'un niveau de couverture de l'espace insuffisant. Une méthode *initialisant toutes les règles* est la procédure de diffusion [73], qui définit, pour un ensemble de règles S , le voisinage v_i de la règle i , $v_i \subset S$, de telle sorte que v_i est l'ensemble des règles dont la prémisse diffère de celle de i par un seul sous-ensemble flou. De cette manière, chaque règle n'a que deux voisins par dimension, et donc le cardinal de v_i est borné : $|v_i| \leq 2^p$.

III.4.1.2 Paramétrisation des règles

Même si l'on peut établir le nombre optimal de partitions de l'espace de discours de chaque variable d'entrée, on peut se retrouver avec le problème du *fléau de la dimension*, car le nombre total de règles du système est égal au produit du nombre des valeurs linguistiques par variable.

Wang et Mendel [161] ont proposé une méthode qui, premièrement, définit une partition floue pour chaque variable linguistique ($k_1, k_2, \dots, k_n$), deuxièmement, associe une règle R_i pour chaque exemple $\mathbf{x}_i$ pour ensuite « fusionner » les règles avec les mêmes prémisses et les mêmes conséquences. La principale faiblesse de cette méthode est la définition *a priori* de la structure des règles. De plus, si la dimension d'entrée est élevée et la quantité des observations est grande, le nombre de règles devient immense. Une solution évitant ce problème, [89], est de partager l'espace en n_r régions et dans chacune lancer la recherche des observations ayant le minimum et le maximum de la sortie et en générer deux règles.

Dans [78], Herrera a utilisé les algorithmes génétiques pour choisir le meilleur ensemble de règles à partir d'un codage par chromosomes (un par règle). Alternativement, les *arbres de décision flous*, dont chaque feuille de l'arbre correspond à une règle impliquant un nombre variable d'entrées, génèrent des règles incomplètes, comme celles générées par les experts, car elles n'appliquent pas toutes les variables disponibles [83], [108]. L'apprentissage renforcé est également utilisé pour ajuster les paramètres des règles flous, [92].

III.4.2 Modèles neuro-flous

Dans les systèmes neuro-flous, le système flou est représenté par un réseau de neurones, ce qui permet l'utilisation de n'importe quelle technique d'apprentissage développée dans le cadre des réseaux de neurones. De ce fait, un système neuro-flou est un FIS qui peut « apprendre ». Un exemple est le système d'inférence flou basé sur des réseaux adaptatifs (adaptive-network-based fuzzy inference system, ANFIS), [87]. Les réseaux adaptatifs (adaptive networks) décrivent, dans un cadre unifié, les systèmes flous et les réseaux de neurones, [90]. Parmi eux, les réseaux à fonctions de base radiale, RBFN, sont mathématiquement équivalents aux systèmes flous à partition par groupage. De même, le perceptron multicouches, MLP, peut être également interprété sous forme de règles floues, [36].

Les RBFN, vus dans la section 2 de ce chapitre, sont un réseau à trois couches : une couche d'entrée, une couche de sortie et une couche cachée. Chaque cellule de la couche cachée est une unité de réception à base radiale qui est ajustée localement ; c'est là où l'analogie avec le partitionnement par groupage : chaque unité de réception correspond à un prototype. Dans la communauté scientifique dédiée à l'étude des systèmes flous, l'algorithme des *c-moyennes floues* permet de définir des clusters ou prototypes à partir des données. Dans cette section nous abordons les techniques de groupage flou à partir des données (fuzzy clustering), tout en remarquant leur analogie avec les réseaux RBFN.

III.4.2.1 Les algorithmes de groupage par données (clustering)

Les c-moyennes floues (FCM)

En 1974, Dunn, [49], a proposé la méthode des *c-moyennes floues* (fuzzy c-means, FCM), méthode de base du groupage flou. Bezdek, [24], en a démontré les propriétés de convergence et a proposé les premiers indices de mesure de la qualité des partitions obtenues. Chacun des ℓ individus appartient à chacun des nc groupes avec un coefficient d'appartenance μ , μ_{ij} étant le degré d'appartenance de l'exemple $\mathbf{x}_i$ au groupe c_j . Définissons d_{ij}^2 comme la distance entre l'exemple $\mathbf{x}_i$ et le groupe c_j . De façon générale, elle est définie comme :

$$d_{ij}^2 = \|\mathbf{x}_i - \mathbf{v}_j\|_A^2 = (\mathbf{x}_i - \mathbf{v}_j)^T \mathbf{A} (\mathbf{x}_i - \mathbf{v}_j),$$

$\mathbf{A}$ étant une matrice symétrique définie positive et $\mathbf{v}_j$ étant le prototype du groupe c_j . Les FCM utilisent $\mathbf{A}=\mathbf{I}$, retrouvant la norme euclidienne standard. Ce cas spécifique peut être vu, dans le cadre du partitionnement partagé, comme un optimisation des FIS avec des fonctions d'appartenance gaussiennes mais également, dans le cadre des RBFN, comme l'utilisation de fonctions d'activation gaussiennes, décrites par une matrice identité, définissant une distance Euclidienne, Figure III.19.

Soit $\mathbf{U}$ la matrice de coefficients μ_{ij} et $\mathbf{V}$ celle des coordonnées des centres $\mathbf{v}_j$. L'objectif de l'algorithme est de trouver $\mathbf{U}$ et $\mathbf{V}$ solution du problème d'optimisation :

$$\begin{array}{ll} \text{Minimiser} & J_{FCM} = \sum_{i=1}^{\ell} \sum_{j=1}^{nc} \mu_{ij}^m D_{ij}^2 \\ \mathbf{U}, \mathbf{V} & \end{array} \quad (III.204)$$

$$\begin{array}{ll} \text{Sous la contrainte} & \sum_{j=1}^{nc} \mu_{ij} = 1, \quad i = 1, \dots, \ell, \end{array} \quad (III.205)$$

où $m \geq 1$, est l'exposant flou (une constante).

La méthode d'optimisation la plus employée pour résoudre (III.204)–(III.205), dite optimisation par alternance, est basée sur les conditions de stationnarité de la fonction Lagrangienne associée au problème d'optimisation. Au départ, une fois le nombre de groupes nc , la valeur de l'exposant flou m et la valeur de la tolérance ε fournis, les coefficients μ_{ik} sont initialisées aléatoirement. Puis, chaque itération k est composée des pas suivants :

 Algorithme III.8. Les c-moyennes floues, FCM.

1. Calcul des prototypes $\mathbf{v}_j$:

$$\mathbf{v}_j^{(k)} = \frac{\sum_{i=1}^{\ell} \mu_{ij}^{(k-1)} \mathbf{x}_i}{\sum_{i=1}^{\ell} \mu_{ij}^{(k-1)}}, \quad j=1, \dots, nc, \quad (\text{III.206})$$

2. Calcul des distances D_{ij} :

$$d_{ij}^2 = \|\mathbf{x}_i - \mathbf{v}_j^{(k)}\|_A^2 = (\mathbf{x}_i - \mathbf{v}_j^{(k)}) \mathbf{A} (\mathbf{x}_i - \mathbf{v}_j^{(k)})^T, \quad (\text{III.207})$$

3. Calcul des degrés d'appartenance μ_{ij} :

$$\mu_{ij} = \frac{1}{\sum_{l=1}^{nc} (d_{ij}/d_{il})^{\frac{2}{m-1}}}, \quad (\text{III.208})$$

4. Répéter jusqu'à $\|\mathbf{U}^{(k)} - \mathbf{U}^{(k-1)}\| < \varepsilon$
-

Les variations autour des FCM

Il existe plusieurs améliorations apportées à l'algorithme de base [73]. Dans l'algorithme *Gustafson-Kessel*, [74], la matrice $\mathbf{A}$ dépend des données et est en fonction, pour chaque groupe, de l'inverse de la matrice de covariance. Dans le cadre des RBFN, c'est équivalent à utiliser des fonctions d'activation gaussiennes, décrites par la matrice de covariance, associant ainsi une distance de Mahalanobis, Figure III.19. Pour le groupe c_j , la matrice de covariance $\mathbf{C}_j$ est :

$$\mathbf{C}_j = \frac{\sum_{i=1}^{\ell} (\mu_{ij})^m (\mathbf{x}_i - \mathbf{v}_j)(\mathbf{x}_i - \mathbf{v}_j)^T}{\sum_{i=1}^{\ell} (\mu_{ij})^m}, \quad (\text{III.209})$$

et la distance du point $\mathbf{x}_i$ au groupe c_j devient :

$$d_{ij}^2 = \|\mathbf{x}_i - \mathbf{v}_j\|_A^2 = (\mathbf{x}_i - \mathbf{v}_j) [\det(\mathbf{C}_j)]^{1/n} \mathbf{C}_j^{-1} (\mathbf{x}_i - \mathbf{v}_j)^T. \quad (\text{III.210})$$

L'algorithme *Gath-Geva*, [59], est équivalent à un mélange de gaussiennes [1], en remplaçant simplement μ_{ik} par une « probabilité » *a posteriori* :

$$P_j = \frac{\sum_{i=1}^{\ell} (\mu_{ij})^m}{\sum_{i=1}^{\ell} \sum_{k=1}^{nc} (\mu_{ik})^m},$$

pour ensuite calculer la distance du point $\mathbf{x}_i$ au groupe c_j :

$$d_{ij}^2 = \frac{1}{P_j} \sqrt{\mathbf{C}_j} \exp\left(1/2 (\mathbf{x}_i - \mathbf{v}_j) \mathbf{C}_j^{-1} (\mathbf{x}_i - \mathbf{v}_j)^T\right). \quad (\text{III.211})$$

La version possibiliste du FCM enlève la contrainte (III.205) et ajoute un terme qui pénalise les degrés d'appartenance trop petits. La fonction des *c-moyennes floues possibilistes* (possibilistic c-means, PCM) est :

$$J_{PCM} = \sum_{i=1}^{\ell} \sum_{j=1}^{nc} \left(\mu_{ij}^m \|\mathbf{x}_i - \mathbf{v}_j\|^2 + \eta_j (1 - \mu_{ij}^m) \right),$$

avec $\eta_j = \frac{\sum_{i=1}^{\ell} \mu_{ij}^m \|\mathbf{x}_i - \mathbf{v}_j\|^2}{\sum_{i=1}^{\ell} \sum_{j=1}^{nc} \mu_{ij}^m}.$

Dans l'algorithme des *modèles de c-régression floue* (fuzzy c-regression model, FCRM), les groupes ne sont, ni des hypersphères (FCM), ni des hyper-ellipses orientées (GK), mais des hyperplans [113]. D'ailleurs, les prototypes, correspondant aux centres des groupes, sont également changés par des hyperplans, donc, pour le groupe c_j : $y_j = \mathbf{w}_j^T \mathbf{x} + b_j$. Les fonctions d'appartenance aux sous-ensembles flous des entrées sont des gaussiennes généralisées :

$$A(x) = \exp\left(-\frac{x-p_1}{p_2}\right)^2.$$

Les paramètres p_1 et p_2 , ainsi que les coefficients $\mathbf{w}_j$ et b_j , sont ajustés par une descente de gradient.

III.4.2.3 Le groupage SVM

Le point le plus sensible dans la construction des systèmes d'inférence floue est toujours la prémisse des règles. L'utilisation des modèles hybrides profite de l'expertise et des observations. Dans le partage par groupage, les techniques de groupage flou permettent de construire, initialement, des ensembles flous multidimensionnels pour, ensuite, les utiliser directement dans le modèle (modèle relationnel ou modèle TS) où après une projection individuelle sur les variables (modèle linguistique).

Dans ce type de technique, la performance dépend du nombre de groupes dans les données, de la forme et du volume de chaque groupe. Dans les algorithmes de groupage par données, la distance au centre du prototype définit sa forme. Par exemple, l'algorithme des c-moyennes floues utilise la distance de norme euclidienne, appropriée pour représenter des formes sphériques. L'algorithme Gustafson-Kessel utilise des matrices de covariance floues pour obtenir une distance de Mahalanobis, représentant des ellipsoïdes orientées. Dans nos travaux, [70], nous introduisons les machines à vecteurs de support pour l'estimation de la densité, SVDE, pour construire des frontières, à partir desquelles on peut calculer des distances orthogonales, Figure III.19.

Pour le groupe c_j , il faut résoudre une variation du problème d'optimisation quadratique (III.120)-(III.122) afin d'introduire la matrice d'appartenance floue, donnant le QP suivant :

$$\begin{array}{ll} \text{Maximiser} & L_D(\mathbf{a}) = -\frac{1}{2} \sum_{i,k=1}^{\ell} h_{ij} \alpha_i \alpha_k \mathbf{x}_i^T \mathbf{x}_k \\ \mathbf{a} & \end{array} \quad (III.212)$$

$$\begin{array}{ll} \text{Sous les contraintes} & \sum_{i=1}^{\ell} \alpha_i = 1, \end{array} \quad (III.213)$$

$$0 \leq \alpha_i \leq \frac{1}{\nu \ell_j} \quad i = 1, \dots, \ell. \quad (III.214)$$

où $h_{ij}=1$ si $\mu_{ij} > 0.15$ et $h_{ij}=0$ autrement.

La fonction (III.123) décrit la distribution apprise, sous forme de frontière ou PDF, et est définie par : les vecteurs de support $\mathbf{x}_{sv}$, leurs poids respectifs α et le seuil b . Comme cette fonction procure des distances négatives pour indiquer que l'exemple évalué réside hors la distribution, nous l'avons modifiée pour n'avoir que des distances positives. De ce fait, la distribution du groupe c_j , définie par la SVM $_j$, calcule la distance du point $\mathbf{x}_i$ au groupe c_j :

$$d_{ij}^2 = \left(\sum_{k=1}^{nsv_j} \alpha_k k(\mathbf{x}_{sv_{j,k}}, \mathbf{x}_i) - 2b_j \right) / 2b_j. \quad (III.215)$$

La valeur de b est calculée par l'équation (III.124).

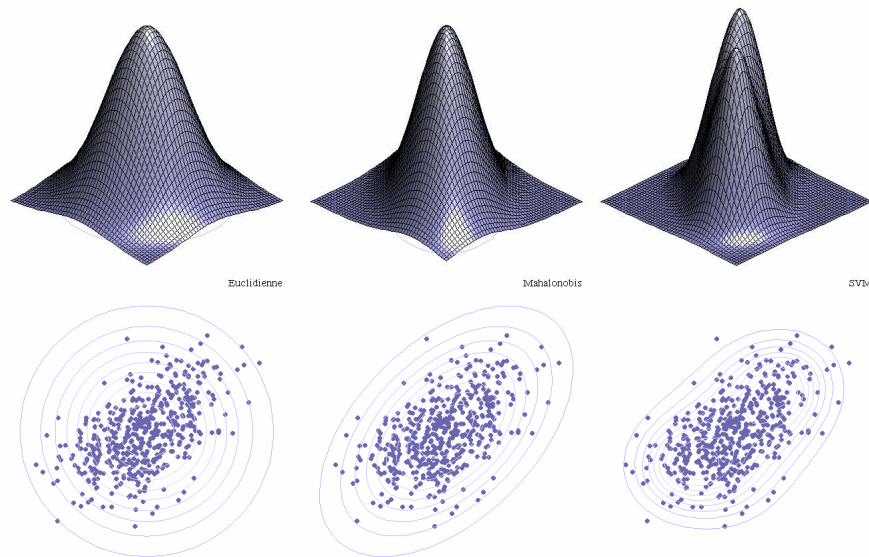


Figure III.19 : Normes des distances : Euclidienne (cercle), Mahalanobis (ellipsoïde orienté) et SVM (ensemble non convexe).

III.4.2.4 Initialisation des algorithmes de groupage

Le nombre de groupes

Le nombre de groupes nc est le paramètre le plus critique, obstacle connu comme problème de validité de groupe (cluster validity problem), au sens où les autres paramètres restant ont une influence plus faible sur la partition résultante. Pour obtenir un groupage des données avec une partition optimale, sans aucune information *a priori* de sa structure, on peut soit exécuter un algorithme de groupage avec un nombre croissant de groupes ($1 \leq nc \leq \ell - 1$) et caractériser les partitions obtenues par des indices, soit utiliser un algorithme qui détermine le nombre de groupes.

Dans la première approche, on fait appel à des *mesures de validité* qui sont des indices scalaires qualifiant la bonne qualité de la partition obtenue. Xie et Beni, [166], minimisent le rapport de la compacité, $C(nc)$, à la séparabilité, $S(nc)$, définies comme :

$$C(nc) = \sum_{i=1}^{\ell} \sum_{j=1}^{nc} \mu_{ij}^m \|\mathbf{x}_i - \mathbf{v}_j\|_A^2 \quad \text{et} \quad S(nc) = \min_{j \neq k} \|\mathbf{v}_j - \mathbf{v}_k\|_A^2.$$

Sugeno, [144], a proposé de retenir le nombre de groupes qui minimise le critère :

$$V(nc) = \sum_{i=1}^{\ell} \sum_{j=1}^{nc} \mu_{ij}^m \left(\|\mathbf{x}_i - \mathbf{v}_j\|_A^2 - \|\mathbf{v}_j - \bar{\mathbf{x}}\|_A^2 \right)$$

où $\bar{\mathbf{x}}$ est le centre de gravité des données, ou son extension flou $\bar{\mathbf{v}}$, [50]. Le premier terme est la variance intra-groupe et le second correspond à la variance inter-groupes. Burrough, [35], utilise un coefficient de partition F et un coefficient d'entropie, H :

$$F = \frac{1}{nc} \sum_{i=1}^{\ell} \sum_{j=1}^{nc} \mu_{ij}^2, \quad \frac{1}{nc} < F < 1, \quad \text{et} \quad H = \frac{1}{nc} \sum_{i=1}^{\ell} \sum_{j=1}^{nc} -\mu_{ij} \ln(\mu_{ij}), \quad 1 - F < H < \ln(nc).$$

Une bonne partition aura un coefficient de partition grand et un coefficient d'entropie petit. Le choix du compromis revient à l'utilisateur.

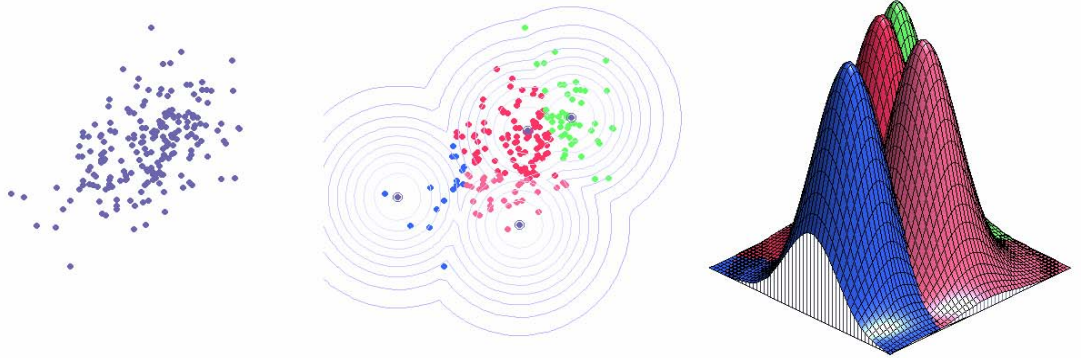


Figure III.20 : La groupage soustractif (de gauche à droite) : l'ensemble aléatoire de 100 points, les groupes identifiés et leurs centres (entourés), les fonctions d'appartenance gaussiennes.

Dans la deuxième approche, l'idée de *la fusion de groupes* est de commencer avec un nombre suffisamment grand de groupes et, successivement, le réduire en fusionnant les groupes « similaires » par rapport à un certain critère. Partant d'une approche opposée, *l'insertion de groupes* commence avec un nombre petit de groupes et, itérativement, insère des groupes où les données ont un degré d'appartenance faible aux groupes existants. Dans cet esprit, Chiu, [39], a proposé le *groupage soustractif* (subtracting clustering) qui considère chaque point $\mathbf{x}_i$ comme un prototype potentiel en lui associant un potentiel Pot_i , par rapport à son voisinage :

$$Pot_i = \sum_{j=1}^{\ell} k(\mathbf{z}_i, \mathbf{z}_j, \sigma_a). \quad (\text{III.216})$$

avec $k(\cdot)$ la fonction kernel gaussienne (III.76). Le point $\mathbf{x}_k$ avec le plus grand potentiel sera le centre du premier groupe et on actualise le potentiel de chaque groupe :

$$Pot_i = Pot_i - Pot_k k(\mathbf{z}_i, \mathbf{z}_k, \sigma_b), \quad \text{with } \sigma_b \approx 1.5\sigma_a \quad (\text{III.217})$$

A nouveau, le point avec le potentiel le plus grand sera le centre du groupe suivant, et le processus s'arrête jusqu'à ce qu'un certain seuil des potentiels soit atteint. La Figure III.20 présente le résultat du groupage soustractif, appliqué à un ensemble aléatoire de 100 points (gauche). Les quatre centres sont

entourés (milieu) et correspondent à quatre règles ayant comme prémisses des fonctions d'appartenance gaussiennes (droit).

La valeur de l'exposant flou

Ce paramètre influence le caractère flou de la partition résultante. Quand m s'approche de la valeur 1, la partition devient forte ($\mu_{ik} \in \{0,1\}$), et les centres $\mathbf{v}_j$ sont des moyennes ordinaires. Lorsqu'il tend vers l'infini, les centres des groupes tendent vers le centre de gravité des données et la partition devient complètement floue. La plupart des auteurs recommandent une valeur fixe de m , habituellement 1.5 ou 2. Néanmoins, dans [38], Chen propose une optimisation de ce paramètre dans l'algorithme FCM où la fonction d'appartenance associée au groupe c_j est :

$$A_j(x) = 1 / \left(1 + \left(\frac{x - \mathbf{v}_j}{\sigma_j} \right)^{b_j} \right),$$

commençant avec une valeur initiale de $m=1.5$ augmentant par pas de 0.1, jusqu'à ce que, dans chaque dimension de l'entrée, il y ait au moins un groupe dont la valeur σ_j soit supérieure ou égale à l'écart type des données d'apprentissage dans cette dimension.

Les centres mobiles flous

Les *centres mobiles flous* (mean tracking algorithm), [162], sont un algorithme de groupage moins coûteux que celui des FCM. Sur chaque axe x_i de l'espace d'entrée, $i=1, \dots, n$, l'utilisateur prédéfinit une fenêtre, associée à un sous-ensemble flou de forme radiale, de centre $c_i : [c_i - a_i, c_i + a_i]$ où a_i est fonction de l'écart type de la distribution sur l'axe x_i . Itérativement, l'algorithme calcule le centre de gravité de la fenêtre de centre c_i , puis, dans un deuxième mouvement, la fenêtre est placée sur son centre de gravité. On peut alors calculer le nouveau centre de gravité et ainsi de suite, jusqu'à stabilisation.

Chaque fenêtre correspond à un groupe et on peut initialiser aléatoirement autant de fenêtres que l'on veut et, éventuellement, fusionner les groupes associés. Si la convergence est assurée, elle est très dépendante de la valeur initiale du centre. Pour éviter les solutions locales, les statisticiens appliquent l'algorithme plusieurs fois au même jeu de données, pour en obtenir des formes fortes. Il n'est pas prévu de faire évoluer la taille de la fenêtre en fonction des données, car les groupes couvriront le même espace.

III.4.3 Identification des modèles TS pour des systèmes MIMO

III.4.3.1 Structure d'un modèle flou TS pour des systèmes MIMO

Supposons le système non-linéaire MIMO inconnu $\mathbf{y} = f(\mathbf{u})$ à n_i entrées $\mathbf{u}_k = [u_{1k}, \dots, u_{n_{ik}}]^T$ et n_o sorties $\mathbf{y}_k = [y_{1k}, \dots, y_{n_{ok}}]^T$, avec $k=1, \dots, \ell$, respectivement. Suivant la notation suivie en [13], soit q l'opérateur de régression **Error! No se pueden crear objetos modificando códigos de campo.** les polynômes $\mathbf{a}$ et $\mathbf{b}$ en q , $\mathbf{a} = a_0 + a_1 q + a_2 q^2 + \dots$ et les deux nombres entiers, $m < n$; alors, la séquence des échantillons retardés y est définie :

$$\{y(k)\}_m^n \stackrel{\text{def}}{=} [y(k-m), y(k-m-1), \dots, y(k-n)]. \quad (\text{III.218})$$

Le système MIMO est décomposé en plusieurs systèmes MISO, ces derniers sous la forme :

$$y_l(k+1) = \mathcal{F}_l(\mathbf{x}_l(k)), \quad l=1, \dots, n_o, \quad (\text{III.219})$$

où le vecteur de régression est donné par :

$$\mathbf{x}_l(k) = [\{y_1(k)\}_{n_{d1}}^{n_{u1}}, \{y_2(k)\}_{n_{d2}}^{n_{u2}}, \dots, \{y_{n_o}(k)\}_{n_{dn_o}}^{n_{un_o}}, \{u_1(k)\}_{n_{d1}}^{n_{u1}}, \{u_2(k)\}_{n_{d2}}^{n_{u2}}, \dots, \{u_{n_i}(k)\}_{n_{din_i}}^{n_{uin_i}}]. \quad (\text{III.220})$$

Le terme n_y est une matrice $n_o \times n_o$ contenant le nombre d'éléments de retard des sorties, les matrices n_u et n_d de dimension $n_o \times n_i$ contiennent le nombre d'éléments de retard des entrées et le nombre d'éléments de retard pur (transport) de l'entrée vers la sortie respectivement. Ainsi, le $l^{\text{ème}}$ modèle $\mathcal{F}_l$ est représenté par un modèle flou Takagi-Sugeno, sous forme (III.199), comme suit :

$$R_{li}: \text{ si } x_{li1}(k) \text{ est } A_{li1} \text{ et } \dots \text{ et } x_{li}(k) \text{ est } A_{lii} \text{ alors } y_{li}(k+1) = \mathbf{A}_{li} \mathbf{y}(k) + \mathbf{B}_{li} \mathbf{x}(k) + c_{li}, \quad \text{avec } i=1, \dots, k_i, \quad (\text{III.221})$$

où A_{li} sont les ensembles flous de la prémisse de la $i^{ème}$ règle, $\mathbf{A}_{li}=[\mathbf{a}_{li1}, \dots, \mathbf{a}_{li n_o}]^T$ et $\mathbf{B}_{li}=[\mathbf{b}_{li1}, \dots, \mathbf{b}_{li n_i}]^T$ sont des vecteurs de polynômes. k_l est le nombre de règles du $l^{ème}$ modèle.

Dans notre approche, nous choisissons des ensembles flous représentés par des fonctions d'appartenance multidimensionnelles $\omega(\mathbf{x}(k))$; **Error! No se pueden crear objetos modificando códigos de campo.** $\rightarrow [0,1]$, où $p_l = \sum_{j=1}^{n_o} n_{y l j} + \sum_{j=1}^{n_i} n_{u l j} + 1$ est la dimension de l'espace des prémisses :

$$R_{li}: \text{si } \mathbf{x}_l(k) \text{ est } A_{li} \text{ alors } y_{li}(k+1) = \mathbf{A}_{li} \mathbf{y}(k) + \mathbf{B}_{li} \mathbf{x}(k) + c_{li}, \quad \text{avec } i=1, \dots, k_l, \quad (\text{III.222})$$

La sortie du $l^{ème}$ modèle est calculée par la moyenne pondérée (III.179), qui prend la forme :

$$y_l(k+1) = \frac{\sum_{i=1}^{K_l} \omega_{li}(\mathbf{x}_l(k)) y_{li}(k+1)}{\sum_{i=1}^{K_l} \omega_{li}(\mathbf{x}_l(k))}. \quad (\text{III.223})$$

Dans la forme conjonctive, $\omega_{li}(\mathbf{x}(k))$ est le produit de tous les degrés d'appartenance individuels. Dans l'espace des produits $\omega_{li}(\mathbf{x}(k))$ est simplement le degré d'appartenance $\mu_{\omega_{li}}(\mathbf{x}(k))$.

III.4.3.2 Méthode d'identification

Pour commencer, l'utilisateur doit définir la structure du modèle, cela veut dire déterminer les matrices n_y , n_u et n_d , sur la base de la connaissance existante *a priori* sur le système à identifier et/ou en faisant une comparaison parmi plusieurs structures candidates, en termes d'erreur de prédiction ou d'autres critères.

Après la définition de la structure, il faut estimer les paramètres des n_o modèles MISO (III.222), les fonctions d'appartenance des prémisses et les polynômes conséquents. Le nombre de clusters est défini, soit par l'utilisateur, soit automatiquement comme abordé dans la section 3 de ce chapitre, comme la méthode d'identification estime chaque modèle MISO indépendamment des autres, nous supprimons l'indice l , à partir de ce point, pour simplifier la notation.

La méthode d'identification générale est composée de cinq étapes, [12] :

1. A partir des séquences d'entrée-sortie $\{\mathbf{x}(k), \mathbf{y}(k)\}_{k=1}^{n_d}$, former le problème de régression non-linéaire de (III.219), en utilisant les paramètres structure définis par l'utilisateur n_x , n_y et n_d .
2. Trouver le nombre optimal de groupes pour chaque modèle MISO.
3. Partitionner les données en un ensemble de sous-modèles locaux en utilisant un algorithme de groupage dans l'espace de produit cartésien $X \times Y$.
4. Déterminer les fonctions d'appartenance à partir des paramètres des groupes.
5. Une fois les fonctions d'appartenance définies, estimer les paramètres conséquents par la méthode des moindres carrés.

Les données de régression

La matrice de régression $\mathbf{X}$ et le vecteur de régression $\mathbf{y}$ sont construits à partir des séquences de données :

$$\mathbf{X} = \begin{bmatrix} \dots \\ \mathbf{x}(k) \\ \dots \\ \mathbf{x}(\ell-1) \end{bmatrix}, \quad \mathbf{y} = \begin{bmatrix} \dots \\ y(k+1) \\ \dots \\ y(\ell) \end{bmatrix}. \quad (\text{III.224})$$

où $\ell \gg p$ es le nombre données disponibles pour l'identification. L'ensemble des données final à traiter dans l'espace des prémisses est la concaténation :

$$\mathbf{Z}^T = [\mathbf{X} \ \mathbf{y}]. \quad (\text{III.225})$$

Les fonctions d'appartenance des prémisses

Le calcul des fonctions d'appartenance multidimensionnelles des prémisses A_i des règles (III.222), utilise les distances d_{ij} de $\mathbf{z}(k)$ aux groupes, pour ensuite pouvoir calculer le degré d'appartenance :

$$\omega_i(\mathbf{z}(k)) = \frac{1}{\sum_{j=1}^c \left[\frac{d(\mathbf{z}_i, f_i)}{d(\mathbf{z}_i, f_j)} \right]} \quad (\text{III.226})$$

Dans le cas des fonctions d'appartenance A_{ij} , $1 \leq j \leq p$, de la forme conjonctive (III.221), les ensembles flous multidimensionnels définis dans la $i^{\text{ème}}$ ligne de la matrice de partition $\mathbf{U}=[\mu_{ik}]$, peuvent être projetés sur chacune des variables $\mathbf{z}_j$ par :

$$\mu_{\Omega_i}(\mathbf{z}_j(k)) = \text{proj}_j(\mu_{ik}) \quad (\text{III.227})$$

où proj est l'opérateur de projection [148]. Les ensembles flous projetés résultants A_{ij} sont alors approximatés par une fonction paramétrique adéquate, afin d'être également capable de calculer $\omega_i(\mathbf{z}_j(k))$ pour des valeurs $\mathbf{z}_j(k)$ non contenues en Z .

Les paramètres conséquents

Il existe plusieurs possibilités pour calculer les paramètres conséquents d'un modèle flou TS. Le cas le plus simple est l'utilisation de la méthode des moindres carrés. Soit θ_i^T le vecteur contenant les coefficients des polynômes conséquents A_i , B_i et le seuil b_i . Soit la matrice $\mathbf{Z}_e$ la représentation de $[\mathbf{Z} \ 1]$ et $\mathbf{W}_i$ la matrice diagonale en $\mathbb{R}^{p \times p}$ contenant les degrés d'appartenance μ_{ik} comme $k^{\text{ème}}$ élément de la diagonale. Si les colonnes de $\mathbf{Z}_e$ son linéairement indépendantes et $\mu_{ij} > 0$ pour $1 \leq k \leq p$, alors :

$$\hat{\theta}_i = [\mathbf{Z}_e^T \mathbf{W}_i \mathbf{Z}_e]^{-1} \mathbf{Z}_e^T \mathbf{W}_i \mathbf{y} \quad (\text{III.228})$$

est la solution de $\mathbf{y} = \mathbf{Z}_e \theta + \varepsilon$ où le $k^{\text{ème}}$ couple de données $(\mathbf{x}_k, y_k)$ est pondéré par μ_{ik} .

III.4.3.3 Exemple d'identification

Comme exemple d'identification, nous avons choisi un des benchmarks de FAMIMO², qui simule un flux continu d'une lagune aérée pour le traitement des eaux usées dans l'industrie du papier. Ce processus est caractérisé par des dynamiques non stationnaires et non linéaires, [21]. Le bioréacteur contient une population microbiologique, appelé biomasse, qui croît dans un mélange de deux types de substrats polluants, énergétique et xénobiotique.

L'objectif de ce problème de commande multivariable consiste à contrôler, à partir de la relation de dilution de l'eau usée, Figure III.21 gauche, le flux de biomasse dans le bioréacteur de sortie, ainsi que la concentration résiduelle des substrats polluants, énergétique et xénobiotique, Figure III.21 droit. Nous disposons d'une base de données pour l'identification et une deuxième pour la validation, de 1000 échantillons chacune pour une fréquence d'échantillonnage de 1 Hz.

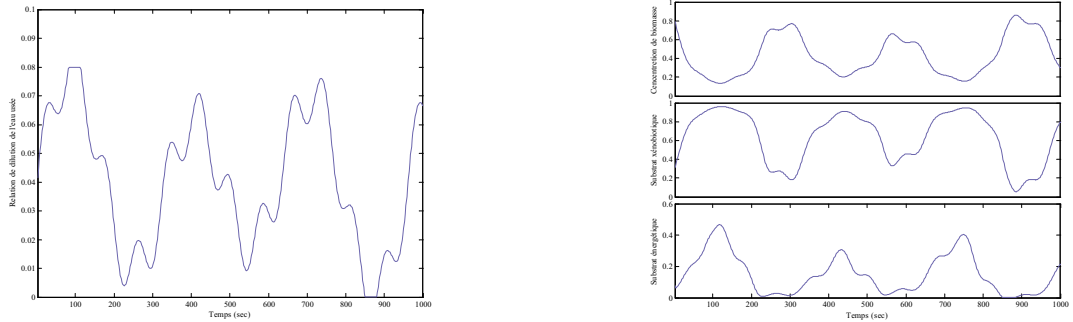


Figure III.21 : Entrée pour l'identification : relation de dilution des eaux usées. Sorties pour l'identification : biomasse, substrat énergétique et substrat xénobiotique.

La structure du modèle est disposée de telle sorte que chaque sortie dépende des autres, comme le montre la matrice n_y . Les vecteurs n_x et n_d spécifient un seul retard sur l'entrée.

$$n_y = \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}, \quad n_u = \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}, \quad n_d = \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} \quad (\text{III.229})$$

² FAMIMO: Fuzzy Algorithms for the control of Multi-Input, Multi-Output Processes. LTR Project 21911 Reactive Scheme.

Pour comparer la performance de notre approche SVM-Flou pour l'identification de règles, nous avons mis en œuvre l'algorithme de groupage flou Gustafson–Kessel, utilisé dans [12] pour l'identification de systèmes flous et l'algorithme de C-moyennes floues. En conséquence, nous avons :

1. Un modèle MIMO flou, FCM-TK, avec un nombre fixe de groupes pour chaque sous modèle, défini *a priori*.
2. Un modèle MIMO flou, GK-TK, avec un nombre fixe de groupes pour chaque sous modèle, défini *a priori*.
3. Un modèle MIMO flou, SVM-TK, utilisant la méthode de groupage soustractif pour calculer le nombre de groupes.

TABLEAU III.12
Comparaison de la précision de prédiction des modèles FCM-TS, GK-TS et SVM-TS (VAF).

	MODELE FLOU FCM-TS	MODELE FLOU GK-TS	MODELE FLOU SVM-TS
Concentration de biomasse	3 groupes	3 groupes	2 groupes, 4sv, 6sv
Substrat xénobiotique	3 groupes	3 groupes	2 groupes, 6sv, 5sv
Substrat énergétique	3 groupes	3 groupes	2 groupes, 4sv, 5sv

Le Tableau III.12 montre les architectures résultantes des trois modèles. Pour chaque sortie, le nombre de vecteurs de support est lié à la complexité du modèle SVM-TS. Un groupe qui contient ns_l vecteurs de support a une matrice de paramètres de taille $(ns_l+1) \times p_l$. Par exemple, pour un groupe avec $ns_l=4$ et $p_l=3$, la matrice de paramètres résultante est de taille 5×3 . Le coût de calcul n'augmente pas considérablement, comparé à la matrice de paramètres de taille 4×3 , correspondant à la matrice de covariance et au centre du prototype, dans les modèles FCM-TS et GK-TS.

Le Tableau III.13 donne les indices de performance *variance accounted for* (VAF) :

$$VAF = 100\% \cdot \left[1 - \frac{\text{var}(\mathbf{y} - \hat{\mathbf{y}})}{\text{var}(\mathbf{y})} \right], \quad (\text{III.230})$$

pour les trois types de modèles construits, en utilisant la base de données d'identification et la base de données de validation. La Figure III.22 présente les sorties pour l'ensemble de validation.

TABLEAU III.13
Comparaison de la précision de prédiction des modèles FCM-TS, GK-TS et SVM-TS (VAF).

		MODELE FLOU FCM-TS	MODELE FLOU GK-TS	MODELE FLOU SVM-TS
Identification	Concentration de biomasse	98.5309 %	99.0082 %	99.7626 %
	Substrat xénobiotique	98.2171 %	99.1278 %	99.7201 %
	Substrat énergétique	96.9935 %	94.5855 %	99.8391 %
Validation	Concentration de biomasse	96.0185 %	96.8809 %	98.9975 %
	Substrat xénobiotique	98.4601 %	98.4731 %	99.5593 %
	Substrat énergétique	97.5849 %	72.5222 %	99.7848 %

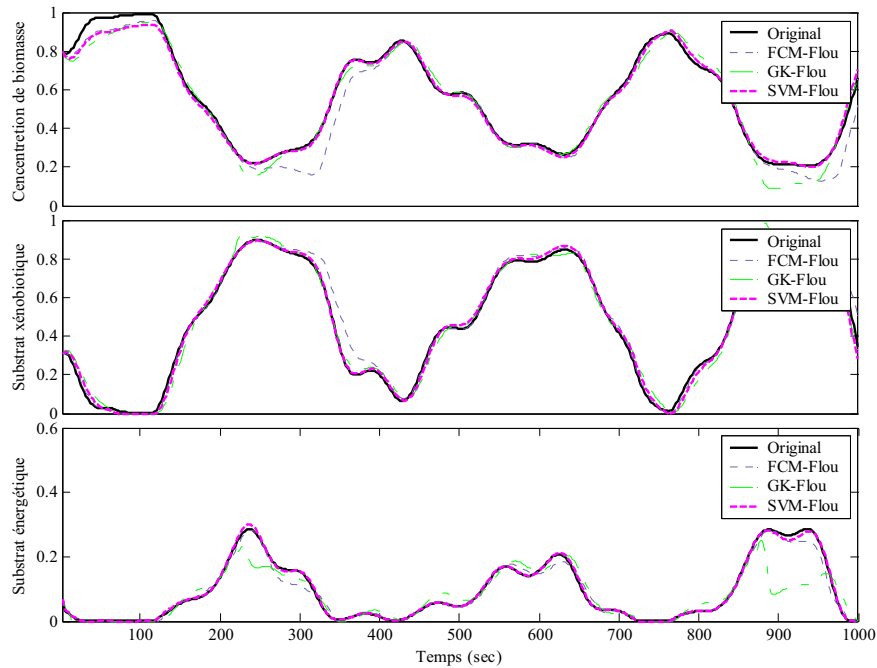


Figure III.22 : Comparaison entre la sortie du processus et les sorties FCM-Flou, GK-Flou et SVM-Flou.

Contrairement à ce que l'on pouvait penser, l'algorithme FCM-TK a, dans certains cas, des meilleurs résultats que l'algorithme GK-TS. Le modèle proposé, SVM-TS, affiche les meilleurs résultats, tant dans la phase d'identification que dans la phase de validation.

III.5 Conclusion

Dans ce chapitre, nous avons étudié différentes approches de modélisation. Selon l'information disponible, la modélisation peut être effectuée par des techniques d'apprentissage (utilisant des observations), de systèmes experts (utilisant de la connaissance) ou les deux (hybrides).

Nous avons abordé le cas des techniques d'apprentissage d'un point de vue neuronal. Après une brève introduction à la théorie de l'apprentissage, nous avons présenté les différentes approches dans un cadre linéaire, pour ensuite passer au cas non-linéaire. Ceci nous a donné les éléments nécessaires pour présenter les machines à vecteurs de support (SVM), approche avec des caractéristiques intéressantes. Nous avons abordé le problème d'apprentissage des SVM en proposant des stratégies d'implémentation plus adaptées à l'algorithme d'optimisation quadratique et à l'algorithme de décomposition pour gérer des problèmes de grande échelle. Nous avons ainsi présenté l'efficacité des implémentations sur des bases des données, largement connues de la communauté scientifique, pour garder ainsi une méthode duale d'ensemble actif qui, comparable aux deux méthodes de point intérieur, présente les meilleures performances.

Dans le cadre de la modélisation à partir de la connaissance, nous avons présenté les systèmes d'inférence flous en décrivant leur structure, les bases pour leur construction et les différents types de modèles que l'on peut construire avec ce type de techniques.

La partie consacrée à la modélisation à partir des données et l'expertise nous a permis d'aborder deux grandes schémas : les modèles flous conventionnels dotés d'un mécanisme d'apprentissage et les modèles flous représentés sous une architecture de réseau de neurones. Dans ce dernier cas, nous avons développé un modèle hybride SVM-flou pour le groupage et nous avons comparé ses performances avec d'autres techniques à travers un exemple d'identification de modèles TS pour des systèmes MIMO.

Chapitre IV

Le problème de la surveillance du conducteur automobile. Les projets AWAKE et PREDIT

Sommaire

IV.1	Introduction.....	107
IV.1.1	Les enjeux de la sécurité routière.....	107
IV.2	Les systèmes d'assistance à la conduite.....	109
IV.2.1	Taxonomie des systèmes d'assistance à la conduite	109
IV.2.3	La tâche de conduite.....	110
IV.2.3.1	Les capteurs de position latérale et le contrôle associé	111
IV.2.3.2	Les capteurs de distance longitudinale et le contrôle associé.....	112
IV.3	La surveillance du conducteur automobile	112
IV.3.1	Généralités	112
IV.3.1.1	Mesures directes du niveau de vigilance	112
IV.3.1.2	Mesures directes de la performance de la conduite.....	114
IV.3.2	Bilan sur les systèmes de surveillance du véhicule/conducteur	116
IV.3.2.1	Quelques systèmes de surveillance de la conduite et de l'état du conducteur.....	117
IV.3.2.2	L'analyse synthétique.....	124
IV.4	Les projets de recherche	125
IV.4.1	Le programme Européen AWAKE (Septembre 2001 – Septembre 2004)	126
IV.4.2	Le projet national Hypovigilance du Conducteur – PREDIT (Mars 2002 – Mars 2004).....	126
IV.4.3	Architectures de fonctionnement	127
IV.4.3.1	L'architecture AWAKE.....	127
IV.4.3.2	L'architecture PREDIT	128
IV.5	Les essais	129
IV.5.1	Moyens expérimentaux	129
IV.5.2	Les campagnes d'essais.....	131
IV.5.2.1	Les essais sur démonstrateurs.....	131
IV.5.2.2	Les essais sur simulateurs.....	133
IV.6	Développement du module de diagnostic de situations à risque	134
IV.6.1	Architecture du module de diagnostic de situations à risque	134
IV.6.2	Module de Diagnostic Événementielle basé sur le temps de sortie de voie.....	135
IV.6.3	EDM basé sur vibreur de bord de route adaptatif	135
IV.6.4	Comparaisons des différents approches EDM	138
IV.7	Détection de l'hypovigilance du conducteur automobile	141
IV.7.1	Architecture du module de détection de l'Hypovigilance.....	142
IV.7.1.1	Architecture du HDM performance.....	142
IV.7.1.2	Analyse temporelle et fréquentielle des signaux	143
IV.7.1.3	Fusion des caractéristiques.....	146

IV.7.1.3	Diagnostic cumulé.....	147
IV.7.2	Evaluation du HDM performance.....	148
IV.7.2.1	Mesures de référence de la vigilance.....	148
IV.7.2.2	Validation.....	149
IV.7.2.3	Résultats	151
IV.8	Conclusion	162

IV.1 Introduction

Ce chapitre présente l'application des éléments et outils pour surveillance des systèmes Homme-Machine, abordés au cours des premiers trois chapitres, à la conduite automobile. Celle-ci est une tâche exigeante des actions et des modifications de contrôle multiples, déterminées par l'environnement dynamique dans lequel elle intervient.

Dans ce chapitre, nous décrivons brièvement les moyens intégrant la supervision du système Homme-Véhicule : les systèmes d'assistance à la conduite, qui ont pour but d'aider le conducteur à mieux comprendre son environnement et l'assister dans sa prise de décision. Ensuite, nous décrivons un modèle générique de la tâche de conduite pour aborder le problème de la perception et décrire les capteurs utilisés. Ensuite, nous décrivons les facteurs qui interviennent dans la surveillance du conducteur automobile et les types de mesures que l'on peut en effectuer et nous présentons quelques systèmes de surveillance du conducteur existants.

Nous décrivons les projets de recherche, AWAKE et PREDIT, dans lesquels ces travaux se situent, avec le but de concevoir un système de prévention des accidents potentiels. Nous présentons l'architecture suivie par chaque projet, en remarquant que nous suivons deux stratégies différentes. Puis, nous décrivons les campagnes d'essais menés afin de développer un ensemble de bases de données en vue de valider les principes du système de diagnostic. Nous présentons les démonstrateurs et les simulateurs faisant parti des moyens techniques de ces essais.

Après cela, nous présentons ensuite nos développements, suivant une architecture de surveillance basée sur deux types d'avertissements : une donné par un module de diagnostic de situations à risque et autre module pour le diagnostic de l'hypovigilance.

Dans un premier temps, nous développons une architecture spécifique pour le module de diagnostic de situations à risque, nous abordons la détection d'événements de sortie de route. Cette approche permet de détecter quand le conducteur est sur le point de sortir de la route et déclencher une alarme pour l'avertir. Nous proposons trois stratégies pour sa mise en œuvre et nous les comparons en utilisant les bases des données des essais.

Dans un deuxième temps, nous décrivons notre approche pour la mise en œuvre du module de diagnostic de l'hypovigilance à partir de mesures de performance de la conduite. Sous une approche multisensorielle permettant d'enrichir l'information disponible pour un diagnostic plus précis. Nous présentons une architecture particulière basée sur une analyse temporelle et fréquentielle des signaux pour, ensuite, en extraire des caractéristiques statistiques et fréquentielle représentant la conduite et finalement fusionner ces dernières afin de calculer un diagnostic final. Nous utilisons les bases de données des essais pour évaluer ce système en termes de *fausses alarmes* et de *non détections*, à partir d'un ensemble de mesures de référence physiologiques et subjectives.

IV.1.1 Les enjeux de la sécurité routière

De tous les moyens de transport, le transport routier est le plus dangereux et le plus coûteux en termes de vies humaines. Ce n'est que très récemment que les accidents de la route ont éveillé de fortes réactions au niveau social [52]. Comment expliquer l'acceptation relative des accidents de la route quand, chaque jour, le nombre total de personnes tuées sur la route, en Europe, est pratiquement le même que dans un crash d'avion moyen ?

Pour souligner l'importance des accidents de la route, rappelons que, seulement, en 2000 les accidents de la route ont tué plus de 40.000 personnes et en ont blessé plus de 1,7 millions dans la

Communauté Européenne. Le groupe d'âge le plus affecté est les 14–25 ans pour lesquels les accidents de la route représentent la première cause de décès. Une personne sur trois sera blessée à un moment de sa vie. Le coût directement mesurable des accidents de la route est de l'ordre de 45 milliards d'euros, tandis que le coût indirect (incluant les dégâts physiques et psychologiques aux victimes et leurs familles) est trois ou quatre fois plus important, mais le coût humain reste incalculable. Le montant annuel est de 169 milliards d'euros, équivalent à 2% du GNP¹ de l'Union Européenne, [52].

Durant ces quatre dernières décennies, les conditions de conduites se sont améliorées grâce aux, campagnes d'éducation et d'information, à la définition de standards d'équipement plus sûrs et à la amélioration des autoroutes. Néanmoins, l'erreur humaine reste la première cause d'accident de la route, citée dans plus de 80% de rapports d'accidents routiers [84].

Dans le cadre des programmes de recherche et développement des systèmes de transport intelligents (intelligent transportation systems, ITS), la sécurité routière est un thème prioritaire. Les ITS permettent de changer le paradigme, de l'aide à la survie des occupants, lors d'un accident, à l'assistance au conducteur. La conception de systèmes embarqués dans les véhicules, les améliorations de l'infrastructure et les systèmes coopératifs véhicule–infrastructure cherchent à assister le conducteur pour éviter des manœuvres dangereuses en minimisant la distraction, en aidant dans des conditions de conduite dégradées et en avertissant (ou en contrôlant) dans des situations d'accident imminent.

L'amélioration des conditions de conduite requiert un niveau de coopération important entre tous les secteurs de l'économie :

- **Les constructeurs automobiles et leurs sous-traitants** doivent travailler ensemble pour développer des produits destinés à rendre le véhicule plus sûr à travers des systèmes électroniques embarqués, avec un temps de réaction faible, capable de compenser les erreurs du conducteur, à travers aussi des bases de données sur cartes numériques pour améliorer, à court terme, les systèmes de navigation et, à long terme, l'ensemble des systèmes à bord.
- **Les sociétés de gestion de l'infrastructure et leurs sous-traitants** se sont engagés à fournir des systèmes d'information et de contrôle du trafic et à créer des réseaux pour interconnecter ces systèmes à des centres de contrôle de trafic, eux-mêmes interconnectés entre eux.
- **Le gouvernement, en collaboration avec des organismes publics et privés**, doit financer la recherche pour mieux comprendre le comportement du conducteur, assurer la sécurité des systèmes d'information et améliorer l'efficacité des technologies avancées de prévention d'accidents.

En France, le Programme national de recherche et d'innovation dans les transports terrestres PREDIT – Ministère des Transports, en Europe, le programme sur les technologies d'information pour la société (information society technologies, IST) – Commission Européenne, aux Etats Unis, la société des transports intelligents d'Amérique (intelligent transportation society of America, ITS), entre autres, ont concentré leur attention et leurs ressources pour évaluer l'efficacité et l'opérabilité des technologies avancées de prévention d'accidents pour les véhicules privées, commerciaux et de transit, à travers plusieurs groupes thématiques de recherche.

Ces efforts ont conduit au développement de plusieurs technologies, divisées en deux grands types de systèmes électroniques embarqués sur véhicule :

- **Systèmes de diagnostic–pronostic**, qui surveillent le fonctionnement des différents composants mécaniques de la voiture : pression d'air des roues, surveillance du moteur, sécurité et stabilité de la cargaison, etc.
- **Systèmes d'assistance à la conduite**, qui aident le conducteur à mieux comprendre et interpréter son environnement et l'assistent dans sa prise de décision : avertissement et contrôle de stabilité de trajectoire, contrôle adaptatif de croisière, etc.

¹ Croissance du produit intérieur brut réel (growth in real per capita, GNP).

Puisque l'intérêt de cette thèse est la surveillance des systèmes homme-machine, nous nous limitons à l'étude des systèmes d'assistance à la conduite. Pour plus d'informations sur les systèmes de diagnostic-pronostic voir [109].

IV.2 Les systèmes d'assistance à la conduite

Dans l'activité de conduite, l'opérateur humain (conducteur) agit comme un contrôleur central qui exécute différentes tâches de stabilisation, de guidage et de navigation. Il constitue une partie relativement sûre et fiable des systèmes homme-machine grâce à sa capacité d'adaptation à une grande variété de situations et à la sophistication de ses capteurs. Néanmoins, les erreurs humaines ne peuvent être totalement exclues car les limitations en perception, traitement et actions sont rapidement atteintes. La voiture de demain doit aider le conducteur à mieux comprendre son environnement et l'assister dans sa prise de décision sans pour autant le déresponsabiliser, Figure IV.1.

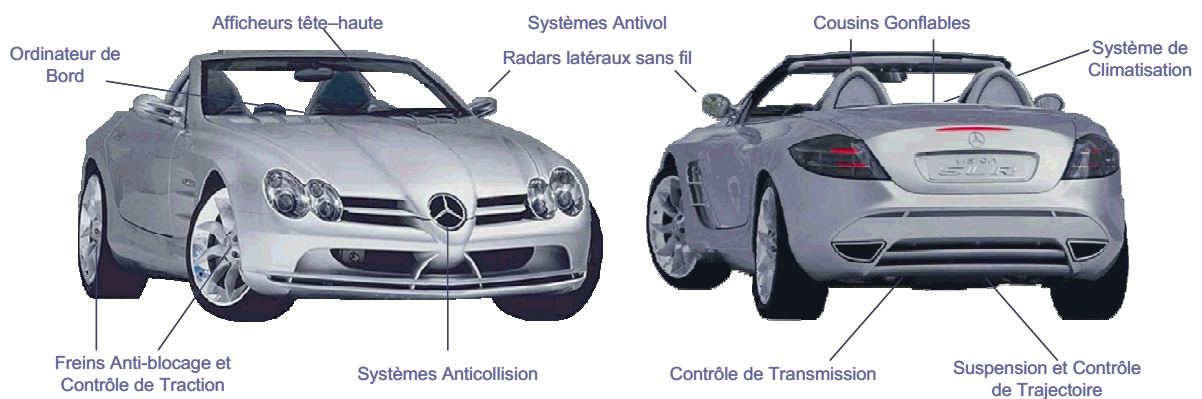


Figure IV.1 : L'automobile et quelques sous-systèmes.

IV.2.1 Taxonomie des systèmes d'assistance à la conduite

Les systèmes d'assistance à la conduite disposent de trois niveaux d'application, allant du système le plus simple (uniquement informatif) au système plus complexe de prise en charge complète du véhicule.

- **Systèmes informatifs**, qui n'ont pas d'intervention directe sur le véhicule et permettent d'avertir lorsque l'interaction homme-véhicule peut devenir dangereuse. Des exemples sont : rétrovision pour détection dans l'angle mort, détection d'obstacles et avertissement de risque assisté par infrastructure, communication véhicule-véhicule ou véhicule-infrastructure, détection prédictive de sortie de route, système d'aide au parking, signalisation dans le véhicule : mécanismes alternatifs pour avertir le conducteur comme les afficheurs tête-haute, les avertissements sonores, visuels et sensitifs, guidage sur route et Information sur trafic : utilisation de la radio-navigation, incluant les systèmes de positionnement global (US global positioning system, GPS) et les services de navigation globale (EU global navigation services, GALILEO), détection de la vigilance du conducteur, déduction des actions du conducteur ou des mesures physiologiques indirectes.
- **Systèmes actifs**, qui interviennent uniquement en cas de détection de danger. Un exemple bien connu est l'ABS qui ne devient actif que lorsque les roues se bloquent, ce qui améliore la tenue de route en cas de freinage d'urgence. D'autres systèmes en développement sont : l'avertissement de vitesse de courbe, l'avertissement de risque de collision, l'avertissement de sortie de voie, le contrôle adaptatif de croisière, le contrôle de traction, le contrôle de stabilité de trajectoire, les coussins gonflables prédictifs, les systèmes anticollision longitudinale et les systèmes de freinage électroniques.
- **Systèmes de contrôle automatique**, entre sécurité et confort, ils comprennent le contrôle adaptatif de croisière intégrant des fonctions anticollision et de communication, le contrôle longitudinal

urbain (stop & go), le contrôle transversal routier, l'évitement de collision avant et arrière, l'évitement de collision aux intersections et les préventions de sortie de voie.

TABLEAU IV.1
Taxonomie des systèmes d'assistance à la conduite.

SYSTEMES INFORMATIFS	SYSTEMES ACTIFS	SYSTEMES DE CONTROLE AUTOMATIQUE
• Rétrovision pour la détection dans l'angle mort • Détection d'obstacles et avertissement de risque assisté par infrastructure • Communication véhicule-véhicule ou véhicule-infrastructure • Détection prédictive de sortie de voie/route • Système d'aide au parking • Signalisation dans le véhicule : mécanismes alternatifs pour avertir le conducteur, comme les afficheurs tête-haute, les avertissements sonores, visuels et sensitifs • Guidage sur route et Information sur trafic : utilisation de la radio-navigation, incluant les systèmes de positionnement global (US global positioning system, GPS) et les services de navigation globale (EU global navigation services, GALILEO) • Détection de la vigilance du conducteur, déduite des mesures de performance ou des mesures physiologiques indirectes	• Avertissement de vitesse de courbe • Avertissement de collision • Avertissement de sortie de voie/route • Contrôle adaptatif de croisière • Contrôle de traction • Contrôle de stabilité de trajectoire (ESP) • Coussins gonflables prédictifs • Système anticollision longitudinal • Systèmes de freinage électroniques, (ABS)	• Contrôle adaptatif de croisière intégrant des fonctions anticollision et de communication • Contrôle longitudinal urbain (stop & go) • Contrôle transversal routier • Evitement de collisions avant et arrière • Evitement de collisions dans des intersections • Préventions de sortie de voie

Quelques-uns de ces systèmes sont individuellement commercialisés sur le marché, mais, pour obtenir un bénéfice plus tangible de cette révolution des systèmes électroniques, il faut obtenir des progrès significatifs dans l'intégration des systèmes, dans la compréhension actuelle des facteurs humains et dans les pratiques et le fonctionnement des institutions qui, actuellement, obstruent le déploiement rapide et élargi des produits.

IV.2.3 La tâche de conduite

La tâche primaire dans la conduite automobile consiste à contrôler le véhicule sans danger, dans des conditions environnementales et de trafic dynamiques, [41], [77]. Le modèle de la tâche de conduite n'est pas facile à trouver à cause des caractères dynamiques et non stationnaires du trafic, de l'environnement et des comportements adaptatifs du conducteur, devant réagir aux contraintes imposées par le trafic et l'environnement. Cependant, Michon, [100], a proposé un modèle de conduite suivant une stratégie de décomposition fonctionnelle à trois niveaux, Figure IV.2. Le niveau le plus haut représente le niveau stratégique, auquel les décisions de navigation sont prises, c'est-à-dire le choix de l'itinéraire. Au niveau intermédiaire, la manœuvre ou les actions du conducteur répondent à des situations instantanées : la dynamique des voitures proches, la vérification des rétroviseurs ou la décision de doubler ou faire demi-tour. Le niveau le plus bas est celui du contrôle, caractérisant les automatismes du conducteur qui s'occupent de changer les vitesses, de contrôler la trajectoire du véhicule et de surveiller le trafic.

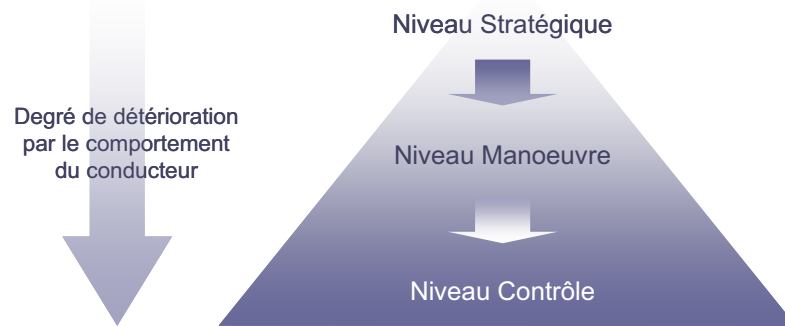


Figure IV.2 : Modèle fonctionnel de conduite automobile.

Dans cette thèse, nous limitons notre étude au niveau du contrôle car il s'agit du niveau le plus affecté par la détérioration du comportement du conducteur. C'est à ce niveau que l'on peut étudier les variations des dynamiques de la voiture, résultat des actions de commande de la part du conducteur.

IV.2.3.1 Les capteurs de position latérale et le contrôle associé

Parmi les tâches de stabilisation, le maintien du véhicule sur la voie (lane control) est celle qui demande le plus d'effort et de temps au conducteur. Même si cette tâche est facile à apprendre et peut être effectuée après une courte période d'apprentissage, elle reste une contrainte considérable, [77].

Afin de concevoir un système de surveillance/contrôle de la trajectoire d'un véhicule, la détection de la position du véhicule sur la voie de circulation devient indispensable. Cette localisation peut être effectuée de différentes façons, [44] :

- Détections des marquages inter-voies ou de bord de route,
- Lecture des bandes magnétiques collées sur les marquages inter-voies,
- Détection d'un fil ferreux non magnétique et non inductif.
- Localisation par GPS différentiel,

Les systèmes de vision par ordinateur constituent la technique de détection des marquages au sol la mieux adaptée pour l'infrastructure routière actuelle. Il s'agit d'une technique largement abordée et qui semble bien maîtrisée dans le cas de bonne visibilité et de bon marquage de la route. Néanmoins, la détection par temps de pluie, neige, orage, nuit, etc. s'avère difficile. L'un des systèmes les plus robustes à de telles conditions est le système de l'Université Carnegie Mellon et de la société AssistWare Technologies Inc., [9], Figure IV.8.

Les expériences menées aux Etats-Unis par le programme de recherche California PATH, pour la lecture des bandes magnétiques collées sur les marquages inter-voies, ont donné des résultats intéressants, [140], mais le fonctionnement de ce type de systèmes est restreint aux autoroutes équipées de marquage magnétique, Figure IV.7.

L'Université de Ohio State a proposé un système embarqué utilisant un radar pour détecter des fils ferreux et, donc, en déduire la position du véhicule sur la voie, [127]. Ceci semble être une technologie intéressante car des essais ont montré l'efficacité d'un suivi, avec un algorithme relativement simple, pour des véhicules roulant environ à 50km/h.

L'Université de Minnesota a démontré la capacité du système GPS différentiel à fournir une localisation précise après une longue phase d'optimisation du paramétrage (plusieurs minutes), [3]. Ce point devra être amélioré pour prétendre à une utilisation en situation réelle.

Pour contrôler la dynamique latérale du véhicule, il existe des modèles de commande sophistiqués, [4], [116], que nous n'allons pas développer car notre objectif est la surveillance des conducteurs

IV.2.3.2 Les capteurs de distance longitudinale et le contrôle associé

Le contrôle longitudinal est une autre tâche centrale pour le guidage du véhicule. Il concerne la surveillance et la régulation de la vitesse d'un véhicule en tenant en compte du tracé de la route et de la vitesse des véhicules précédents ou obstacles. Le conducteur doit surveiller le trafic et déterminer les distances, la vitesse de son véhicule et les vitesses relatives avec suffisamment de précision et de régularité. Ceci requiert un grand niveau de vigilance et focalise toute l'attention du conducteur. Des capteurs, combinés à des systèmes d'assistance, peuvent aider le conducteur dans le contrôle de la distance et de la vitesse.

Le contrôle longitudinal peut utiliser diverses technologies comme des télémètres laser ou des radars avec une portée de 150 à 200 mètres, en combinaison avec des algorithmes sophistiqués de suivi et de classification d'obstacles, [44]. Des techniques utilisant la vision par ordinateur sont déjà disponibles et continuent à se développer en tant que dispositif unique ou en combinaison avec des télémètres laser ou des radars, [29]. En général, le radar est préféré car il est peu sensible aux conditions climatiques. L'objectif consiste à mesurer la distance qui sépare l'avant du véhicule contrôlé et l'arrière du véhicule précédent, et la vitesse à laquelle les deux véhicules se rapprochent, afin de régler la bonne vitesse du véhicule contrôlé. Cette information est ensuite traitée pour alerter le conducteur en cas de situation critique. Les stratégies d'avertissement utilisent des icônes couleur et/ou des pictogrammes.

Une application directe de cette technologie est le régulateur de vitesse adaptatif (adaptive cruise control, ACC) qui couple un calculateur à l'accélérateur et aux freins, permettant à la voiture de réguler elle-même la distance de sécurité qui doit la séparer de la voiture précédente. Evidemment, le conducteur conserve toujours la maîtrise du véhicule, par l'usage de l'accélérateur et du frein, et décide la mise en route de ce système d'assistance à n'importe quel moment, pendant une conduite sur autoroute.

IV.3 La surveillance du conducteur automobile

IV.3.1 Généralités

La surveillance de la vigilance du conducteur au volant visant à la détection de la dégradation de celle-ci, afin d'avertir le conducteur ou de prendre des mesures d'extrême urgence (arrêt total du véhicule), sont des fonctions très attendues pour les véhicules de demain. Il s'agit, en effet, d'un enjeu très important en matière de sécurité primaire, domaine de la sécurité routière qui vise à éviter l'accident. Le développement de tels systèmes représente actuellement un véritable défi pour les constructeurs automobiles, tant les difficultés sont grandes dès lors que l'être humain est au coeur du système que l'on cherche à mettre au point.

Il existe plusieurs facteurs, internes et externes, qui font diminuer le niveau de **vigilance** dans l'exercice de la conduite automobile. Comme facteurs internes nous avons : la fatigue, l'inattention, le sommeil, la prise de médicaments, l'alcool et les drogues, [105]. Comme facteurs externes : la monotonie favorisée par une densité de trafic réduite, une chaussée rectiligne et, en général, le manque de stimuli extérieurs pour maintenir éveillé le conducteur, [129].

IV.3.1.1 Mesures directes du niveau de vigilance

La baisse de vigilance engendre chez le conducteur automobile des comportements corporels observables. Ces activités comportementales (**tâches secondaires**) collatérales à la conduite (**tâche principale**) peuvent être groupées en cinq classes [29], [129] :

1. Variations posturales,
2. Echanges verbaux, recherche de dialogue, échanges d'informations non liées à la tâche,
3. Activités ludiques impliquant la manipulation d'objets,

4. Gestes autocentrés d'une ou deux mains vers le corps,
5. Autres activités non verbales détectables dans l'expression du visage.

Ces activités augmentent progressivement avec la **durée** de la tâche et aussi avec la monotonie car elles sont un mécanisme de défense du sujet qui lutte pour rester éveillé (nous voyons ces activités comme des tâches secondaires que le conducteur exécute pour maintenir son niveau de performance et de vigilance). Cependant, elles ne sont pas toujours efficaces et, au-delà d'un certain seuil de fatigue, plus la performance se dégrade plus ces activités comportementales croissent.

Néanmoins, la **perception** de ces activités comportementales, pour la conception d'un système embarqué, est difficile à réaliser. Le principal problème est l'accès à des informations caractéristiques de ces activités comportementales. Différents capteurs peuvent être envisagés pour cela, [29] :

- *Capteur de pression sur le volant.* Il s'agit d'un ruban adhésif intégrant un capteur de pression installé du côté droit et du côté gauche du volant. L'hypothèse est que la variation de la pression en fonction du temps peut être liée à l'état de vigilance.
- *Capteur de posture sur le siège.* De nombreux véhicules sont maintenant équipés en série de capteurs « nappe » intégrés dans le siège, permettant d'estimer le poids du passager ainsi que la distribution de ce poids sur le siège. En fait, ces nappes comportent de nombreux petits capteurs répartis sur toute leur surface qui permettent une mesure distribuée (encre résistive, capteurs capacitifs, . . .).
- *Systèmes de vision stéréoscopique.* Capteur plus sophistiqué consistant à reconstituer en trois dimensions l'occupation volumique du conducteur à l'intérieur de l'habitacle, déjà commercialisé par SIEMENS pour l'activation de coussins gonflables (airbags).

L'évaluation de l'efficacité de ces capteurs dans la perception de ces activités comportementales est en cours de développement dans le cadre du projet PREDIT, [29], en ce qui concerne le capteur du pression sur le volant, et dans le cadre du projet AWAKE, [7], pour les deux autres capteurs.

Le capteur de mouvement des paupières

Actuellement, la seule activité comportementale abordée dans la conception d'un système temps réel est l'analyse visuelle du battement des paupières, témoin quasi direct de l'état physique du conducteur, [134]. Ainsi, a été développé un système de vision des mouvements des paupières, système de vision non contraignant permettant la mesure d'ouverture des yeux d'un conducteur de véhicules routiers. L'analyse de l'occurrence et de la durée des clignements fournie par le capteur de mouvement des paupières (eye lid sensor, ELS) peut conduire à un diagnostic de somnolence. SIEMENS, l'ENSEEIH et le LAAS ont développé, il y a environ cinq ans, une version du capteur ELS. Ce prototype a été testé, sur simulateur, dans le cadre du programme européen SAVE et lors d'essais sur route, dans des conditions de faible luminosité, [77]. La version actuelle de ce système permet un fonctionnement temps réel (50 Hz). C'est sûrement une voie très riche même s'il reste des problèmes de mise en œuvre, liés aux variations de l'éclairage du visage du conducteur, aux mouvements de sa tête, au port de lunettes, etc. L'intérêt de cette technique fait qu'elle bénéficie d'efforts de recherche et développement très importants pour la production de systèmes de vision, devenus aujourd'hui très performants à des coûts raisonnables.

Les travaux menés par Wierville, [147], ont permis de dégager trois **mesures principales** dérivées des clignements longs :

- PERCLOS est la somme de la durée des clignements longs, par période de une ou trois minutes.
- EYEMEAN est la durée moyenne des clignements longs, par période d'une ou trois minutes.
- EYEMEAS est la moyenne du carré des durées des clignements, par période de une ou trois minutes.

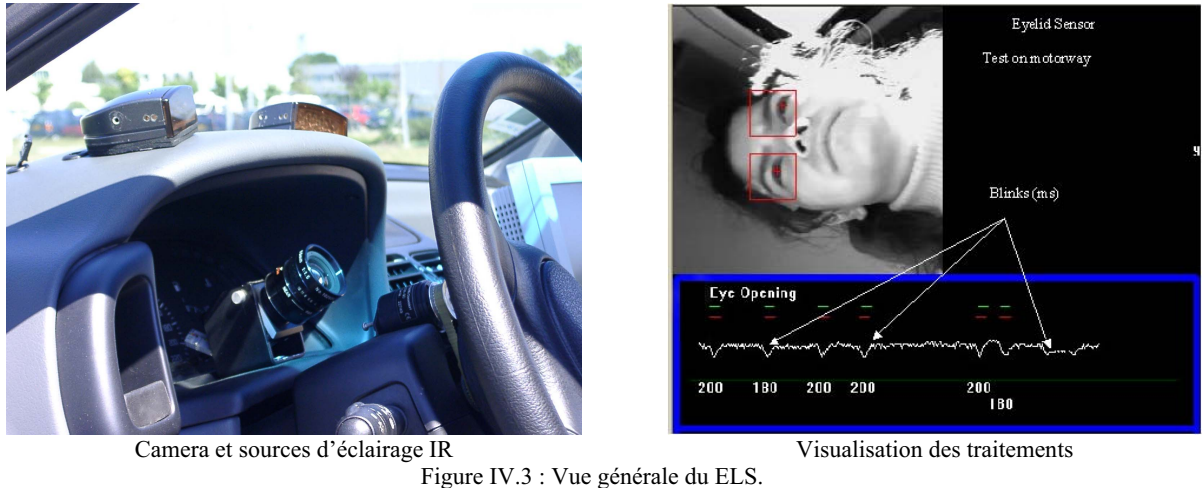


Figure IV.3 : Vue générale du ELS.

Ces mesures ont montré une bonne corrélation avec les performances de conduite et l'activité électro-encéphalographique sur des périodes de six minutes. PERCLOS est actuellement considéré comme une mesure de référence de la somnolence. Par contre, la non détection d'un clignement influence fortement sur la mesure. Des mesures alternatives qui combinent fréquence et durée des clignements sont explorées. Nos partenaires de la société Siemens-VDO ont proposé une mesure à trois niveaux (dont le niveau 0 : vigilance), sur des périodes de trois minutes, glissant par pas d'une minute, [29] ; le niveau 1 de somnolence est déclenché dès l'apparition d'un certain nombre de clignements longs ; le niveau 2 est déclenché à l'apparition de fermetures des paupières. Cette approche est directement issue de la méthode d'analyse des EEG et EOG et déclenche un niveau d'alerte de somnolence dès l'apparition de clignements longs.

Les mesures physiologiques utilisées pour la validation

Les signaux d'activité cérébrale (EEG) sont probablement les signaux les plus fiables pour caractériser l'état de **vigilance** ou de sommeil. Pourtant, leur utilisation dans le cadre de la conduite automobile ne semble pas envisageable à cause des difficultés techniques de mesure d'EEG dans le véhicule, à caractère fortement intrusif et à la présence de plusieurs phénomènes physiologiques jouant un rôle décisif et nécessitant l'avis d'un expert, [7]. Il faut toutefois préciser qu'il n'y a pas de relation biunivoque établie entre les caractéristiques de ces signaux et l'état d'endormissement à cause d'une grande variabilité intra individuelle et aux décalages temporels entre les apparitions des signaux cérébraux typiques du sommeil et l'endormissement, [134]. Nous utilisons, cependant, cette métrique, quand elle est disponible, pour valider notre approche.

IV.3.1.2 Mesures directes de la performance de la conduite

Côté véhicule, l'activité de conduite est représentée par les dynamiques des signaux « mécaniques » reflétant l'activité du conducteur sur les actionneurs (mouvements du volant, position du véhicule sur la voie, vitesse du véhicule, etc.). C'est à partir de ces mesures que l'on peut détecter une dégradation de la performance de conduite, liée au niveau de vigilance.

Mouvements du volant et variabilité de l'angle du volant

Plusieurs travaux ont montré que la variabilité de l'angle du volant est un indicateur de somnolence chez le conducteur, [17], [41], [167], après un prétraitement adéquat pour enlever les effets dépendant de la route. Ce prétraitement peut être réalisé, par exemple, en enlevant la moyenne de la variabilité de l'angle du volant associée à chaque mile pour réduire la variation associée aux courbes.

Parmi les mesures extraites des variations de l'angle du volant ayant, d'une certaine manière, une corrélation avec la somnolence, se trouvent : l'**écart type** de l'angle du volant (*stsw*), la **variance** de la vitesse de l'angle du volant, etc., [41], [77]. En [167], les auteurs recommandent l'utilisation de la **densité de la puissance spectrale** de l'angle du volant en relation à la fatigue. Ces études ont suggéré l'introduction d'autres mesures statistiques et fréquentielles, [29], [133], [134].

L'introduction des **micro-corrections** sur le volant est due aux facteurs environnementaux comme des petits défauts de la route et le vent. Le conducteur tend à réduire le nombre de micro-corrections de l'angle du volant quand la somnolence augmente. L'indice VHAL est une mesure utilisée par certains auteurs pour caractériser les variations de l'angle du volant, [28]. Elle correspond à la relation entre le nombre de mouvements rapides et le nombre de mouvements lents (MH/SH), et utilise le carré de la première dérivée de l'angle du volant, Figure IV.4. L'idée est qu'un conducteur choisit une stratégie de conduite simple dans laquelle il compense seulement les grandes déviations de la trajectoire du véhicule. Le VHAL diminue quand la fatigue augmente.

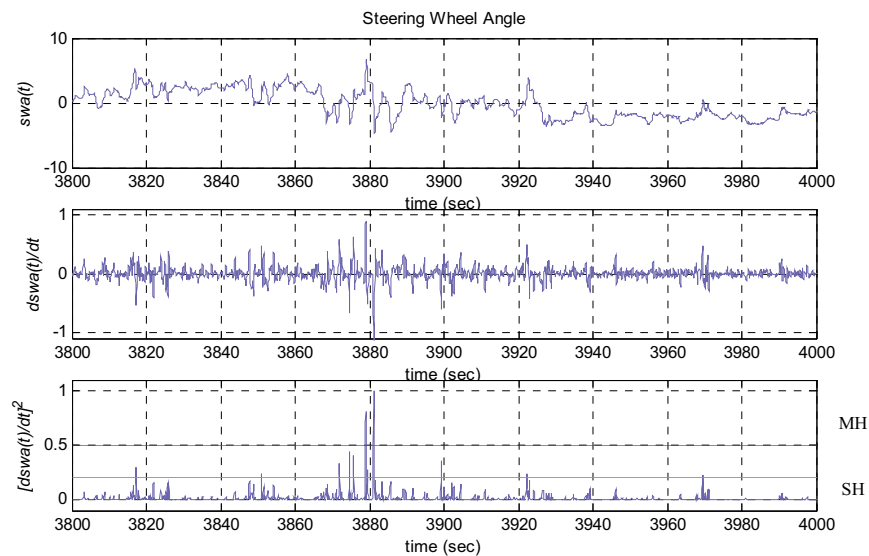


Figure IV.4 : Extraction de l'indice VHAL. SH représente la zone de mouvements lents et MH représente la zone de mouvements rapides

Position Latérale

Dingus et son équipe, [46], ont testé la corrélation entre différents indicateurs de la dégradation des performances du conducteur tels que la fermeture des yeux et la position latérale. Leur résultat montre que les franchissements des lignes de bord de la route ont les corrélations les plus importantes avec la mesure PERCLOS. Le Tableau IV.2 présente les résultats de cette étude.

TABLEAU IV.2

Corrélation entre les mesures de l'activité oculaire et les mesures de la position latérale en relation avec la dégradation des performances du conducteur. (LANEX: nombre d'échantillons durant le franchissement, LANEDEVV : **variance** de la position latérale, LANEDEVSQ : variance pondérée de la position latérale (poids plus important si plus loin du centre de la voie par une fonction au carré), LANEDEV4 : variance fortement pondérée de la position latérale (poids plus important si plus loin du centre de la voie par une fonction du quatrième ordre))

	EYEMEAN	EYEMAS	PERCLOS
LANEX	0.47	0.54	0.62
LANEDEVV	0.50	0.55	0.60
LANEDEVSQ	0.55	0.59	0.60
LANEDEV4	0.36	0.40	0.40

En [17], [41] et [167], il a été conclu que la variabilité du maintien du véhicule (lane tracking), connu comme l'**écart type** de la position latérale (*stlp*), est un indicateur valide de performance du conducteur et de fatigue, car lié la probabilité de franchissement de la route. Comme pour l'angle du volant, ces informations ont suggéré l'étude d'autres grandeurs statistiques et fréquentielles, [134].

Le temps de sortie de voie

Godthelp, [63], a proposé la mesure du temps de sortie de voie (time to lane crossing, *tlc*), représentant le temps prédit pour qu'un véhicule, sur la trajectoire actuelle, quitte la route. Le *tlc* est une mesure que l'on peut classer de **performance tâche primaire**, très importante pour évaluer la performance de la conduite, [41]. Néanmoins, il est difficile de le calculer dans des conditions réelles de conduite, d'où diverses approximations proposées pour les besoins du traitement temps-réel, [165].

Dans un étude de Verwey et Zaidel, [158], ont trouvé que le *tlc* peut être utilisé pour prédire la baisse de la performance du conducteur. Le *tlc* minimum peut indiquer la dégradation progressive de la performance et peut être utilisé pour avertir le conducteur de la détérioration de sa performance, avant qu'un accident se produise.

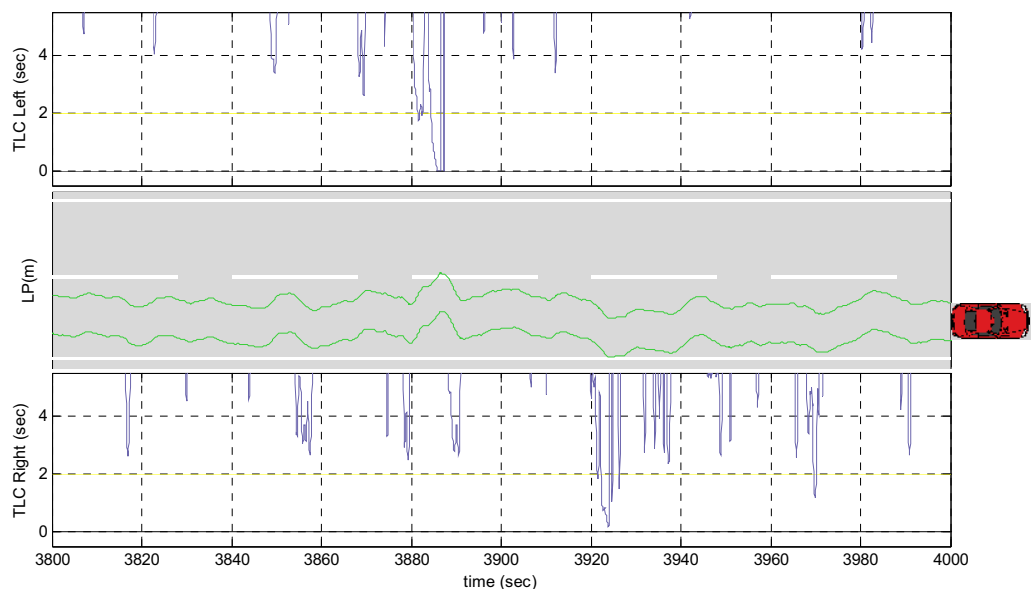


Figure IV.5 : Le temps de sortie de voie, *tlc*. haut) *tlc* gauche, milieu) position latérale, bas) *tlc* droit.

IV.3.2 Bilan sur les systèmes de surveillance du véhicule/conducteur

Le but de cette section est d'identifier les systèmes non intrusifs de détection et prédiction de l'hypovigilance du conducteur qui ont été testés et validés et qui sont actuellement sur le marché automobile ou qui ont un plan de commercialisation. Ces technologies peuvent être classifiées en une ou plusieurs de ces quatre catégories, [45] :

- Technologies « prête à exécuter » (Readiness-to-perform) et « apte au devoir » (fitness-for-duty).
- Technologies de surveillance ambulatoires pour mesurer la vigilance de l'opérateur.
- Technologies basées sur le comportement du véhicule.
- Technologies de surveillance embarquées, en ligne, de l'état de l'opérateur.

Les technologies « prête à exécuter » et « apte au devoir » essaient d'évaluer la capacité de vigilance d'un opérateur **avant** qu'il commence à effectuer son travail. Le but principal est d'établir si l'opérateur

est apte pour exécuter ses tâches sur une durée déterminée, ou au début d'une période supplémentaire de travail. Ce type de système n'est pas relevant de la surveillance en ligne, car il utilise les trois types de métriques (physiologiques, subjectives et de performance) à la fois.

Les technologies de surveillance ambulatoires pour mesurer la vigilance de l'opérateur utilisent des technologies portables pour instrumenter le conducteur et mesurer son activité cérébrale, oculaire, cardiaque, etc. (*métriques physiologiques* principalement). Ce type de système est utile pour aider à la conception et à validation d'un système de surveillance embarqué.

Les technologies basées sur le comportement du véhicule mesurent l'action du conducteur à travers la surveillance des dispositifs mécaniques et électroniques sous contrôle de l'opérateur (*métriques de performance* principalement), comme les variations du volant et de la vitesse ou la déviation de la trajectoire du véhicule, qui sont supposés démontrer des changements identifiables par rapport à son état normal, quand un conducteur est fatigué, [17].

Les technologies de surveillance embarquées, en ligne, de l'état physique de l'opérateur mettent en oeuvre des techniques et des algorithmes pour une surveillance temps-réel de l'opérateur. Elles utilisent des informations physiologiques, liées aux yeux, à la tête, au cœur, à l'activité cérébrale, au temps de réaction, etc.

IV.3.2.1 Quelques systèmes de surveillance de la conduite et de l'état du conducteur

Système de détection de somnolence et d'alerte du conducteur de Toyota

Lors du ITS 2001 à Sydney, Australie, le constructeur Toyota a présenté son système de détection de somnolence du conducteur. Le système utilise une camera montée sur le rétroviseur pour analyser la durée de fermeture des paupières, Figure IV.6. Le système avertit le conducteur à travers l'écran de navigation du véhicule et avec des messages sonores, [18].



Figure IV.6 : Système de détection de somnolence de Toyota.

Système d'avertissement au conducteur inattentif/somnolent de Nissan

Le système repose sur une caméra installée dans le tableau du bord et utilise un algorithme de traitement d'images pour analyser les images du visage et détecter les phases d'inattention ou de somnolence, [18]. De plus, le constructeur offre un système de support d'aide au maintien de la trajectoire mais qui n'est pas utilisé pour la détection de la baisse de vigilance. Le capteur de position latérale associé à ce système utilise deux technologies : vision par ordinateur ou magnétique, cette dernière si la route est équipée de marqueurs magnétiques, Figure IV.7.

Système d'alerte de somnolence du conducteur SafeTRAC

Le système SafeTRAC est basé sur un algorithme breveté conjointement par l'Université Carnegie Mellon et la société AssistWare Technologies Inc., [9], avec le support de NTHSCA-USDOT. Il combine un système de détection préventif de somnolence et un système d'alerte instantané de sortie de voie,

Figure IV.8. Le système mesure la position du véhicule avec un procédé de vision frontal utilisant une caméra CCD installée derrière le rétroviseur.

Selon les informations techniques fournies par le constructeur, le système est capable de déterminer la position du véhicule sur la voie dans plus de 97% du temps, et sous plusieurs conditions climatiques : pluie, neige, soleil, nuit, etc. La seule combinaison posant des problèmes est pluie et nuit. Dans ce cas, un signal sonore averti le conducteur de l'indisponibilité de ce dispositif.

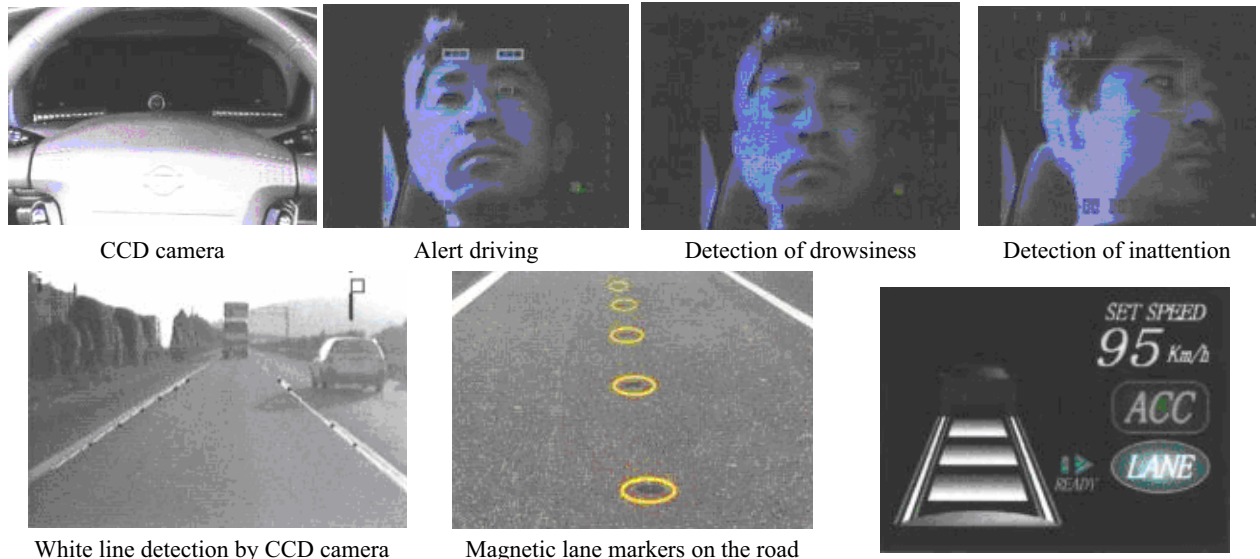


Figure IV.7 : Système de détection de somnolence et d'inattention et système de support de maintien du véhicule sur la voie de Nissan

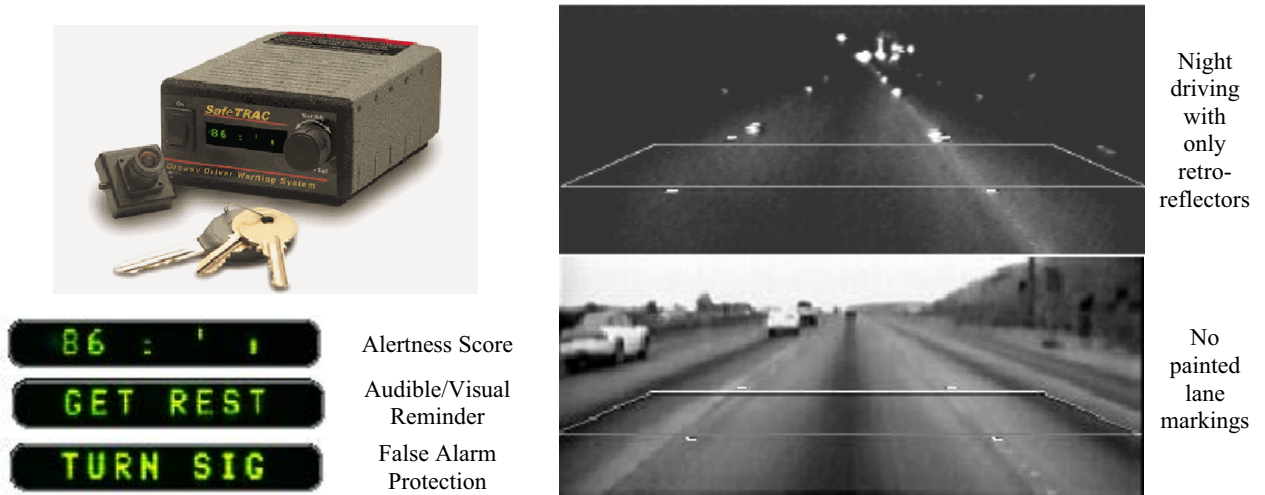


Figure IV.8 : Le système SafeTRACK.

Système de détection de la somnolence de Daimler-Chrysler

Le système de détection de la somnolence du conducteur est basé sur l'évolution de la position du véhicule dans la voie de circulation (information obtenue par analyse des images fournies par une caméra installée derrière le pare-brise), l'angle du volant et un capteur des clignement des yeux, Figure IV.9.

Le système alerte également le conducteur en cas de risque de sortie de voie et d'inattention. Dès que la valeur de TLC est au-dessous d'une certaine valeur, *a priori* identifiée comme une valeur associée à une conduite normale, le système génère des sons de vibreurs de bord de route (rumble strips), [18]. Le

système calcule les paramètres d'un filtre de Kalman pour contrôler les mouvements du volant, utilisant la position latérale, le taux d'inclinaison du véhicule, la courbure de la route, le ratio de variation de courbure de la route, le tilt de l'angle de la caméra.

Le système a été testé durant huit mois, en conditions réelles, sur les autoroutes hollandaises avec le soutien du Ministère de Transports Hollandais. 30 poids lourds de trois constructeurs de camions, dont Daimler-Chrysler, ont été équipés avec ce système. Les camions ont servi au transport de marchandises et n'ont pas été considérés comme véhicules d'essais [49].



Figure IV.9 : Système de détection de somnolence de DC.

Système de suivi du regard du conducteur de Daimler-Chrysler

Le système de suivi du regard du conducteur est basé sur une caméra montée sur le tableau du bord. Un logiciel de traitement d'images détecte les yeux, le nez et la bouche pour ensuite suivre le regard en s'appuyant sur un modèle de la tête humaine, Figure IV.10. Les paramètres principaux du système incluent six degrés de liberté de la tête et leurs vitesses et peut fournir les paramètres biométriques : distance entre les yeux, distance entre les extrêmes de la bouche, hauteur des yeux par rapport au nez et hauteur du nez par rapport à la bouche.

L'utilisation de ce système, conçu par Daimler-Chrysler, est prévue dans le projet AWAKE afin de fournir des informations discriminantes pour améliorer ou compléter, le diagnostic des autres systèmes embarqués de détection d'hypovigilance.

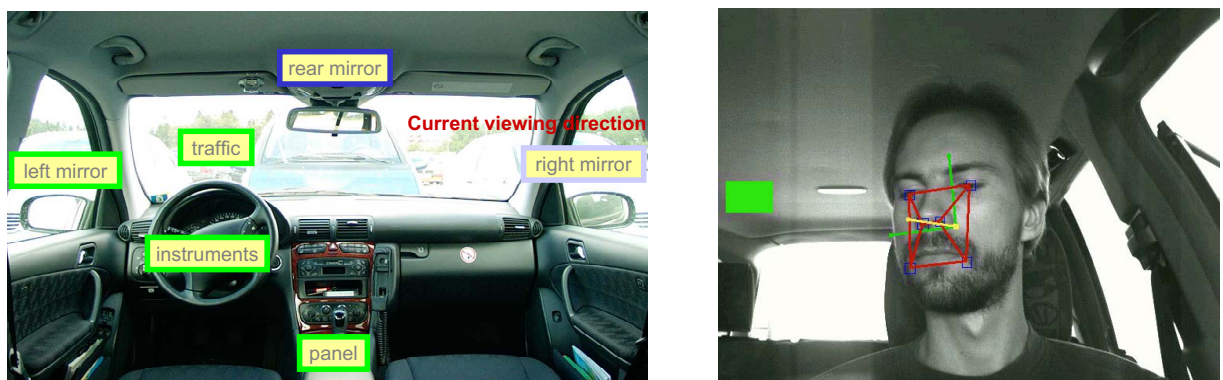


Figure IV.10 : Capteur de suivi du regard de DC.

Système intelligent d'aide au conducteur de Honda

Le système intelligent d'aide au conducteur (Intelligent Driver Support, HiDS) de Honda système est composé de deux sous systèmes, [159] : un système de maintien du véhicule sur la voie (lane-keeping assistance system, LKAS) et un régulateur de vitesse intelligent pour autoroute (intelligent highway cruise control, IHCC).

Le LKAS, basé sur une caméra C-MOS montée sur le pare-brise, calcule et **applique** la correction nécessaire sur le volant pour maintenir le véhicule sur la voie. Le système fonctionne à partir de 65 km/h sur les lignes droites et sur les courbes plus grandes ou égales à 230 m (seulement autoroute). Le IHCC utilise un radar d'onde millimétrique pour détecter les véhicules jusqu'à 100 m. de distance frontale. Selon le niveau de risque de la situation, le système alerte le conducteur par des vibrations sur la ceinture de sécurité ; s'il y a un risque imminent de crash, le système applique la puissance totale de freinage, Figure IV.8. Honda commercialisera ce système en Asie, à partir de l'été 2004.

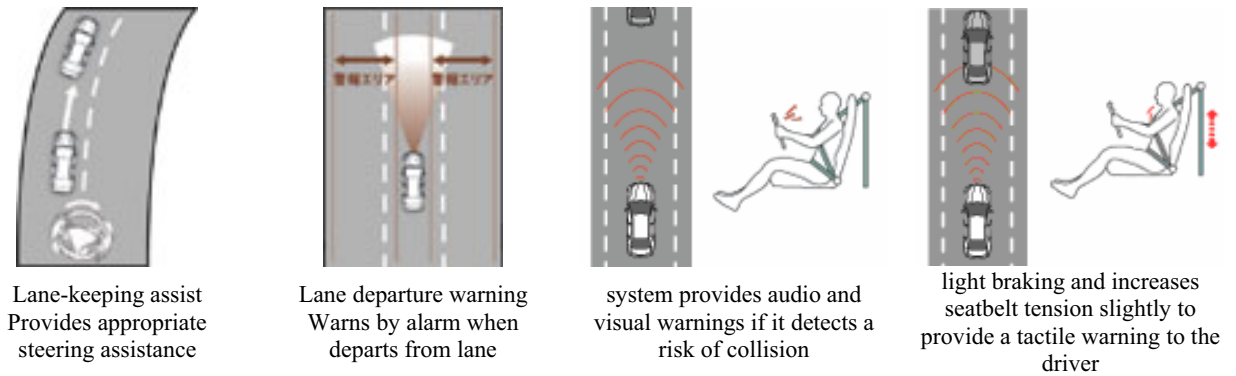


Figure IV.11 : Le système intelligent d'aide au conducteur de Honda.

Système de surveillance PERCLOS

Le système PERCLOS a été développé à l'Université Carnegie Mellon pour détecter la fatigue. Le système utilise deux caméras CCD en configuration stéréo, pour le traitement d'images temps réel. A partir de l'analyse du mouvement des paupières, le système calcule la proportion de temps où les yeux sont fermés durant un intervalle de temps donné (percentage of closure, PERCLOS), [18], [147].

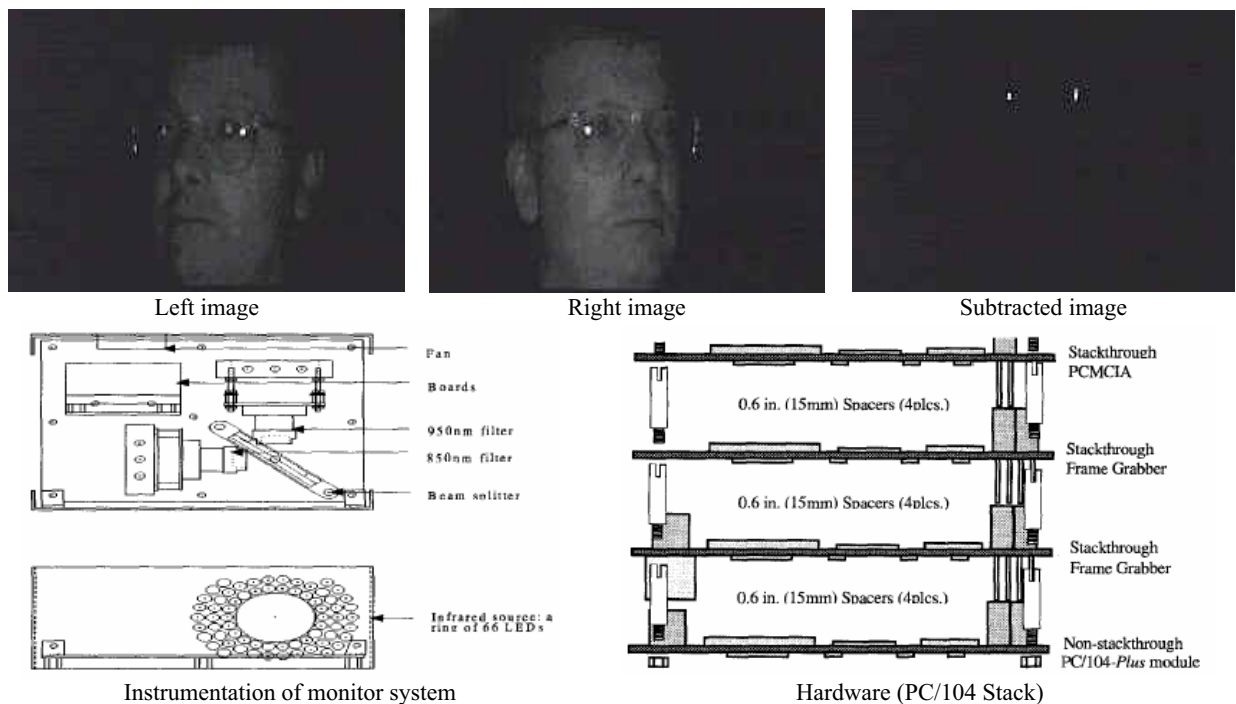


Figure IV.12 : Système de détection PERCLOS.

Le système PERCLOS est techniquement reconnu comme une référence standard pour mesurer la somnolence et valider les systèmes de détection de baisse de vigilance. Les tests conduits par l'Université

de Pennsylvanie, dans un simulateur de conduite, font apparaître une bonne corrélation entre les valeurs fournies par PERCLOS et un test de vigilance psychomotrice (PVT) inter et intra individus.

Système de détection de la somnolence Ford

Ford a lancé une étude sur la somnolence du conducteur automobile d'Octobre 2003 à Février 2004 utilisant son simulateur de conduite VIRTEX, Figure IV.13. Une camera surveille les mouvements de l'œil gauche et un algorithme calcule le pourcentage de fermeture des yeux. Les sorties de voie génèrent des sons de vibreurs mais d'autres possibilités sont envisagées comme un flash rouge dans l'afficheur tête haute ou des vibrations sur le volant, même des actions de corrections comme le système de Toyota. Ford a déjà annoncé la commercialisation du premier avertisseur de sortie de voie (lane departure warning, LDW), [173].



Figure IV.13 : Simulateur de conduite VIRTEX de Ford et le premier système LDW commercialisé massivement en Amérique.

Le système de sécurité intégré ISS© de la société Delphi

L'équipementier automobile Delphi mène diverses recherches sur les systèmes de sécurité active pour l'automobile. Par exemple, durant l'année 2000, un programme de recherche sur la sécurité a été financé par le gouvernement américain et des essais anticollision ont été menés conjointement entre le constructeur automobile General Motors et l'administration américaine chargée de la sécurité sur les autoroutes, afin d'apporter des solutions d'avenir aux consommateurs avec les systèmes anticollision, [134]. Ainsi, Delphi travaille sur différents prototypes intégrant différents équipements dont le véhicule à systèmes de sécurité intégrés ISS© basé sur le développement des synergies entre sécurité active et passive [53].

La recherche sur les systèmes ISS© est réalisée sur un véhicule-laboratoire en situation réelle de conduite, avec la poursuite de plusieurs buts : 1) Déterminer la relation entre les données physiologiques telles que le rythme cardiaque, le rythme respiratoire et les mouvements des yeux, et l'état de somnolence, d'éveil ou d'attention du conducteur. 2) Améliorer la compréhension de la gestion de la vigilance et des répercussions sur le facteur humain, particulièrement pour établir les messages de différents types et degrés d'alerte. 3) Savoir comment parvenir efficacement à rétablir l'attention du conducteur et à le prévenir d'un fort potentiel d'accident, sans interrompre sa vigilance.

Le groupe de recherche des systèmes ISS© divise la conduite en cinq états distincts: 1) Normal. 2) Etat d'alerte. 3) Collision pouvant être évitée. 4) Collision inéluctable. 5) Post-collision.

Pendant la phase normale et celle d'alerte, le véhicule enregistre son environnement, par exemple l'état de la route et les objets autour du véhicule. Le système évalue de manière dynamique l'évolution des conditions de circulation autour du véhicule-hôte. Si une éventualité d'accident est décelée, il déclenche un signal d'alarme suffisamment tôt pour rétablir l'attention du conducteur. Celui-ci sera en mesure de faire une manœuvre d'évitement. Le système choisira la façon d'alerter le conducteur à partir

de la mesure de sa respiration, de la dilatation de son œil ou de son rythme cardiaque. Il saura, ainsi, quel degré d'aide apporter pour gérer la situation en toute sécurité.

Une première génération d'ISS© intégrant plusieurs équipements de sécurité, fonctionnant ensemble pour optimiser la sécurité dans toutes les situations de conduite, avait déjà été exposée au Mondial de l'automobile de Paris 2000. La deuxième génération d'ISS© se rapproche de la réalité de la production de série en intégrant des équipements déjà existants, comme le régulateur de vitesse adaptatif Forewarn™ et le multimédia embarqué Communiport®, ainsi que d'autres qui seront mis en production cette décennie.

De nombreuses interfaces homme-machine ont été intégrées, y compris les systèmes de reconnaissance vocale et de synthèse de la parole. Le véhicule est équipé d'un système capable d'alerter le conducteur par des signaux sonores, sensitifs ou visuels. Les consignes du système de navigation, les problèmes de fonctionnement du véhicule et les alertes anticollision peuvent être projetées en couleurs sur le pare-brise.

ISS© intègre un élément, que l'on trouve de plus en plus dans les systèmes de surveillance de la vigilance, le *système de suivi du regard*. Ce système traite des images pour la détection des yeux et l'analyse des mouvements de paupières, Figure IV.14. Si le système détecte que le conducteur est distrait, ou qu'il devient somnolent, une alarme est déclenchée. En même temps, le système peut interagir avec d'autres systèmes comme celui de régulation de vitesse et celui d'anticollision. Le *système de suivi du regard* opère avec un dispositif d'éclairage de jour et de nuit et est compatible avec l'utilisation de lunettes de soleil [18].

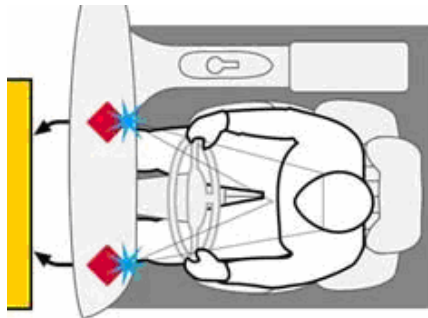


Figure IV.14 : Système de suivi du regard de Delphi.

Ainsi la deuxième génération de véhicules, ISS© développe des équipements qui requièrent de nouveaux capteurs :

1. pour plusieurs mesures biométriques visant à évaluer l'état d'alerte du conducteur : analyse des yeux, enregistrement du rythme respiratoire et un nouveau système de surveillance du rythme cardiaque. Ce dernier mesure les variations de la pression sanguine grâce aux différences de potentiel générées par un courant de faible intensité passant dans les mains du conducteur posées sur le volant. L'évaluation de l'état du conducteur va influencer les systèmes d'alerte en adaptant les signaux d'alerte et la disponibilité des éléments du système multimédia Communiport®.
2. pour la détection précoce avant collision fournie par l'association d'un radar à champ large 76GHz et des capteurs basés sur une reconnaissance optique. Cet ensemble est capable de prévoir l'imminence d'un accident, comme un choc éventuel avec un piéton et de prendre les mesures pour l'éviter.
3. pour une alerte arrière utilisant un radar à champ étroit 17GHz et une caméra braquée sur l'arrière, qui fournit une image sur l'écran multimédia.
4. pour une alerte de changement de voie utilisant un radar 24GHz de détection de proximité qui prévient le conducteur si un obstacle est situé dans l'angle mort du rétroviseur, complétée par une alerte de changement de voie, qui se déclenche quand le conducteur s'aventure hors de sa route.

Les travaux de PSA Peugeot Citroën.

L'approche choisie par PSA Peugeot Citroën est de privilégier le rôle de l'utilisateur. La préoccupation est de diagnostiquer l'activité de conduite dans la perspective d'aider le conducteur lorsque cette activité se dégrade quelle qu'en soit la cause, [101]. Elle s'appuie sur une analyse du comportement dynamique du véhicule à travers des mesures de la vitesse, de l'angle du volant et du positionnement sur la voie de circulation. La détection ne se limite pas à la détection de l'assoupissement mais comprend, aussi, la détection de tout type de baisse de vigilance.

L'évaluation de la performance du système est fait selon une approche éthologique sur la base de l'argumentation suivante : Ce système a pour objectif d'établir un diagnostic sur une tâche réalisée par un conducteur, dans le but de l'aider, sachant que la pertinence du diagnostic sera jugée par le conducteur lui-même. Néanmoins dans ces démarches, l'acceptabilité du système de la part du conducteur reste le facteur le plus important, puisque c'est le conducteur, lui-même, qui conditionne l'utilisation du système selon la perception subjective de son état et de sa conduite et selon sa propre estimation du danger comme ultime critère de validation du système.

Les ingénieurs de PSA font aussi des études ergonomiques sur les interfaces homme-machine d'un système de surveillance de trajectoire latérale qui vise à alerter le conducteur, en cas de dérive latérale dans sa voie, [107]. Les interfaces d'alerte sont de différents types : visuelle, sonore, haptique ou tactile. Ils se sont intéressés à l'évolution du regard du conducteur et à son action sur le volant, lors de l'apparition des alertes. Ils ont essayé de caractériser les actions sur le volant, mais sans réussite sur ce dernier point.

D'après leurs études, il semble que les meilleures interfaces soient les interfaces haptiques (vibrations siège et volant) si on retient les critères de capacité à attirer l'attention, le pouvoir d'alerte, le confort du conducteur et la facilité à distinguer.

Les travaux au LAMIH, Université de Valenciennes

Le LAMIH, à l'Université de Valenciennes, a mené des essais de conduite monotone de longue durée sur simulateur, pour mettre en évidence des modifications temporelles du comportement du conducteur [119]. Une approche multisensorielle utilise des données quantitatives hétérogènes comme celles provenant des capteurs mécaniques : angle du volant, position latérale et vitesse. Sont également exploitées des informations quantitatives telles que : les expressions du visage, les mouvements corporels, aussi que des informations physiologiques : EEG, ECG et EOG.

Pour obtenir une base de données homogène, on utilise des découpages temporels sur des fenêtres fixes de l'ordre de la minute et des découpages spatiaux : tronçon de ligne droite ou de longueur donnée (10km, 20km, etc.) Dans ce traitement des données, après codage, une analyse factorielle des correspondances est réalisée. Le travail se poursuit sur un plan expérimental pour rendre plus réaliste les situations introduites dans le simulateur. Pour l'instant, des méthodes de classification automatique pour diagnostiquer et anticiper les phénomènes de fatigue et de baisse de vigilance n'ont pas été mises en œuvre.

Un autre travail mené par cette équipe, sur la base des essais de conduite sur le simulateur, a été la conception d'un système pour identifier les différents styles de conduite. Ils ont proposé l'utilisation d'une analyse de correspondance multiple (MCA) et d'une analyse discriminante (DA) pour les variables référant à la position du véhicule sur la voie et les actions du conducteur [100].

Les travaux du LESCOT à l'INRETS

Le LESCOT à l'INRETS a développé un dispositif capable d'estimer, en temps réel, le niveau de disponibilité du conducteur dans l'objectif de gérer la diffusion des informations sonores dans l'habitacle,

par exemple l'arrivée d'un appel sur le téléphone portable, dans le cadre du projet européen CEMVOCAS [11].

L'objectif du système est de réduire la charge d'information délivrée au conducteur, en gérant les instants où cette information peut être donnée, c'est à dire l'instant où celle-ci ne gêne plus l'attention du conducteur. Ce système est basé sur une unité de reconnaissance des situations de conduite (DSRU) à partir des capteurs bas coût déjà disponibles dans les véhicules, comme les capteurs de vitesse, d'angle du volant, de la position de pédales du frein et de l'accélérateur. Trois situations sont identifiées : verte quand le conducteur est disponible pour recevoir un message, rouge s'il ne l'est pas et orange en situation intermédiaire.

Le module de reconnaissance de situations de conduite (DSRU) utilise des réseaux de neurones pour corréler les informations fournies par les capteurs avec l'opinion des experts et le conducteur, sur les différentes situations de conduite. Le système identifie correctement 80% des situations et, dans 95% des cas, les conducteurs se sentaient non perturbés par l'arrivée des messages gérés par le système. Dans le cas où il n'y a pas de gestion des messages, les conducteurs déclarent être perturbés entre 43 et 58% des cas.

Le système est personnalisable et prend en compte le temps nécessaire pour assister le conducteur et la façon la plus appropriée de lui faire parvenir le message. Le système propose une technologie adaptative qui accorde la situation courante de conduite et les capacités du conducteur à chaque instant, [19].

Etudes à l'Institut Polytechnique et à l'Université de l'Etat de Virginia

Des études réalisées au Laboratoire de Simulation et Analyse des Véhicules à l'Institut Polytechnique et à l'Université de l'Etat de Virginia, en 1994 et en 2000, [20], [164], sur simulateur avec des chauffeurs de camions, ont porté sur la détection de la baisse de vigilance. Elles concluent qu'il n'y a aucune mesure unique capable de détecter la fatigue au volant. Une combinaison de plusieurs paramètres semblent, dans certains cas, prometteuse.

Des analyses de régression appliquées à l'autoévaluation de la fatigue et à la séparation temporelle entre états d'alerte du conducteur conduisent à de mauvaises performances de prédiction, attribuées d'une part, à la variabilité entre conducteurs et, d'autre part, à plusieurs facteurs externes (autres que la fatigue) qui perturbent la vigilance du conducteur.

IV.3.2.2 L'analyse synthétique

Le problème de la surveillance du conducteur automobile est un problème très complexe et ce tour d'horizon met en évidence la diversité des approches abordées par la communauté scientifique et industrielle pour commencer à avoir des « produits » relativement simples.

Nous avons constaté la présence des systèmes basés sur des *mesures physiologiques*, utilisant, en général, le traitement d'images (donc non intrusifs) pour mesurer le pourcentage de fermetures des paupières, la direction du regard entre autres. D'autre part, les systèmes basés sur des mesures des dynamiques du véhicule permettent de mesurer le comportement du conducteur à travers la surveillance des dispositifs mécaniques et électroniques sous contrôle de l'opérateur (*mesures de performance*). Le Tableau IV.3 montre les caractéristiques des systèmes décrits selon le type des informations d'entrée et leur niveau applicatif.

Loin d'être deux approches concurrentes, les techniques basées sur des mesures physiologiques et celles basées sur des mesures de performance sont deux approches qui sont complémentaires.

TABLEAU IV.3
Caractéristiques de quelques systèmes de surveillance-assistance du conducteur.

SYSTEME	MESURES PHYSIOLOGIQUES	MESURES DE PERFORMANCE	COMBINAISON DES MESURES	NIVEAU D'APPLICATION
Système de détection de somnolence et d'alerte du conducteur de Toyota	✓	-	-	Système informatif
Système d'avertissement au conducteur inattentif / somnolent de Nissan	✓	✓	✓	Système informatif
Système de détection de la somnolence de Daimler-Chrysler	✓	✓	✓	Système informatif
Système de suivi du regard du conducteur de Daimler-Chrysler	✓	-	-	Système informatif
Système de support intelligent au conducteur de Honda	-	✓	-	Système informatif et de commande
Système de surveillance PERCLOS	✓	-	-	Système informatif
Système de détection de la somnolence Ford	✓	✓	En cours de développement	Système informatif et de commande

IV.4 Les projets de recherche

Des travaux ont commencé au LAAS, il y a plus de quinze ans, dans le but de concevoir un système de prévention des accidents potentiels. Le LAAS a participé activement aux programmes Européens DREAM, DETER, PROMETHEUS, SAVE et récemment AWAKE, dans lesquels diverses orientations de mesures multidimensionnelles et de diagnostic par apprentissage ont été proposées.

- Dans le projet DREAM, une approche par réseaux de neurones a été appliquée à la détection des états de danger et de non danger en traitant les non-respects aux règles du code de la route. Cette idée a été reprise dans PROCHIP/PROMETHEUS, [23]. Dans ces travaux, le concept de mesures de performance multiples, comme le mouvement du volant et la position du conducteur sur son siège, a été avancé. Une contribution importante a été l'emploi d'une méthode neuronale, appelée Offset, de construction automatique du réseau de neurones souhaitable pour des systèmes embarqués personnalisés.
- Le projet DETER a continué ces efforts par l'étude approfondie de la corrélation des mesures de performance avec les états du conducteur caractérisés par des mesures psychophysiologiques : une analyse en composantes principales réalise le prétraitement de l'information avant apprentissage des réseaux de neurones. Des expériences préliminaires ont été réalisées sur des scénarios particuliers, tels que l'approche des carrefours et des ronds-points, [79].
- Dans le projet SAVE, [77], la LAAS a participé à la conception d'un système embarqué temps-réel sur le démonstrateur CopiTech, ce qui a mis en évidence l'existence de plusieurs éléments décisifs pour le diagnostic temps-réel de l'hypovigilance. Du côté méthodologique, l'extraction de caractéristiques, l'apprentissage et la décision finale ont été abordés par des méthodes comme l'analyse en composantes principales et indépendantes, le développement de plusieurs techniques neuronales, dont les GRBF, et statistiques pour le diagnostic.

Complémentairement à ces travaux menés pour la détection de l'hypovigilance, le LAAS a également participé :

- A la conception et à la réalisation du simulateur PAVCAS (Poste d'Analyse de la Vigilance en Conduite Automobile Simulée), dans le cadre d'une collaboration interne CNRS entre le LAAS et CeBHA (actuellement CEPA). Ce simulateur est utilisé pour plusieurs études sur les facteurs

humains affectant la conduite automobile, dans le cadre de PREDIT, [29], et du projet Européen AWAKE, [7].

- Au développement de méthodes de diagnostic, en partenariat avec la société ACTIA, notamment pour le diagnostic de pannes automobiles, [109], et la société Siemens-VDO, pour les applications automobile, dont le diagnostic de l'hypovigilance du conducteur, [29].
- Au lancement du projet régional HypoVigil Midi-Pyrénées. Ainsi, le LAAS a su assurer la suite du projet SAVE, sous l'égide de l'IERSET qui a mis en place une plate-forme de recherche composée de plusieurs partenaires. Des méthodes plus adaptées pour l'analyse des signaux y ont été testées, [134].

IV.4.1 Le programme Européen AWAKE (Septembre 2001 – Septembre 2004)

Le projet Européen AWAKE², System for effective Assessment of driver vigilance and Warning According to traffic risk Estimation (IST-2000-28062), est un programme de recherche de Information Society Technologies, IST, de la Commission Européenne. L'objectif du projet AWAKE est d'augmenter la sécurité routière en réduisant le nombre et les conséquences d'accidents de trafic provoqués par l'hypovigilance du conducteur. Afin d'atteindre cet objectif, AWAKE prévoit le développement d'un système fiable et non intrusif, surveillant le conducteur et l'environnement, pour détecter des phases d'hypovigilance en temps réel, en travaillant sur des paramètres multiples.

En cas de détection d'hypovigilance, le système fournit au conducteur un avertissement, réglable à plusieurs niveaux, adapté à l'état estimé d'hypovigilance du conducteur et aux conditions d'environnement et de trafic. Ce système doit fonctionner sûrement et efficacement dans tous les scénarios de conduite.

Le consortium est constitué des partenaires européens suivants :

- Instituts de recherche : CNRS-LAAS, CNRS-CEPA, HIT, VTI, TNO, BIVV/CARA, IAT, ICCS, TUD, COAT.
- Equipementiers automobile : SIEMENS-VDO, ACTIA, AUTOLIV, NAVTECH.
- Constructeurs automobile: DAIMLER-CHRYSLER, FIAT.
- Association de conducteurs : AIT/FIA.



Le système proposé par AWAKE est ciblé vers tous les conducteurs. Néanmoins, il existe certaines catégories de conducteurs pour lesquels un tel système s'avère plus adapté. Dans AWAKE, ces groupes sont nommés groupes d'utilisateurs primaires. Ils ont été sélectionnés à partir de la littérature sur l'étude des accidents. Afin d'identifier les besoins des utilisateurs, une analyse plus précise a été faite à travers une enquête auprès de plus de 500 personnes de 10 pays Européens, [18].

Le consortium AWAKE a organisé également un colloque (AWAKE 1st International Workshop) pour rassembler au niveau international des spécialistes de différents domaines, pour traiter, dans un premier temps, des besoins des utilisateurs et des aspects légaux ; dans un deuxième temps, pour identifier les paramètres les plus sensibles pour la surveillance de la vigilance du conducteur, les méthodes d'évaluation associées et, finalement, le type de stratégie pour interagir avec un conducteur hypovigilant.

IV.4.2 Le projet national Hypovigilance du Conducteur – PREDIT (Mars 2002 – Mars 2004)

Le projet « Facteurs Biologiques et Environnementaux de la Dégradation de la Vigilance chez les Conducteurs et Système de Détection Automatique des Baisses de Vigilance »³ est une action du PREDIT (Programme national de REcherche et D'Innovation dans les Transports terrestres).

² <http://www.awake-eu.org/>

³ Nous rappelons que nous utilisons l'acronyme PREDIT pour faire référence à ce projet en particulier

D'une part, l'analyse des facteurs de dégradation de la vigilance a comme objectif de renforcer et d'approfondir les connaissances relatives aux « facteurs de dégradation de la vigilance et de la sécurité dans les transports », facteurs qui peuvent être des variables isolées ou en interaction avec des phénomènes endogènes (état psycho-physiologique, trait de personnalité....etc.) ou exogènes (ambiance thermique, configuration routière...etc.) aux conducteurs de véhicule. Cette tâche se développe selon trois orientations :

- **une action informative** : orientée vers le grand public, elle a pour but d'avertir sur les quatre causes principales de somnolence (pathologies, prise de médicaments, dette de sommeil, perturbations circadiennes).
- **une action éducative** : proposant quelques pratiques ponctuelles destinées à relever, par une action mentale ou physique, la baisse de vigilance notée à certaines heures de la journée mais, aussi, pour certaines classes de conducteurs : par exemple, la pratique d'une activité physique chez les personnes âgées qui serait susceptible d'améliorer la vigilance diurne et le sommeil nocturne.
- **une action préventive** : par la formation à la prise de conscience et à la prise en charge, par les conducteurs, de leurs signes personnels de baisse de vigilance. Il nous semble important de sensibiliser les conducteurs aux difficultés posées par le vieillissement, par la privation de sommeil, par la prise de médicaments et par la perturbation des rythmes chronobiologiques et chronopsychologiques propres à chacun.

D'autre part, la conception d'un système de diagnostic de la baisse de vigilance d'un conducteur prévoit le développement d'une approche globale portant sur l'analyse :

- des fautes de conduite que le conducteur commet,
- des comportements anormaux par rapport à un comportement normal, identifié et personnalisé,
- des évolutions et des modifications de ses caractéristiques comportementales et de son style de conduite.

Le projet est constitué par les partenaires français suivants :

- Instituts de recherche : CNRS-LAAS, CNRS-CEPA, INRETS, CHU de Toulouse, CHU de Caen, ONERA, IERSET.
- Equipementiers automobile : SIEMENS-VDO, ACTIA.



IV.4.3 Architectures de fonctionnement

IV.4.3.1 L'architecture AWAKE

Le système de diagnostic est formé de plusieurs sous-systèmes. Leur interaction est montrée dans la Figure IV.15. Les trois sous-systèmes sont : le module de diagnostic de l'hypovigilance (hypovigilance diagnosis module, HDM), le module d'estimation de risque dû au trafic (traffic risk estimation, TRE) et le système d'alarme au conducteur (driver warning system, DWS). Ces sous-systèmes sont sous le contrôle du gestionnaire hiérarchique (hierarchical manager, HM). Le HM est le seul sous-système pouvant communiquer avec le DWS et donne des informations, sous requête, aux autres sous-systèmes.

Le HDM réalise un diagnostic l'hypovigilance du conducteur en temps réel, par un traitement multivariable des signaux mécaniques et physiologiques. A partir de cette information, le module calcule trois diagnostics : diagnostic physiologique, diagnostic stochastique et diagnostic déterministe, lesquels sont fusionnés pour envoyer une seule sortie au HM.

Le TRE détecte le niveau de risque, lié à la situation actuelle dû au trafic. Pour cela, le système fait une analyse temps réel de l'évolution des informations provenant des capteurs embarqués (vitesse, taux

d'inclinaison du véhicule – yaw-rate –, angle du volant), des capteurs d'entourage (radar frontal, GPS, détection de la position latérale) et des capteurs environnementaux (température, visibilité, pluie).

Le DWS fournit des signaux de sortie sonores, visuels et tactiles pour avertir le conducteur à travers différentes modalités qui dépendent du niveau de vigilance du conducteur et du niveau de risque lié au trafic. Le HM utilise ces entrées provenant du HDM et du TRE pour déterminer la modalité d'avertissement adéquate.

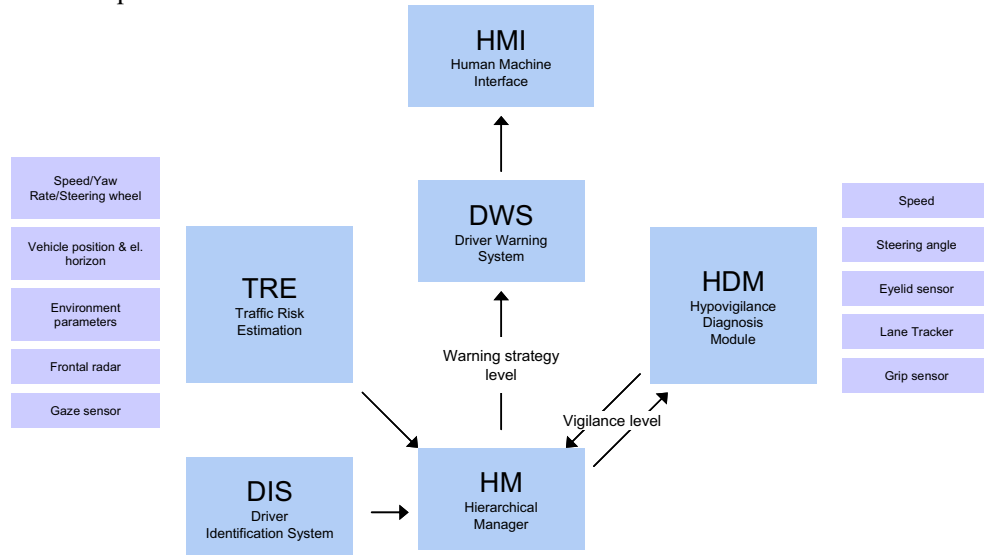


Figure IV.15 : Schéma fonctionnel AWAKE.

IV.4.3.2 L'architecture PREDIT

L'architecture du module de surveillance PREDIT propose d'associer plusieurs approches pour améliorer le diagnostic de l'hypovigilance : Un module de diagnostic statistique, HDM⁴, un module de diagnostic événementiel, EDM, et un module chargé de la fusion des informations des autres modules, FM.

Le HDM analyse, sur un horizon de temps relativement long (de l'ordre de la minute), les évolutions et dégradations des paramètres caractérisant le conducteur et son style de conduite, pour un trajet donné. Le EDM est basé sur l'observation des fautes de conduite par rapport à un référentiel défini a priori. Ces fautes pourront être caractérisées par des comportements dangereux (risque de sortie de voie). Le FM utilise ces deux diagnostics pour fournir une information qualitative au conducteur.

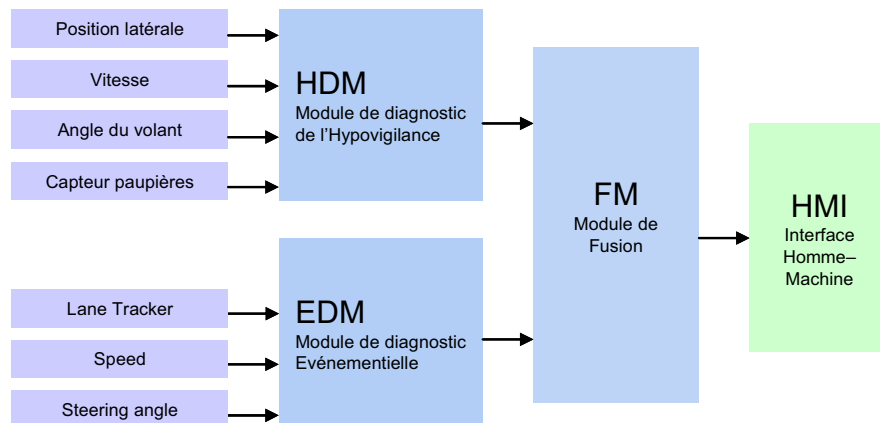


Figure IV.16 : Schéma fonctionnel PREDIT.

⁴ Nous utilisons le même acronyme pour identifier le module de diagnostic de l'hypovigilance.

IV.5 Les essais

Des essais de conduite, sur démonstrateur et sur simulateur, ont été réalisés au cours de plusieurs campagnes d'essais menées sur deux véhicules démonstrateurs et sur deux simulateurs. Les objectifs de ces campagnes sont multiples, [7], [29] :

- Le développement d'un ensemble de bases de données complet et expertisé, pour des conditions environnementales variées et pour différents conducteurs, en vue de valider les principes du système de diagnostic et la pertinence des informations utilisées pour élaborer le diagnostic,
- la validation sur véhicule du système de diagnostic complet, dans des conditions de conduite réelle.

IV.5.1 Moyens expérimentaux

Les moyens expérimentaux du projet PREDIT comprennent :

- Un véhicule d'essais, COPITECH, Figure IV.17, appartenant au LAAS-CNRS et muni de capteurs comportementaux, [77]. Ce véhicule, équipé partiellement lors de précédents projets est utilisé pour la première campagne d'essais. Une mise à jour de l'équipement du véhicule est envisagée pour homogénéiser l'architecture avec nos partenaires.
- Un véhicule d'essais, LAGUNA, Figure IV.18, fourni par la société Siemens-VDO pour les deux campagnes autoroutières suivantes. Il dispose de capteurs de deuxième génération et d'une architecture matérielle et logicielle modulaire et évolutive, [29].
- Le simulateur PAVCAS (Poste d'Analyse de la Vigilance en Conduite Automobile Simulée) installé au CEPA-CNRS à Strasbourg, Figure IV.19. Ce dernier est constitué d'un habitacle avant et des commandes de bord complètes d'une 605 Peugeot, fixée sur trois plateformes superposées mobiles, permettant de restituer l'essentiel des mouvements du véhicule, en réponse aux actions effectuées sur les commandes. L'ensemble du dispositif est installé dans une chambre climatique où les ambiances thermique, hygrométrique, acoustique et lumineuse sont contrôlées.



Figure IV.17 : Véhicule d'expérimentation CopiTech, LAAS-CNRS.



Figure IV.18 : Véhicule d'expérimentation Laguna, Siemens-VDO.



Figure IV.19 : Simulateur de conduite PAVCAS, CEPA-CNRS.

Ces trois supports font partie des moyens expérimentaux du projet AWAKE, qui comprennent en plus deux autres simulateurs :

- Le simulateur VTI, installé à Linköping, en Suède. Il s'agit d'une base qui peut recevoir des véhicules réels et « créer » des sensations réelles de conduite chez le conducteur.
- Le simulateur IAT, de l'Université de Stuttgart, Allemagne. Il est constitué d'une Mégane entière, sans plateforme mobile. L'objectif principal de ce simulateur de conduite est d'étudier la pertinence des dispositifs d'interaction Homme-Véhicule en termes d'ergonomie et d'acceptabilité.

Dans le cadre du projet AWAKE, la phase expérimentale comprend l'intégration du système AWAKE complet sur trois véhicules pour sa validation :

- Une Fiat Stylo, dans la catégorie des véhicules de ville. CRF a été responsable de son équipement.
- Une Mercedes S500, pour les voitures de luxe. Daimler-Chrysler, à Berlin, est le responsable de son équipement.
- Un camion Mercedes Actros, pour les véhicules utilitaires. CARA et ACTIA ont eu en charge l'instrumentation du véhicule.



Simulateur VTI



Simulateur IAT

Figure IV.20 : Simulateurs de conduite automobile également participant dans le projet AWAKE.



Démonstrateur Fiat Stylo



Démonstrateur Mercedes S500



Démonstrateur ACTROS

Figure IV.21 : Les véhicules démonstrateur AWAKE.

IV.5.2 Les campagnes d'essais

IV.5.2.1 Les essais sur démonstrateurs

Les essais sur autoroute ont été divisés en 3 campagnes, [29], dont la première a débuté en septembre 2001, avec CopiTech, la deuxième en novembre 2002 avec la Laguna et la dernière en février 2003 avec la Laguna. Remarquons que les essais sur autoroute représentent plusieurs points d'intérêt mais aussi des inconvénients. L'intérêt principal est, bien sûr, qu'il s'agit d'essais en conduite normale. Par contre, les mesures sont beaucoup plus bruitées (trafic, météo) et il n'est pas possible, pour des raisons de sécurité, de laisser conduire le sujet dans un état fortement hypovigilant ; c'est l'inconvénient majeur de ces essais.

Les objectifs de la campagne N° 1 sont de constituer une base de données expertisée en conditions de conduite prolongée, sur autoroute, par l'enregistrement de paramètres comportementaux véhicule et physiologiques du conducteur, dans le but de tester et valider les différents algorithmes de traitement de données et de diagnostic.

Les objectifs de la campagne N° 2 sont similaires à ceux de la première campagne : il s'agit de constituer une base de données comportementales et physiologique expertisée. Les différences principales résident en l'utilisation d'un nouveau véhicule d'essais : une Renault Laguna équipée par Siemens-VDO et l'utilisation d'une nouvelle échelle d'autoévaluation.

La campagne N°3 finalise les deux campagnes précédentes. Le protocole et les sujets sont sensiblement les mêmes. Elle a pour but de valider le diagnostic en ligne et d'évaluer l'acceptabilité du système par les conducteurs.

Sur ces campagnes, nous avons utilisé les bases des données correspondantes à 11 conducteurs. Ces conducteurs seront notés par la suite C01 à C10 (entre parenthèses les initiales des conducteurs) : C01 (EP), C02 (AE), C03 (CR), C04 (MG), C05 (JP), C06 (TK), C07 (BV), C08 (BG), C09 (SB) et C10 (CC).

On distingue 2 sortes d'essais :

- Les essais instrumentés pour lesquels des électrodes, permettant de mesurer activité cérébrale (EEG) et oculaire (EOG), sont placés sur la tête du conducteur. Ces essais sont soumis à une autorisation préalable du CPPRB⁵. Le trajet est défini dans le protocole et ne peut être modifié.
- Les essais non instrumentés qui ne requièrent aucune autorisation et pour lesquels le trajet est libre.

Les conducteurs C01, C02, C03 ont réalisé 3 essais instrumentés chacun. Les conducteurs C04, C05, C06, C07 ont réalisé 1 essai instrumenté et 1 essai non instrumenté. Les conducteurs C08, C09 et C10 ont participé à des essais non instrumentés.

Le protocole

Les essais se déroulent sur autoroute, avec le véhicule CopiTech et la Laguna. Chaque parcours est composé de deux allers-retours Toulouse–Carcassonne/Lézignan, France.

L'expérimentation commence au CHU-Toulouse-Rangueil (Unité de la Circulation Cérébrale, Service de Neurologie) à 9h/9h30 où le sujet est instrumenté. La technique consiste à coller des capteurs sur le cuir chevelu et autour des yeux afin de mesurer activité cérébrale (EEG) et oculaire (EOG).

Le sujet part ensuite au volant du véhicule. Le conducteur effectue un trajet Toulouse–Carcassonne par l'autoroute le matin, soit 190 km, puis un aller-retour Toulouse–Lézignan par la même autoroute l'après midi, soit 144 km ; soit une durée totale de conduite d'environ 5 à 6 heures. Il est accompagné, au cours de ce voyage, par un médecin qui suit en temps réel l'évolution de sa fatigue, un technicien médical EEG et un superviseur chargé de s'assurer du bon déroulement de l'essai.

L'expérimentation s'achève au CHU-Toulouse-Rangueil.

Les sujets et les consignes

Nous avons fait le choix d'utiliser un petit nombre de sujets mais de répéter les essais afin d'obtenir des temps de conduite plus importants par sujet pour mieux évaluer la répétitivité des caractéristiques de conduite inter sujet. Les conducteurs étaient des volontaires composés d'hommes et de femmes répondant à des critères spécifiques (âge, ayant subi une visite médicale, sans privation de sommeil, etc....) sélectionnés par l'équipe médicale. Pour des raisons de facilité de sélection, les sujets des essais sur autoroute sont tous des étudiants ou de jeunes chercheurs du LAAS-CNRS.

Les sujets avaient pour consigne de conduire sur la voie de droite et de ne dépasser que sur autorisation du responsable de l'essai, sauf, bien sûr, situation d'urgence. Cette consigne a pour effet de

⁵ Comité consultatif de protection des personnes dans la recherche biomédicale (http://ile-de-france.sante.gouv.fr/sante/fen_cc01.htm)

limiter rapidement la vitesse du véhicule à celle d'un poids lourd soit environ 80 km/h. Ainsi, une grande partie du parcours est effectuée dans le sillage d'un camion à une distance suffisante pour que la conduite ne soit pas modifiée par les turbulences générées par celui-ci.

Conduire à une vitesse aussi réduite, dans un contexte autoroutier, a un effet hypovigilant qui est de plus amplifié par une consigne de silence et par le maintien de l'habitacle à une température élevée. Il était, en effet, demandé au conducteur et aux passagers de ne pas parler et la température du véhicule était maintenue à 23°C.

Il s'agissait donc de limiter au maximum le nombre de perturbations externes à la tâche de maintien du véhicule dans la voie et de favoriser l'apparition de signes d'hypovigilance.

L'expertise

La détermination du niveau d'hypovigilance de référence du sujet (dans un référentiel à quatre niveaux) est obtenue par l'analyse des enregistrements EEG à quatre niveaux et EOG (***mesures physiologiques***), enrichie par l'autoévaluation de son état par le conducteur utilisant l'échelle KSS (***mesures subjectives***), mesures définies dans le premier chapitre. Il faut remarquer que l'évolution de cette autoévaluation est beaucoup plus lissée que celle des EEG mais il est apparu, au cours de nos essais, que ces mesures (autoévaluation et EEG) étaient assez mal corrélées.

L'analyse des enregistrements EEG et EOG consiste en un premier temps à découper par page de 16 secondes les signaux. Cette période correspond à la période de visualisation du logiciel, ce qui simplifie et accélère le dépouillement.

On a pu constater, sur les 12 essais réalisés, que les sujets atteignaient très rarement le niveau 3. Ceci semble dû à la nature même de l'essai puisque celui-ci est arrêté lorsque le niveau de vigilance est estimé critique par le neurophysiologue. De plus, nous avons constaté que l'apparition de bouffées d'ondes thêta et de mouvements oculaires lents étaient rares. Ainsi, les indices qui participent principalement à la détermination des niveaux de vigilance sont les durées des bouffées alpha et l'apparition de clignements longs.

IV.5.2.2 Les essais sur simulateurs⁶

La conduite automobile simulée constitue une démarche expérimentale de choix dans les études portant, tant sur la sécurité routière, que dans celles orientées vers la définition des capacités de l'opérateur humain. L'avantage essentiel de la simulation réside dans le fait que le conducteur n'est pas exposé aux dangers multiples de la conduite réelle et que la reproduction du contexte routier est strictement identique pour l'ensemble de la population de sujets, [129]. On sait, toutefois, que la situation de simulation est parfois vécue comme étant artificielle et souvent ludique et, qu'en l'absence de tout risque réel, la disponibilité attentionnelle du conducteur peut fluctuer de façon très importante selon la complexité de la tâche et surtout de son réalisme.

Essais sur le démonstrateur PAVCAS

L'objectif des essais réalisés sur le démonstrateur PAVCAS était de constituer une base de données expertisées, pour étudier l'impact de la privation de sommeil sur plusieurs conducteurs différents, [29].

Essais sur le démonstrateur VTI

Les essais réalisés sur le démonstrateur VTI avaient pour but de constituer une base de données expertisées pour une conduite prolongée sur autoroute des conducteurs de poids lourds.

⁶ Remarque : Par des raisons d'espace, nous ne présentons pas, dans notre travail de thèse, les résultats de ces essais.

IV.6 Développement du module de diagnostic de situations à risque

L'utilisation des systèmes d'assistance au maintien du véhicule sur la voie (LKAS), qui préviennent le conducteur d'une possible sortie de voie, pourrait amener en une réduction de 35% d'accidents, selon le type de route, [40]. En France, les accidents liés à des sorties de route mono-véhicule (single-vehicle road departure, SVRD) représentent 35% des accidents. En 2000, plus de 3000 morts ont été la conséquence de ce type d'accidents, [8].

Une application particulière des LKAS est le système d'avertissement de sortie de route (road departure warning, RDW, systems). Ce système détecte quand le conducteur est sur le point de sortir de la route et déclenche une alarme pour avertir le conducteur suffisamment tôt pour qu'il puisse reprendre le contrôle du véhicule et engager une action corrective.

Les systèmes de RDW les plus simples n'utilisent qu'un seuil définissant la distance latérale minimale, comme les vibreurs des bords latéraux de la route (rumble strips, RBS), [117], sans donner aucune prédiction sur l'état du véhicule. Une autre approche utilise le temps de sortie de voie (*tlc*). En général, le *tlc* a plus d'intérêt que les vibreurs car il avertit quand on considère que le conducteur est en situation de danger potentiel. Cependant, cette prévision peut être erronée parce que son calcul tient compte seulement de la trajectoire du véhicule, négligeant, donc, les actions du conducteur. En conséquence, le nombre de fausses alarmes est en général plus important que celles données par les vibreurs. Batavia, [15], a proposé une combinaison du *tlc* et de l'action des vibreurs du bord de la route, en ajoutant un bord de route virtuel (virtual rumble strip, VRBS), ajusté par un processus d'apprentissage. Dans une étude plus récente, [117], Pilutti a présenté un système de RDW basé sur un VRBS floue, utilisant la distance latérale et la vitesse latérale du véhicule. Dans nos travaux, nous avons proposé une extension de cette dernière approche afin d'introduire les actions du conducteur sur le volant, [71].

IV.6.1 Architecture du module de diagnostic de situations à risque

Le module de détection de situations à risque possède à une architecture classique pour une détection temps-réel, Figure IV.22. Il utilise la mesure et la mise en forme des signaux provenant des capteurs mécaniques pour ensuite :

- réaliser un prétraitement des signaux, filtrant les signaux et en extrayant des informations pertinentes,
- détecter des situations à risque à partir de ces informations et envoyer un avertissement au conducteur en cas de situation de risque potentiel.

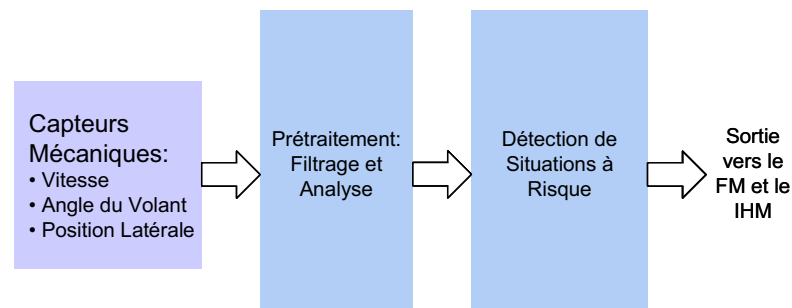


Figure IV.22 : Schéma fonctionnel du module de détection de situations à risque.

Ce module applique, dans la première phase, un filtrage des variables mécaniques, à travers des filtres passe-bas avec une fréquence de coupure de 1.2 Hz pour la position latérale, 2 Hz et 0.5 Hz pour la vitesse. Nous avons fixé ces fréquences après une analyse du spectre en fréquence de chaque signal, sur l'ensemble des bases de données des essais faits sur démonstrateur.

IV.6.2 Module de Diagnostic Événementielle basé sur le temps de sortie de voie

Le temps de sortie de voie peut être calculé hors ligne, [165]. Néanmoins, un problème important avec le tlc est la complexité de son calcul temps-réel sur véhicule. Par ailleurs, cette variable est calculée différemment pour les tronçons droits et pour les courbes. Winsum, [165], a étudié la corrélation entre le tlc réel (sur simulateur) et deux approximations ; une utilisant la première dérivée $\dot{y}$ de la distance latérale y (tlc_1) :

$$tlc_{1Droit} = \begin{cases} \frac{y}{\dot{y}}, & \text{Si } \dot{y} > 0 \text{ (véhicule part vers la droite)} \\ tlc_{\max} & \text{autrement} \end{cases} \quad (IV.1)$$

$$tlc_{1Gauche} = \begin{cases} \frac{W_R - (y + W_V)}{-\dot{y}}, & \text{Si } \dot{y} < 0 \text{ (véhicule part vers la gauche)} \\ tlc_{\max} & \text{autrement} \end{cases} \quad (IV.2)$$

et l'autre utilisant la première $\dot{y}$ et la deuxième dérivée $\ddot{y}$ de la distance latérale y (tlc_2) :

$$tlc_{2Droit} = \begin{cases} \frac{y}{\dot{y} + \ddot{y}}, & \text{Si } \dot{y} + \ddot{y} > 0 \text{ (véhicule part vers la droite)} \\ tlc_{\max} & \text{autrement} \end{cases} \quad (IV.3)$$

$$tlc_{2Gauche} = \begin{cases} \frac{W_R - (y + W_V)}{-(\dot{y} + \ddot{y})}, & \text{Si } \dot{y} + \ddot{y} < 0 \text{ (véhicule part vers la gauche)} \\ tlc_{\max} & \text{autrement} \end{cases} \quad (IV.4)$$

où W_R est la largeur de la route, W_V est la largeur du véhicule et $tlc_{\max}$ est une limite maximum du tlc (dix secondes par exemple). Il a trouvé que le tlc_2 est la meilleure approximation.

L'étude de ces deux approximations, menée au LAAS par Hernández, [77], a permis d'établir l'approximation tlc_1 comme mesure plus adaptée du tlc pour des besoins temps-réel. Ainsi, pour une conduite standard sur autoroute, on peut spécifier de façon adéquate les dynamiques latérales du véhicule en utilisant uniquement la distance latérale y et sa première dérivée $\dot{y}$. Il faut préciser que, dans cette thèse, la distance latérale est définie entre les roues droites et le bord intérieur de la ligne blanche, la plus proche du côté droit du véhicule. D'autres repères peuvent être choisis.

La stratégie d'avertissement courant utilisant le tlc est basée sur la détection de la présence d'une valeur du tlc inférieure à un seuil donné. Afin d'adapter la sensibilité de l'évaluation du tlc à l'angle du volant, nous pouvons ajouter une contrainte de temps, [117]. De cette manière, l'estimation de la valeur du tlc doit être inférieure au seuil d'avertissement (0.1–1.0 sec.) pendant 0.3 secondes pour qu'un avertissement soit émis. Ici, le seuil d'avertissement réduisant au minimum les alarmes fausses est choisi à 0.5 sec.

IV.6.3 EDM basé sur vibreur de bord de route adaptatif

Actuellement, l'approche la plus abordée en prévention de sortie de route mono-véhicule (SVRD) est l'utilisation de vibreurs des bords latéraux de route (RBS) : quand un véhicule franchit les limites de la route, ses roues font contact sur les RBS, ce qui fait vibrer le véhicule, faisant un bruit qui **alerte** le conducteur pour qu'il puisse engager une action corrective.

Cette stratégie d'alerte au conducteur a prouvé son efficacité, ayant un bon niveau d'acceptabilité par les conducteurs, [15]. Néanmoins, les vibreurs requièrent une infrastructure qui n'est pas toujours disponible sur les autoroutes. Par ailleurs, leur emplacement impose la définition d'une distance, typiquement entre 0.15 m et 0.45 m des limites extrêmes des deux côtés de la route. Si cette distance tend vers zéro (vibreurs sur la ligne blanche), elle peut garantir les détections de sortie de route mais il y aura une augmentation du nombre de fausses alarmes. Si cette distance est trop loin de la ligne (s'approchant

du bord de la route), il n'y aura pas de fausses alarmes, mais ceci impose au conducteur de réagir en un temps très restreint.

De façon similaire, les systèmes RBS embarqués sur véhicule activent un signal d'avertissement quand le véhicule dévie d'un certain seuil. La principale différence c'est que le seuil des systèmes RBS embarqués peut être variable, donnant ainsi la possibilité d'améliorer la stratégie d'avertissement.

Essentiellement, les vibreurs virtuels de bord de route (VRBS) évaluent la situation actuelle et décident si le seuil du vibreur doit être élargi (W), réduit (T), ou laissé tel quel (N), selon la loi :

$$VRBS = RBS + VRBS_{adj}. \quad (IV.5)$$

$VRBS_{adj}$ est finement modifié pour produire le VRBS et une alarme sera déclenchée si la distance latérale excède VRBS.

Comme discuté dans la section consacrée au *tlc*, la distance latérale y et sa première dérivée $\dot{y}$ représentent de manière suffisante les caractéristiques dynamiques latérales du véhicule pour une conduite sur autoroute. Utilisant cette information, Pilutti, [117], a proposé un système à base de règles floues comme mécanisme d'ajustement du seuil du vibreur. De cette manière, un avertissement peut être émis si le véhicule dérive lentement et se place trop près du bord de la route, mais aussi si le véhicule s'approche rapidement du bord de la route, même s'il en est encore loin. Le Tableau IV.4 présente l'ensemble des règles de cette implémentation, que nous baptisons $VRBS_1$.

TABLEAU IV.4
REGLES POUR DETERMINER $VRBS_{adj}$ UTILISANT y ET $\dot{y}$ COMME ENTREES.

		$\dot{y}^2$		
		<i>petite</i>	<i>moyenne</i>	<i>grande</i>
y	<i>proche</i>	W	N	T
	<i>loin</i>	N	N	T

Néanmoins, il existe plusieurs facteurs supplémentaires à prendre en compte, comme le niveau de vigilance, les conditions climatiques, le risque de sortie de route (en montagne par exemple), etc. Dans nos travaux, nous introduisons, de manière simple, le niveau de vigilance du conducteur. Pour ce faire, nous analysons les actions du conducteur sur le volant comme réaction de celui-ci pour corriger la position du véhicule sur la route. Pour cela, nous faisons appel au concept de **temps de réaction**, qui est le temps entre la présentation d'un stimulus et l'exécution d'une réponse (mesure de **performance tâche primaire**). Bien sûr, l'effet des actions sur le volant sur la distance latérale du véhicule n'est pas facile à évaluer. Dans notre approche, nous ajoutons une troisième variable, le **temps de correction**, t_c , qui correspond au retard entre la correction du conducteur sur le volant et la correction de la dynamique latérale du véhicule, pour une trajectoire approchant le côté droit ou gauche de la route. Le Tableau IV.5 présente la base des règles augmentée en tenant compte du temps de correction. Nous appelons cette implémentation $VRBS_2$.

TABLEAU IV.5
REGLES POUR DETERMINER $VRBS_{adj}$ UTILISANT y , $\dot{y}$ ET t_c COMME ENTREES.

			$\dot{y}^2$		
			<i>petite</i>	<i>moyenne</i>	<i>grande</i>
t_c	<i>court</i>	<i>proche</i>	W	N	T
		<i>loin</i>	N	N	T
	<i>long</i>	<i>proche</i>	N	T	T
		<i>loin</i>	N	N	T

Le mécanisme d'inférence, discuté dans le chapitre III, utilise une base de règles si-alors de la forme :

si (y est loin) et ($\dot{y}^2$ est petit) **alors** (VRBS_{adj} est nul). (IV.6)

Ces règles représentent des heuristiques sur l'utilisation du VRBS. La base de règles floue est complète et couvre tout l'espace d'entrée.

Chaque variable d'entrée est fuzzifiée avant d'être traitée dans la base de règles floues. Ainsi, la distance latérale y est couverte par deux fonctions d'appartenance, représentant les ensembles flous *proche* et *loin*, Figure IV.23 a). La première dérivée de la distance latérale $\dot{y}$ est partitionnée en trois fonctions d'appartenance : petite, moyenne et grande, Figure IV.23 b). La troisième variable, le temps de correction, t_c , du VRBS_2 , a deux fonctions d'appartenance : *court* et *long*, Figure IV.23 c). Tous les ensembles flous d'entrée ont des fonctions d'appartenance gaussiennes. L'espace d'entrée U est défini pour travailler dans l'intervalle défini par la partition floue de chaque variable d'entrée. Toute valeur sortant de ces intervalles est saturée à la limite minimum ou à la limite maximum, selon la situation.

La sortie du système d'inférence floue, VRBS_{adj} , est couverte par trois fonctions d'appartenance de type trapézoïdales, représentant les ensembles flous *réduire*, *ne rien faire*, et *augmenter*, Figure IV.24 a). Le niveau de sortie restreint le VRBS à ajuster le vibreur avec une tolérance de $[-0.5, 0.5]$ m.

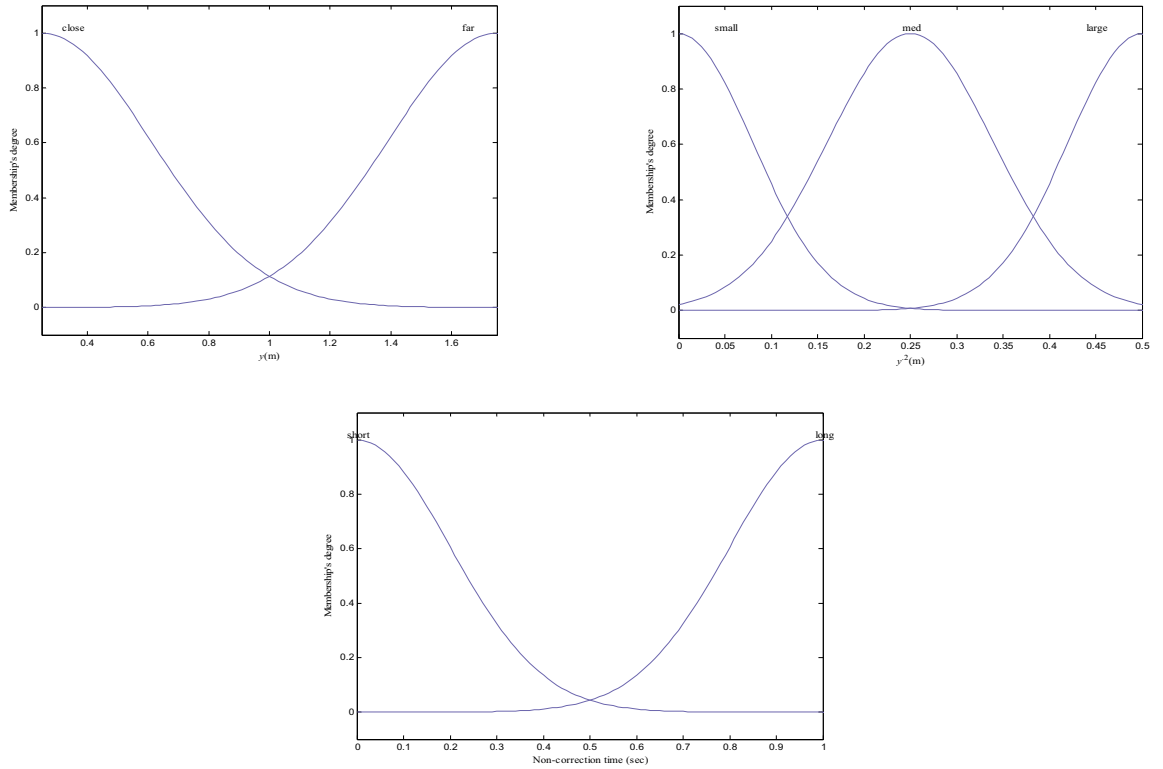


Figure IV.23 : a) Fonctions d'appartenance d'entrée pour la distance latérale et y (μ_{close} et μ_{far}). b) Fonctions d'appartenance d'entrée pour la première dérivée de la distance latérale $\dot{y}$ (μ_{small} , μ_{med} et μ_{large}). c) Fonctions d'appartenance d'entrée pour le temps de correction t_{nc} (μ_{short} et μ_{long}).

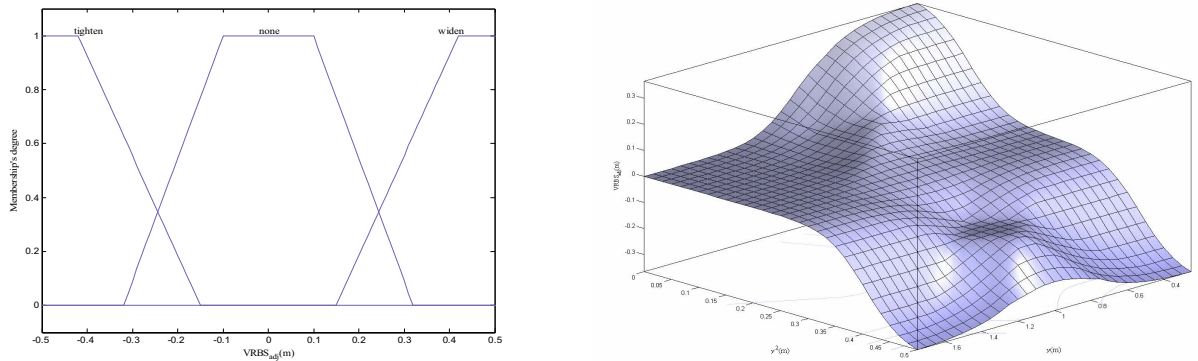


Figure IV.24 : a) Fonctions d'appartenance de sortie pour VRBS_{adj} (μ_{tighten} , μ_{none} et μ_{widen}). b) Exemple de réglage du VRBS utilisant y et y^2 seulement.

Le mécanisme d'inférence de type Mamdani, présenté dans le chapitre II, est utilisé sur cette base de règles floues, le produit algébrique étant retenu comme t-norme, pour l'intersection d'ensembles flous. L'implication des règles utilise l'opérateur minimum et le résultat de chaque implication est agrégé par un opérateur max. La défuzzification est de type centroïde. La Figure IV.24 b) illustre la sortie du système VRBS par rapport à tout l'espace d'entrée formé par y et y^2 .

IV.6.4 Comparaisons des différents approches EDM

L'évaluation de la validité de la stratégie d'avertissement est un sujet très important, pour déterminer lequel des événements de franchissement des marquages externes est suffisamment critique pour être jugé comme indication de sortie de route (RDE) pour avertir en conséquence.

Méthode d'évaluation

Dans [116], les auteurs proposent une approche pour évaluer le RDW qui consiste à prendre en compte le changement du comportement du conducteur avant et après l'avertissement, à travers des dynamiques de l'angle du volant. Néanmoins, leur conclusion est que cette métrique est très sensible aux courbes et n'indique pas si l'avertissement a été généré en premier. Une deuxième idée est de proposer les distances des vibreurs de bord de route comme métrique pour comparer la performance des différentes stratégies d'avertissement. L'implantation des vibreurs sur les bords de route se fait à des distances entre 0.15 m. et 0.45 m. par rapport aux marquages externes de la route. Nous décrivons cette métrique nous adoptons pour évaluer les trois approches pour la construction du EDM.

L'évaluation basée sur les distances entre les vibreurs et les marquages externes est utilisée pour former un ensemble d'avertissements de *validation* v . De même, chaque stratégie d'implémentation du RDW va former un ensemble d'avertissements d'évaluation e . Ainsi, les avertissements dans e qui coïncident avec ceux de v sont classifiés comme des *bonnes détections* H ; les *fausses alarmes* F correspondent à des avertissements dans e qui ne sont pas un élément de l'ensemble de validation H et, finalement, les avertissements dans v qui ne sont pas éléments dans e sont classifiés comme *non détections* M .

Nous utilisons deux ensembles, représentés par les vibreurs placés à l'intérieur des marquages externes de la route de 0.15 m. et de 0.3 m.

Résultats

Pour cette analyse des trois approches pour la construction du EDM, nous avons utilisé les bases de données correspondant aux essais. Le Tableau IV.6 présente le nombre d'événements de franchissement

des marquages externes gauche et droit pour l'ensemble des conducteurs, par rapport aux pneus gauches-droits, excédant 0, 0.15, 0.3 et 0.45 m. Puisqu'il n'y a aucun événement de franchissement du côté gauche, nous traitons seulement les événements de franchissement du côté droit.

TABLEAU IV.6

Nombre d'événements de franchissement des marquages externes à gauche et à droite par les conducteurs C1-C8.

<i>Conducteur</i>	<i>Côté Droit</i>				<i>Côté Gauche</i>			
	0 m	0.25 m	0.3 m	0.45 m	0 m	0.25 m	0.3 m	0.45 m
C01	1	1	0	0	0	0	0	0
C02	19	0	0	0	0	0	0	0
C03	21	4	0	0	0	0	0	0
C04	2	0	0	0	0	0	0	0
C05	0	0	0	0	0	0	0	0
C06	3	0	0	0	0	0	0	0
C07	0	0	0	0	0	0	0	0
C08	6	1	0	0	0	0	0	0
C09	0	0	0	0	0	0	0	0
C10	3	0	0	0	0	0	0	0

La première comparaison se fait sur le EDM basé sur *tlc*. Le Tableau IV.7 présente la comparaison entre les ensembles d'avertissements $v_{0.15}$ et $v_{0.3}$ et l'ensemble d'avertissements $e_{0.15}$ et $e_{0.3}$, correspondant au *tlc*, le Tableau IV.8 présente la comparaison entre les ensembles $v_{0.15}$ et $v_{0.3}$ et les ensembles $e_{0.15}$ et $e_{0.3}$, de l'approche VRBS₁ et le Tableau IV.9 présente la comparaison entre les ensembles $v_{0.15}$ et $v_{0.3}$ et les ensembles $e_{0.15}$ et $e_{0.3}$, de l'approche VRBS₂.

TABLEAU IV.7

Evaluation du EDM basé sur le *tlc* utilisant les ensembles de validation $v_{0.15}$ et $v_{0.3}$.

<i>Conducteur</i>	<i>Distance à 0.15 m</i>				<i>Distance à 0.3 m</i>			
	$n(v_{0.15})$	$n(H)$	$n(M)$	$n(F)$	$n(v_{0.3})$	$n(H)$	$n(M)$	$n(F)$
C01	1	1	0	0	0	0	0	0
C02	0	0	0	0	0	0	0	0
C03	4	4	0	2	0	0	0	0
C04	0	0	0	0	0	0	0	0
C05	0	0	0	0	0	0	0	0
C06	0	0	0	0	0	0	0	0
C07	0	0	0	0	0	0	0	0
C08	1	1	0	3	0	0	0	4
C09	0	0	0	0	0	0	0	0
C10	0	0	0	0	0	0	0	0

TABLEAU IV.8

Evaluation du EDM basé sur le VRBS₁ utilisant les ensembles de validation $v_{0.15}$ et $v_{0.3}$.

<i>Conducteur</i>	<i>Distance à 0.15 m</i>				<i>Distance à 0.3 m</i>			
	$n(v_{0.15})$	$n(H)$	$n(M)$	$n(F)$	$n(v_{0.3})$	$n(H)$	$n(M)$	$n(F)$
C01	1	0	0	0	0	0	0	0
C02	0	0	0	0	0	0	0	0
C03	4	0	4	0	0	0	0	0
C04	0	0	0	0	0	0	0	0
C05	0	0	0	0	0	0	0	0

C06	0	0	0	0	0	0	0	0
C07	0	0	0	0	0	0	0	0
C08	1	1	0	6	0	0	0	4
C09	0	0	0	0	0	0	0	0
C10	0	0	0	1	0	0	0	0

TABLEAU IV.9

Evaluation du EDM basé sur le VRBS₂ utilisant les ensembles de validation $v_{0.15}$ et $v_{0.3}$.

<i>Conducteur</i>	<i>Distance à 0.15 m</i>				<i>Distance à 0.3 m</i>			
	$n(v_{0.15})$	$n(H)$	$n(M)$	$n(F)$	$n(v_{0.3})$	$n(H)$	$n(M)$	$n(F)$
C01	1	1	0	0	0	0	0	0
C02	0	0	0	0	0	0	0	0
C03	4	1	3	3	0	0	0	1
C04	0	0	0	0	0	0	0	0
C05	0	0	0	0	0	0	0	0
C06	0	0	0	0	0	0	0	0
C07	0	0	0	0	0	0	0	0
C08	1	1	0	6	0	0	0	4
C09	0	0	0	0	0	0	0	0
C10	0	0	0	1	0	0	0	0

Nous montrons sur un exemple, extrait de l'un des essais du conducteur C3, les capacités d'adaptation de ces implémentations dans la Figure IV.25. Dans cet exemple, la distance du vibreur est 0.15 m. Une alarme potentielle existe à la seconde 7179.3, que le *tlc* détecte également avec une prédiction de 0.01 secondes. La différence entre les implémentations du VRBS₁ et du VRBS₂ peut être appréciée. Néanmoins, il existe un avertissement à la seconde du VRBS₂ sans qu'aucun autre avertissement soit déclenché. Cette détection est liée à un manque de correction de durée plus importante. La Figure IV.26 détaille l'intervalle autour de cet instant.

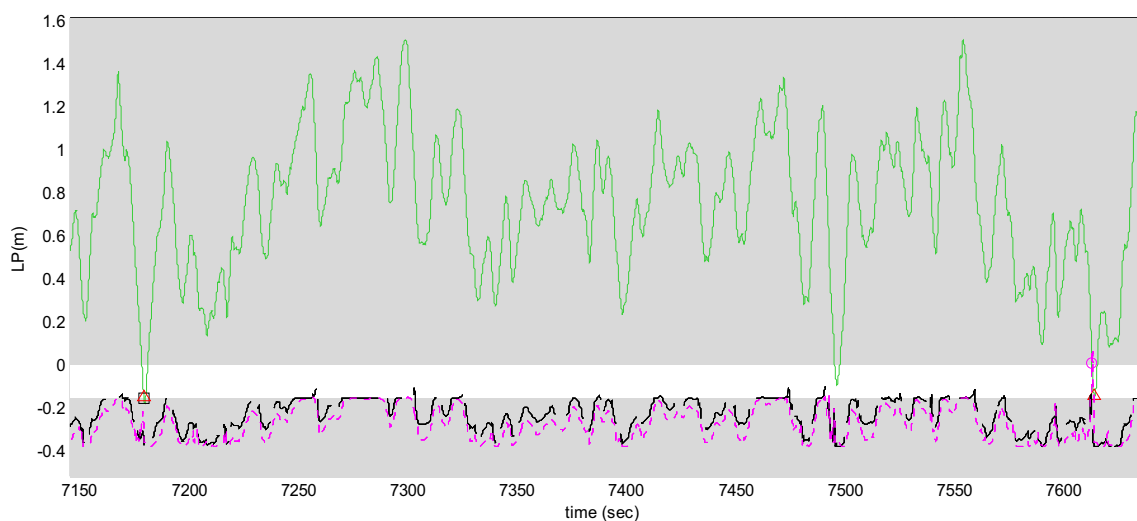


Figure IV.25 : Exemple de la trajectoire du véhicule, montrant l'effet du VRBS₁ (ligne en tiré long) et du VRBS₂ (ligne en tiré court) pour le conducteur C3. Le carré représente un avertissement du vibreur placé à 0.15 m. et le triangle un avertissement du *tlc*.

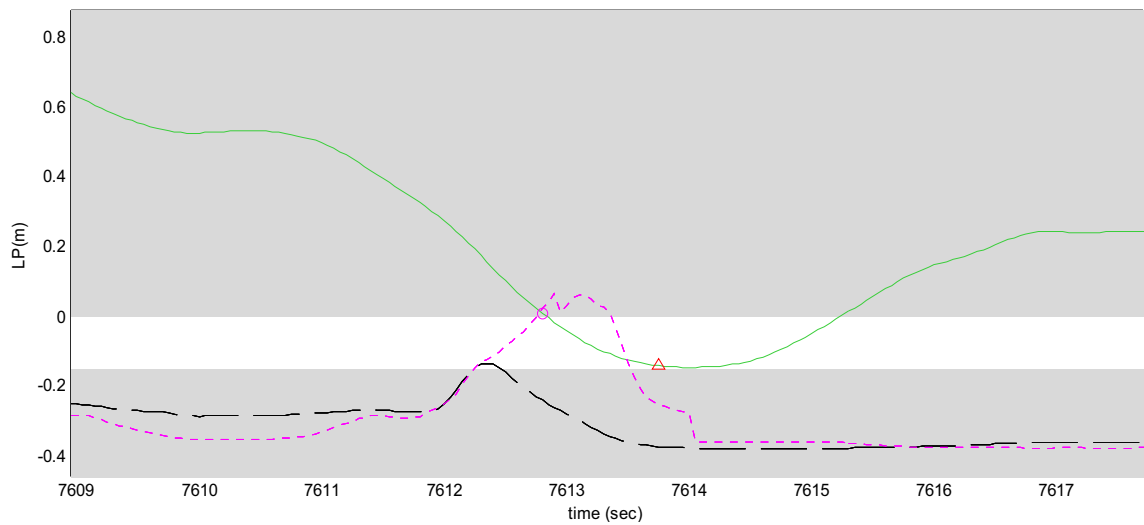


Figure IV.26 : Le VRBS₂ (ligne en tiré court), comparé à l'évolution du VRBS₁ (ligne en tiré long), génère un avertissement (cercle) pour le conducteur C3. Il n'y a pas d'avertissement du vibreur à 0.15 m. mais une détection du *tlc* (triangle).

IV.7 Détection de l'hypovigilance du conducteur automobile

L'un des axes de recherche les plus importantes ces dernières années, dans le domaine des transports, a été la détection de l'hypovigilance du conducteur automobile, car celle-ci est la première cause directe ou indirecte d'accidents routiers. Un conducteur, dans un état hypovigilant, ne tentera pas d'action correctrice avant une collision.

Nombreux sont les travaux menés concernant le développement d'un système embarqué de détection de la dégradation de la vigilance du conducteur, à partir de l'observation de son activité de conduite seule, [17], [18], [44], [55],... La surveillance du système *conducteur-véhicule-environnement* à partir des dynamiques de la voiture est très complexe, [77], car elle porte également sur l'influence de phénomènes extérieurs tels que virages, vent latéral, chaussée irrégulière, etc. Une *approche multisensorielle* s'impose donc, permettant d'enrichir l'information disponible et atténuer le « bruit » porté par chacun des signaux. Toutefois, il convient de restreindre le choix des capteurs aux seuls réellement porteurs d'une information pertinente.

De plus, il faut tenir compte des *variabilités inter-individuelles et intra-individuelles*, car les habitudes de conduite changent d'un conducteur à l'autre et varient, pour même conducteur, sous l'influence de facteurs externes (conditions climatiques, de trafic, etc.) et personnels (motivation pour le parcours, condition physique, etc.). Dans de telles conditions, il n'est pas possible de définir une conduite type relative aux différents états des conducteurs (somnolence, fatigue, effets de l'alcool, etc., comme abordé dans le projet Européen DETER à travers de la définition de plusieurs classes) ou, tout simplement, d'un état vigilant et d'un état hypovigilant (hypothèse à deux classes du projet Européen SAVE), [17], [77].

La *validation* d'un tel type de système se heurte à d'importantes difficultés. L'apparition de signes physiologiques caractéristiques de l'hypovigilance n'est pas nécessairement synchrone avec l'apparition de la dégradation de la performance de la conduite. En effet, un conducteur fatigué, mais pas somnolent, peut tout à fait contrôler son activité de conduite avant d'atteindre un état extrême, [44], [134]. Ceci explique les faibles taux de corrélation entre ces mesures.

Pour l'ensemble de ces approches, les principales considérations retenues sont :

- Une fréquence de mise à jour du diagnostic de l'ordre de la minute, de quelques minutes ou la conséquence d'un événement particulier. En effet, une mise à jour plus fréquente serait une distraction inutile, voire dangereuse, pour le conducteur.
- Un diagnostic comportemental de l'état de vigilance du conducteur, par rapport à l'évolution de la conduite, vis-à-vis d'*une classe* de référence apprise, dite *normale* ou *vigilante*.

IV.7.1 Architecture du module de détection de l'Hypovigilance

Le module de détection de l'état de vigilance du conducteur (HDM) est basé sur la fusion d'informations provenant des capteurs « physiologiques », décrivant le comportement du conducteur, et de capteurs « mécaniques » reflétant les mouvements du véhicule. Les signaux correspondants sont analysés séparément dans le HDM qui se décompose en deux sous-modules, Figure IV.27 :

- « HDM performance », analysant les informations provenant des capteurs mécaniques, pour surveiller l'évolution de la conduite,
- « HDM physiologique », obtenant des mesures directes de l'état de vigilance du conducteur.

La sortie de ces modules est ensuite fusionnée pour aboutir à une décision finale.

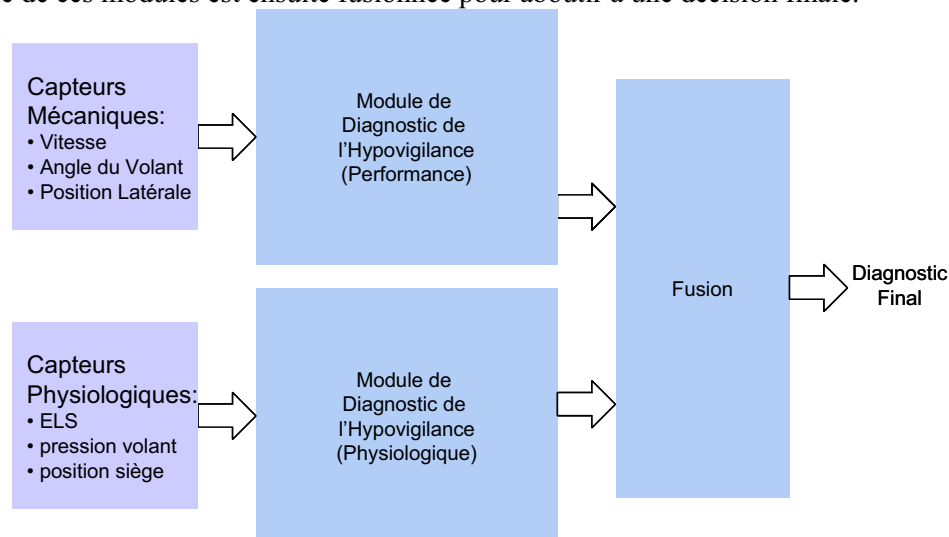


Figure IV.27 : Schéma fonctionnel du module de détection de l'Hypovigilance.

IV.7.1.1 Architecture du HDM performance

Le besoin de personnalisation du système pour chaque conducteur amène à définir un protocole spécial pour construire un modèle approprié. Pour cela, nous avons sélectionné les 20 premiers minutes de voyage comme période de référence pour paramétrer un modèle de conduite vigilante (une seule classe), car le niveau de vigilance peut diminuer significativement après cette phase, [129]. Pour des besoins temps-réel, cette contrainte semble acceptable, même si elle doit être mieux définie. Après cette phase « d'adaptation » au conducteur, le HDM analyse le comportement durant toute la durée du trajet.

Dans l'architecture AWAKE, le HDM demande le ID au module d'identification du conducteur (DIM) et vérifie, dans la base de données, s'il existe un modèle associé au conducteur. Si c'est le cas, le HDM utilise le modèle. Si non, le système aura comme consigne la construction d'un modèle, en espérant avoir les meilleures conditions pour que le conducteur soit vigilant (début du trajet, heure du jour, etc.). Le HDM physiologique pourra assurer le fonctionnement du HDM pendant l'acquisition de la période de référence, [62].

La qualité du modèle construit dépend de la pertinence des informations d'entrée. Pour identifier ces informations, nous avons choisi de limiter l'analyse de l'évolution de la conduite au maintien du véhicule sur la voie. De cette manière, les paramètres mécaniques associés sont : la position latérale, l'angle du

volant et la vitesse. Le choix des caractéristiques à extraire des signaux est un pas fondamental mais pas évident. Nous en avons déjà énoncé quelques unes dans la section VI.3.1.2. que nous pouvons compléter avec les mesures statistiques et fréquentielles décrites dans le chapitre II. Pour calculer ces caractéristiques, nous pouvons utiliser des fenêtres d'analyse de taille fixe ou variable. L'objectif de ces dernières est d'analyser le signal dans des intervalles de temps où le comportement observé est stable, afin d'éviter, au maximum, le mélange des caractéristiques représentant une conduite normale et une conduite anormale.

Ainsi, le travail du HDM performance peut être décomposé en différentes tâches, [30] :

- mesure et mise en forme des signaux provenant des différents capteurs,
- prétraitement pour l'extraction des caractéristiques
- diagnostic (le cœur du système) et estimation en ligne de l'évolution de l'état du conducteur,
- décision finale qui fournit une information au conducteur au travers d'une interface homme-machine adaptée.

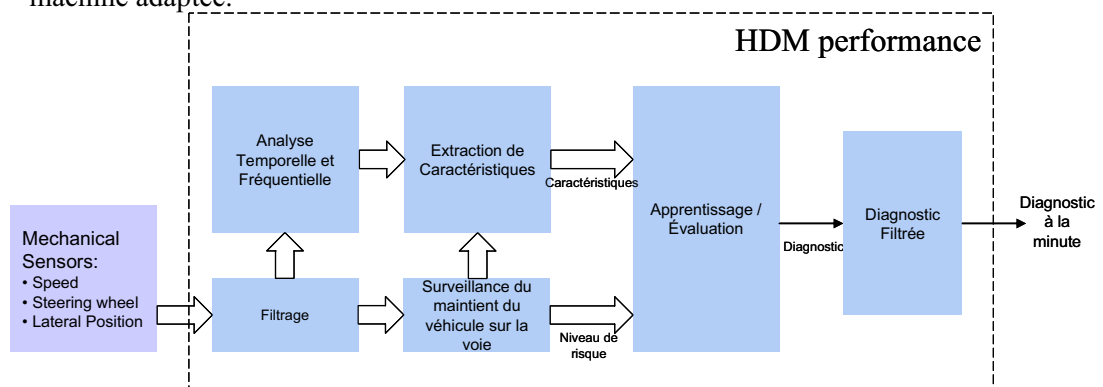


Figure IV.28 : Schéma fonctionnel du module HDM performance.

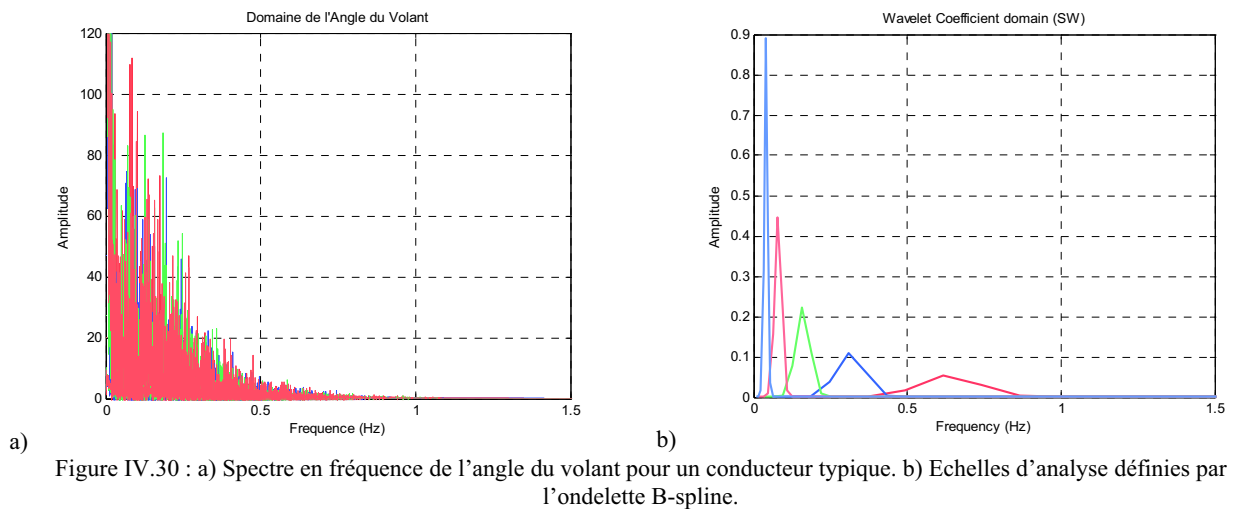
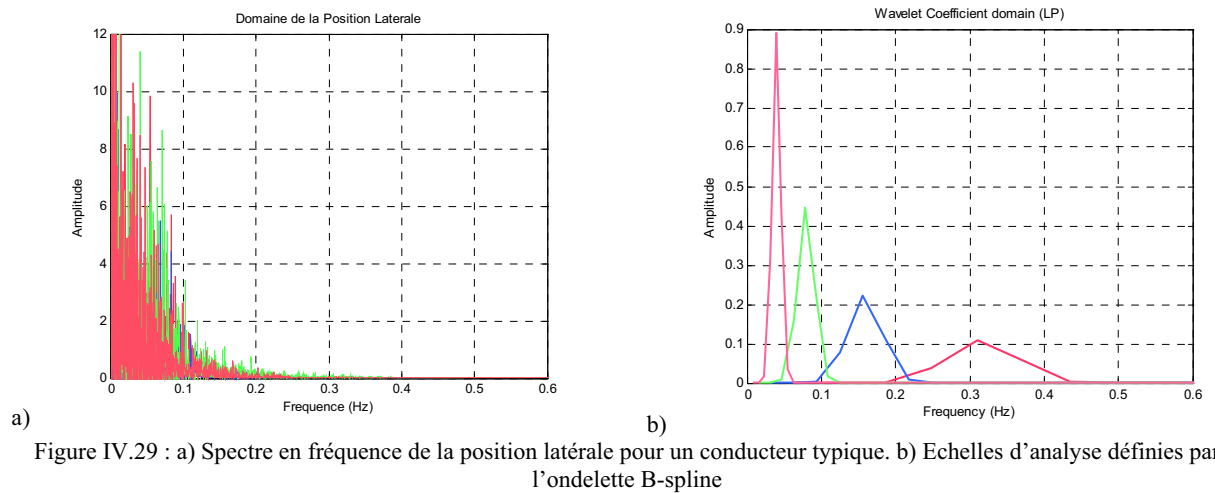
IV.7.1.2 Analyse temporelle et fréquentielle des signaux

Dans [134], Santana a introduit l'analyse en ondelettes pour faire une analyse en temps et en fréquence de la position latérale et de l'angle du volant. Cette analyse permet de détecter les instants où il existe un changement « brusque » du comportement du signal, ce qui induit la définition des fenêtres d'analyse de taille variable où le signal présente une dynamique stable.

Définition des échelles d'analyse

Pour cette implémentation hors-ligne, Santana a introduit l'ondelette dyadique B-spline pour la décomposition du signal, sur huit fréquences (échelles) pour l'angle du volant et huit autres pour la position latérale. Les fréquences choisies dans le cas de la position latérale se situent entre 0.04 Hz et 0.625 Hz. Pour l'angle du volant, les échelles sont définies entre 0.04 Hz et 1.25 Hz.

L'implantation d'un système embarqué pose des problèmes de calcul des huit échelles de coefficients par ondelettes. En vérifiant le comportement du spectre en fréquence de la position latérale et de l'angle du volant, nous avons réduit le nombre d'échelles d'analyse à quatre pour la position latérale, chaque intervalle ayant une fréquence centrale située en 0.04 Hz, 0.08 Hz, 0.16 Hz et 0.32 Hz respectivement, Figure IV.29, et cinq pour l'angle du volant, avec les fréquences centrales 0.04 Hz, 0.08 Hz, 0.16 Hz, 0.32 Hz et 0.64 Hz, Figure IV.30.



Variables statistiques et fréquentielles sur des fenêtres de taille variable

Dans notre approche, [133], [134], l'information est décomposée en séquences où le signal a un comportement homogène. Puis, dans chaque intervalle, nous calculons des caractéristiques statistiques et fréquentielles, souvent appelées variables artificielles.

Pour déterminer les instants où il existe une **rupture** (ou singularité) en temps et en fréquence, nous utilisons une analyse en ondelettes, comme défini dans le chapitre II. La mise en ligne de l'algorithme de détection des ruptures réalise une transformée en ondelettes à la minute, sur les deux minutes de conduite précédant l'instant d'analyse. Ceci permet le chevauchement de l'analyse en ondelettes d'une minute à l'autre, sans que les discontinuités affectent la détection de ruptures. La Figure IV.31 (d'après [134]) montre un exemple de localisation des changements sur la dynamique de la position latérale et l'angle du volant. Pour détecter les ruptures, il faut calculer le produit (graphe 3 et 6) des coefficients en ondelettes (graphe 2 et 5). Le changement du comportement du signal correspond à un incrément important du produit des coefficients.

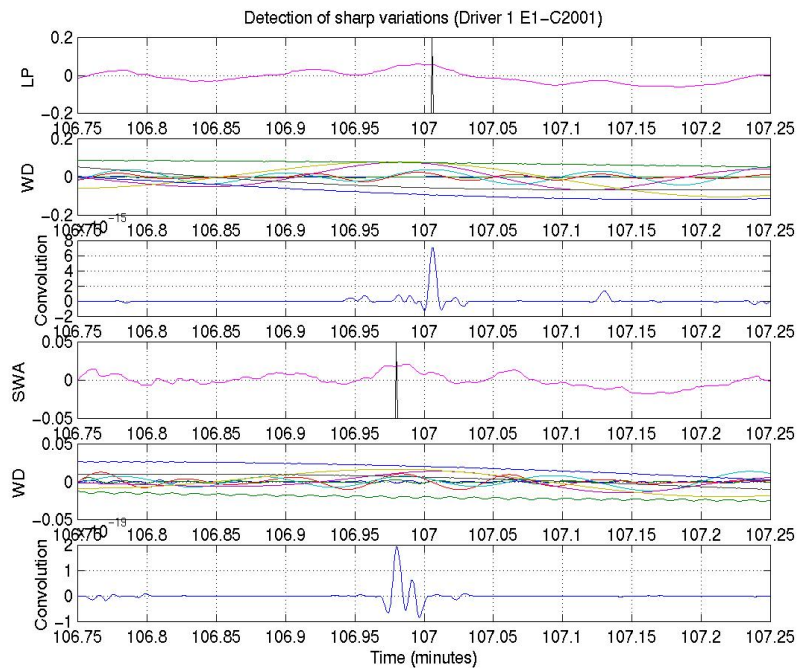


Figure IV.31 : Détection de ruptures dans la position latérale (LP) et l'angle du volant (SWA) utilisant l'ondelette B-Spline.

TABLEAU IV.10

Variables statistiques et fréquentielles extraites pour la caractérisation du mode de conduite.

1. Mean LP	1. Skewness LP	41. Mean SWA	61. Skewness SWA
2. Mean LP1	2. Skewness LP1	42. Mean SWA1	62. Skewness SWA1
3. Mean LP2	3. Skewness LP2	43. Mean SWA2	63. Skewness SWA2
4. Mean LP3	4. Skewness LP3	44. Mean SWA3	64. Skewness SWA3
5. Std LP	5. Kurtosis LP	45. Std SWA	65. Kurtosis SWA
6. Std LP1	6. Kurtosis LP1	46. Std SWA1	66. Kurtosis SWA1
7. Std LP2	7. Kurtosis LP2	47. Std SWA2	67. Kurtosis SWA2
8. Std LP3	8. Kurtosis LP3	48. Std SWA3	68. Kurtosis SWA3
9. Std d(LP)	9. Central frequency (cog) LP	49. Std d(SWA)	69. Central frequency (cog) SWA
10. Std d(LP)1	10. Central frequency (cog) LP1	50. Std d(SWA)1	70. Central frequency (cog) SWA1
11. Std d(LP)2	11. Central frequency (cog) LP2	51. Std d(SWA)2	71. Central frequency (cog) SWA2
12. Std d(LP)3	12. Central frequency (cog) LP3	52. Std d(SWA)3	72. Central frequency (cog) SWA3
13. Max d(LP)	13. Bandwidth LP	53. Max d(SWA)	73. Bandwidth SWA
14. Max d(LP)1	14. Bandwidth LP1	54. Max d(SWA)1	74. Bandwidth SWA1
15. Max d(LP)2	15. Bandwidth LP2	55. Max d(SWA)2	75. Bandwidth SWA2
16. Max d(LP)3	16. Bandwidth LP3	56. Max d(SWA)3	76. Bandwidth SWA3
17. Nrg LP	17. root mean square LP	57. Nrg SWA	77. root mean square SWA
18. Nrg LP1	18. root mean square LP1	58. Nrg SWA1	78. root mean square SWA1
19. Nrg LP2	19. root mean square LP2	59. Nrg SWA2	79. root mean square SWA2
20. Nrg LP3	20. root mean square LP3	60. Nrg SWA3	80. root mean square SWA3
81. mean Correction time R	82. mean Correction time L	83. STD Correction time R	84. STD Correction time L

Les variables artificielles calculées dans chaque intervalle permettent d'obtenir une information significative sur la position latérale et l'angle du volant. Le Tableau IV.10 présente les indices calculés pour l'analyse des différentes variables pertinentes. Chaque variable est calculé en utilisant trois bandes de fréquence intermédiaires définies par les coefficients en ondelettes ; par exemple, pour l'écart type de la position latérale (variable 5), nous calculons également l'écart type des coefficients en ondelettes de fréquence centrale 0.08 Hz (variable 6), pour 0.16 Hz (variable 7) et pour 0.32 Hz (variable 8).

IV.7.1.3 Fusion des caractéristiques

La classe normale est représentée par une fusion des caractéristiques extraites de la phase d'analyse temporelle et fréquentielle des signaux provenant des capteurs mécaniques. Pour la construction de cet ensemble de vecteurs, nous avons utilisé la méthode des machines à vecteurs de support pour l'estimation de la densité (SVDE), présentée dans le chapitre III, [68], [69].

Comme discuté dans la section III.2.3.4, le processus d'apprentissage d'une SVM se ramène à résoudre le problème quadratique suivant :

¡Error! No se pueden crear objetos modificando códigos de campo. (VI.7)

¡Error! No se pueden crear objetos modificando códigos de campo. (IV.8)

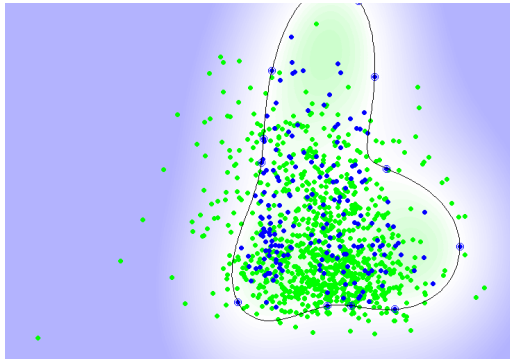
¡Error! No se pueden crear objetos modificando códigos de campo. (IV.9)

La fonction de décision associée est définie comme :

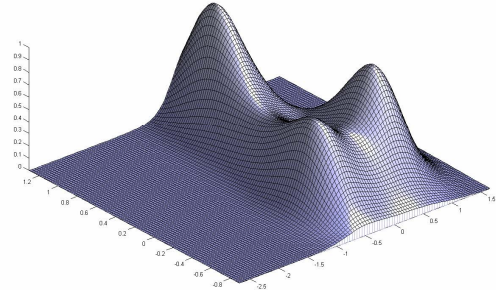
¡Error! No se pueden crear objetos modificando códigos de campo. (IV.10)

où les exemples $\mathbf{x}_i$ sont des vecteurs formés par les caractéristiques obtenues pendant le période de référence. Le calcul de la valeur de b est fait selon l'expression suivante :

¡Error! No se pueden crear objetos modificando códigos de campo. (IV.11)



a)



b)

Figure IV.32 : Exemple à deux dimensions de la construction de la classe normale. a) Borne apprise par les SVM b) Modèle de densité associée.

La fonction de décision (IV.11) donne une valeur positive si l'exemple de conduite correspond à une conduite vigilante et une valeur négative dans le cas contraire. Notre but étant d'évaluer une conduite anormale, nous modifions la fonction de décision (IV.11) pour avoir une fonction donnant des valeurs bornées entre [0,1] liées à une conduite anormale :

$$f(\mathbf{x}) = \max \left(\frac{b - \sum_{i=1}^{MSV} \alpha_i k(\mathbf{x}_i, \mathbf{x})}{b}, 0 \right). \quad (IV.12)$$

Si $f(\mathbf{x})$ est zéro, $\mathbf{x}$ correspond à une conduite normale et si $f(\mathbf{x})$ est un, $\mathbf{x}$ correspond à une conduite anormale. La Figure IV.32 montre un exemple pour des caractéristiques à deux dimensions.

Comme le but est de donner une évaluation à la minute, nous calculons un diagnostic à la minute en utilisant une moyenne des évaluations de tous les intervalles adjoints à la minute précédente, pondérée par la durée de l'intervalle associée à chaque évaluation :

$$h = \frac{\sum_{j=1}^{k_v} d_j f_j(\mathbf{x})}{\sum_{j=1}^{k_v} d_j} \quad (\text{IV.13})$$

IV.7.1.3 Diagnostic cumulé

La difficulté de comparer le diagnostic instantané fourni par le système avec le niveau de vigilance fourni par le médecin, principalement à cause du décalage, nous a conduit à calculer un diagnostic cumulé, [134], dont l'objectif est de fournir une tendance globale de la dégradation de la performance de conduite. Pour calculer cette tendance, on considère que toute sortie de la zone normale de l'espace des caractéristiques constitue une alerte, dont l'effet sur le diagnostic cumulé dépendra de l'amplitude et de la durée, c'est à dire de la taille de la fenêtre d'analyse courante. Ainsi, une alerte isolée aura un effet peu important sur le diagnostic cumulé, qui sera, par contre, fortement influencé par plusieurs alertes consécutives.

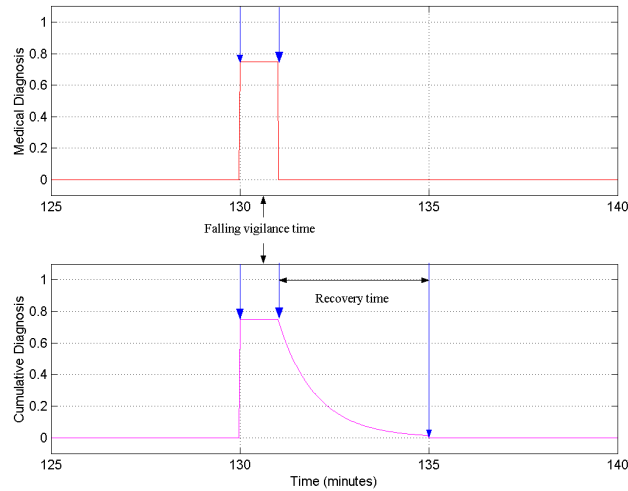


Figure IV.33 : Hypothèse de récupération de la vigilance lors de une phase de baisse vigilance observée.

Le modèle utilisé pour représenter l'effet d'une alerte dans le temps, Figure IV.33, est une exponentielle dont la constante de temps dépend de la durée de l'alerte. On trouve ce type de phénomène sur certains systèmes biologiques, par exemple dans le cas de la rétine soumise à une transition abrupte de lumière, [134].

Les deux évolutions possibles du diagnostic cumulé sont : une **phase d'accumulation** si le diagnostic instantané est différent de zéro, et une **phase de récupération** si le diagnostic instantané est égal à zéro. Dans la phase d'accumulation, le diagnostic cumulé courant est égal au diagnostic cumulé précédent, augmenté de la valeur du diagnostic instantané pondéré par la durée de la fenêtre d'analyse. Le temps cumulé de non vigilance est actualisé, en ajoutant la durée de la fenêtre d'analyse courante au temps nécessaire au diagnostic cumulé précédent pour retourner au niveau initial. Dans la phase de récupération, la fraction récupérée est calculée à partir de l'expression (IV.14). Dans cette expression, le temps de récupération courant est égal au temps de récupération précédent augmenté de la durée de la fenêtre d'analyse courante. k_d est la constante de temps de récupération. Par exemple, pour un temps cumulé de non vigilance de 3 secondes, $k_d = 1$ conduirait à un temps nécessaire pour revenir au niveau de vigilance initial d'environ 12 secondes. Avec $k_d = 2$, ce temps ne serait que de 6 secondes.

$$\text{;Error! No se pueden crear objetos modificando códigos de campo..} \quad (\text{IV.14})$$

Le temps est exprimé en secondes. Lorsque la fraction de récupération atteint la valeur 1, alors, le temps cumulé de non vigilance et le temps de récupération sont remis à zéro.

IV.7.2 Evaluation du HDM performance

L'étape la plus critique dans le processus de développement est la sélection de l'ensemble des caractéristiques capables de différencier une conduite vigilante d'une conduite hypovigilante. Ceci est d'autant plus vrai que nous n'avons aucune information *a priori* de cet ensemble.

IV.7.2.1 Mesures de référence de la vigilance

La référence médicale du niveau de vigilance est donnée toutes les 16 secondes, quand elle est disponible (référence *physiologique*). Elle est basée sur l'analyse fréquentielle à quatre niveaux, présentée dans le chapitre I. La référence *subjective*, est donnée par une autoévaluation utilisant l'échelle KSS, également définie dans le chapitre I.

Pour qualifier le modèle de vigilance, nous avons transformé ces références en des références plus simples :

TABLEAU IV.11
Références simplifiées EEG et KSS.

		Niveau EEG	Niveau KSS
0	Normal	0	≤ 5
1	Indéterminé	1	5 – 6
2	Anormal	2–3	$7 \leq$

Le Tableau IV.12 présente une combinaison des expertises EEG et KSS simplifiés qui génère une référence unique. Le niveau –1 indique l'absence de cette référence.

Ces trois références sont utilisées pour déterminer le nombre de non détections et de fausses alarmes, considérant que si une détection arrive quand le conducteur est qualifié vigilant, cette détection est alors classifiée comme fausse alarme (*F*). Dans le cas opposé, il y aura une non détection si le conducteur est déclaré hypovigilant et il n'y a pas de détection. A un certain niveau de vigilance, la présence ou l'absence d'une alarme n'est pas critique car le conducteur peut accepter un niveau d'alarme, même s'il n'est pas suffisamment somnolent.

TABLEAU IV.12
Référence combinée basée en l'EEG et le KSS.

EEG \ KSS	0	1	2
	0	1	1
0	0	1	1
1	0	1	2
2–3	2	2	2
–1	0	1	2

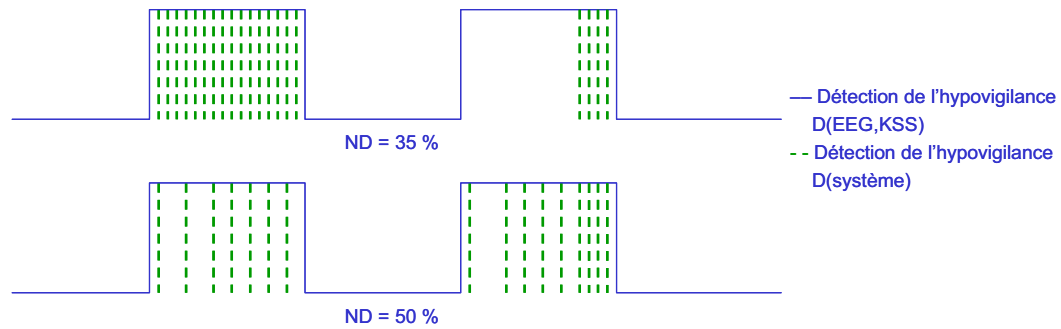


Figure IV.34 : Influence du nombre de non détections par rapport à la fréquence

De cette manière, un bon système de détection de l'hypovigilance aura un faible taux de fausses alarmes (0%–5%). Néanmoins, un pourcentage de non détections de 50% peut être un résultat satisfaisant si les détections apparaissent durant toutes les phases d'hypovigilance, Figure IV.34. Cette valeur semble raisonnable car durant les phases d'hypovigilance le conducteur peut conduire d'une manière toute à fait normale.

IV.7.2.2 Validation

Pour cette phase de validation, nous avons testé trois approches :

- l'analyse de la corrélation directe entre les variables artificielles elles-mêmes et une deuxième analyse entre les variables artificielles et les références (en rappelant des problèmes de décalage entre l'apparition de l'hypovigilance et l'apparition du changement de style de conduite).
- le test des combinaisons à deux, trois et quatre variables artificielles pour minimiser le nombre de fausses alarmes et le nombre de non détections, complété par une analyse visuelle du comportement des variables par rapport aux signaux de référence.

Analyse des corrélations

L'analyse de la corrélation entre les différentes variables statistiques et fréquentielles est essentielle pour ne garder que les informations qui, réellement, ont une contribution sur la détection de l'hypovigilance.

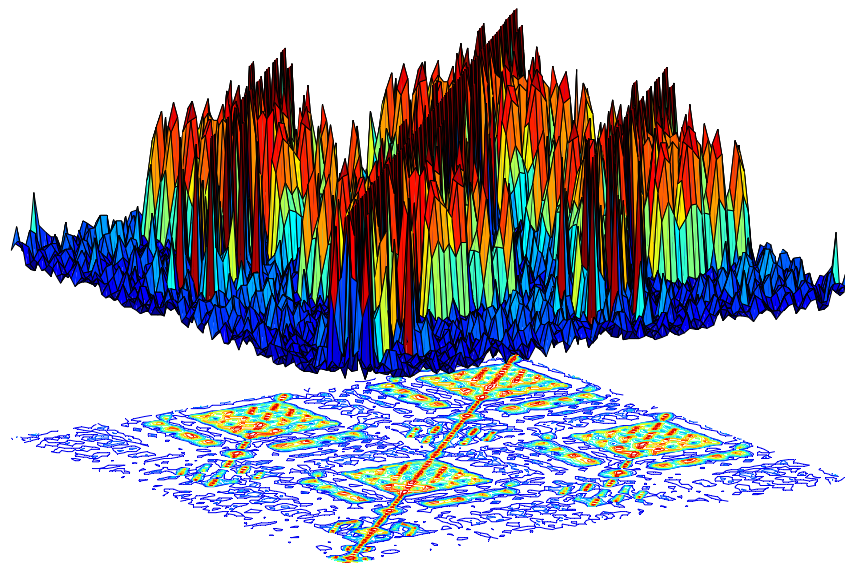


Figure IV.35 : Coefficient de corrélation de l'ensemble formé par les caractéristiques de conduite et les références (bornes extrêmes). La ligne diagonale représente le niveau 1 (corrélation d'une variable par rapport à elle-même).

L'analyse de la corrélation entre les variables artificielles montre qu'il existe une forte dépendance (supérieur à 0.46) entre les variables 5 à 20, 26 à 28, 45 à 60 et 66 à 68 du Tableau IV.10. La Figure IV.35 est un exemple de corrélation entre ces variables sur les bases de données du conducteur C01. Après une analyse en composantes principales sur les conducteurs C01–C08, nous avons gardé l'information suivante :

TABLEAU IV.13
Variables statistiques et fréquentielles choisies après analyse de la corrélation.

1. Std LP	11. Bandwidth LP	19. Std SWA	29. Bandwidth SWA
2. Skewness LP	12. Bandwidth LP1	20. Skewness SWA	30. Bandwidth SWA1
3. Skewness LP1	13. Bandwidth LP2	21. Skewness SWA1	31. Bandwidth SWA2
4. Skewness LP2	14. Bandwidth LP3	22. Skewness SWA2	32. Bandwidth SWA3
5. Skewness LP3	15. root mean square LP	23. Skewness SWA3	33. root mean square SWA
6. Kurtosis LP	16. root mean square LP1	24. Kurtosis SWA	34. root mean square SWA1
7. Central frequency (cog) LP	17. root mean square LP2	25. Central frequency (cog) SWA	35. root mean square SWA2
8. Central frequency (cog) LP1	18. root mean square LP3	26. Central frequency (cog) SWA1	36. root mean square SWA3
9. Central frequency (cog) LP2		27. Central frequency (cog) SWA2	
10. Central frequency (cog) L3		28. Central frequency (cog) SWA3	
37. mean Correction time R	38. mean Correction time L	39. STD Correction time R	40. STD Correction time L

L'analyse de la corrélation montre qu'aucune variable n'est associée aux références avec un coefficient de corrélation supérieur à 0.38 pour la référence physiologique, 0.22 pour la référence subjective et 0.29 pour la référence unique, dans le meilleur des cas, pour l'ensemble des conducteurs C01–C08. La Figure IV.35 montre les coefficients de corrélation de l'ensemble formé par les caractéristiques de conduite et les références. Ces dernières se trouvent aux extrêmes de la figure.

Test des combinaisons et analyse visuelle

Sur l'ensemble des 40 variables restantes, nous avons choisi la meilleure combinaison à deux, à trois et à quatre variables, comme proposé par Santana, [134]. Il faut remarquer que les variables sélectionnées sont celles qui représentent le meilleur comportement pour tous les conducteurs.

Le Tableau IV.14 montre les variables finalement choisies ayant une influence dans le diagnostic.

TABLEAU IV.14
Variables statistiques et fréquentielles choisies après analyse visuelle.

1. Std LP	2. Central frequency (cog) LP	3. Kurtosis SWA	4. mean Correction time R
5. Skewness LP	6. Bandwidth LP	7. Central frequency (cog) SWA	8. mean Correction time L
9. Skewness LP1	10. root mean square LP	11. Bandwidth SWA	
12. Kurtosis LP	13. Std SWA	14. root mean square SWA	

IV.7.2.3 Résultats

A partir des essais réalisés lors des campagnes 1 et 2, nous avons fait une analyse de l'évolution de la conduite pour chacun des conducteurs. Nous utilisons les diagnostics de référence et les diagnostics système afin de calculer le pourcentage de fausses alarmes, le pourcentage de non détections et le pourcentage de détections système dans une zone indéterminé pour les trois références.

Chaque figure représente tous les essais faits par un seul conducteur. Les lignes verticales en magenta représentent l'instant où se trouvent la fin d'un trajet et le début d'un autre. L'axe des abscisses représente le temps en minutes. En ordonnées, nous avons, d'une part, les références :

- l'autoévaluation du conducteur (KSS) à 10 niveaux,
- l'évaluation physiologique (EEG), entre 0 et 3 quand cette information est disponible, -1 quand l'information est invalide ou nulle (par exemple, pour le conducteur C01 il existe des informations EEG invalides entre la minute 180 et la minute 240, Figure IV.36),
- la référence combinée (EK) à 3 niveaux, suivant l'information représenté dans le Tableau IV.12,

d'autre part, les diagnostics système :

- le diagnostic système moyenné à la minute, entre 0 et 1 pour les états vigilant et hypovigilant respectivement et -1 lors de la phase d'apprentissage ou d'une perte des données d'acquisition,
- le diagnostic système cumulé, entre 0 et 1 pour les états vigilant et hypovigilant respectivement.

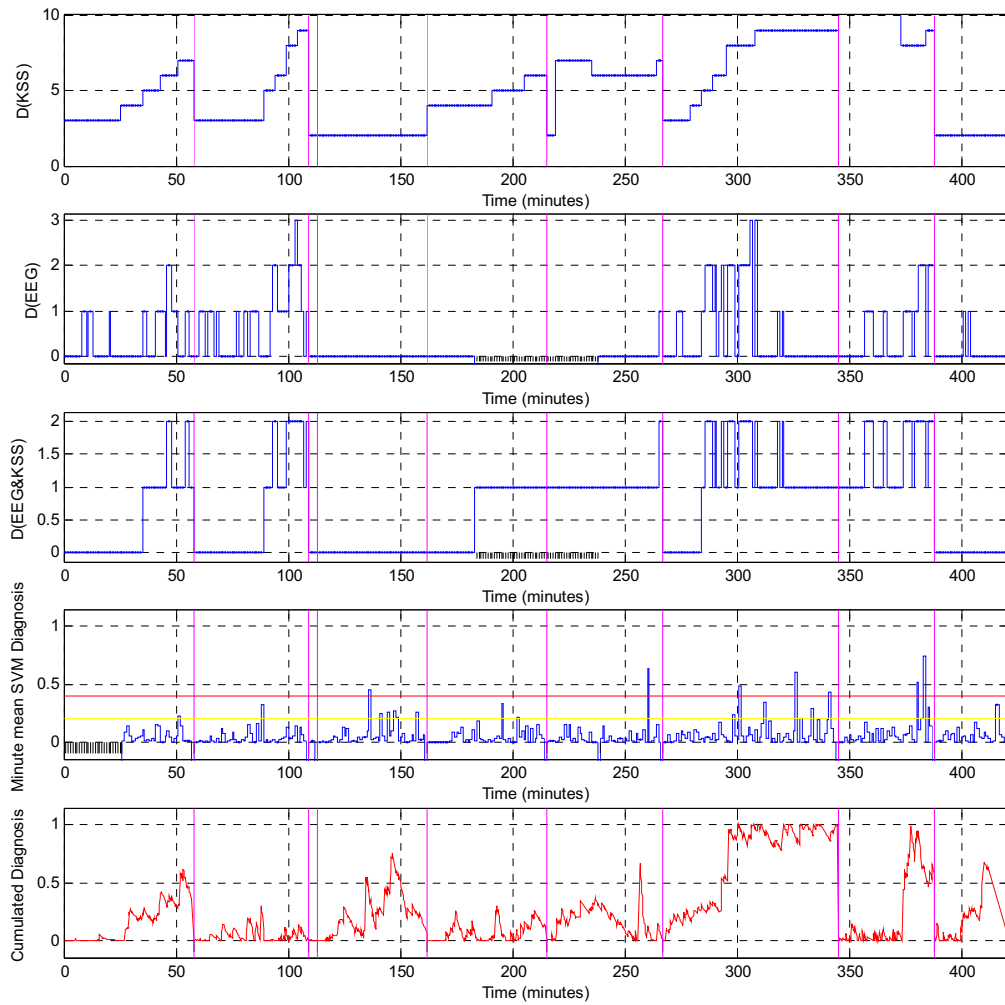


Figure IV.36 : Résultats pour le conducteur C01.

TABLEAU IV.15
Résultats pour le conducteur C01.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	4.669%	3.825%	4.094%	90.32%	91.47%	94.23%	0%	3.333%	3.911%
Diagnostic cumulé	0%	0%	0%	12.9%	6.977%	11.54%	0%	0%	1.334%

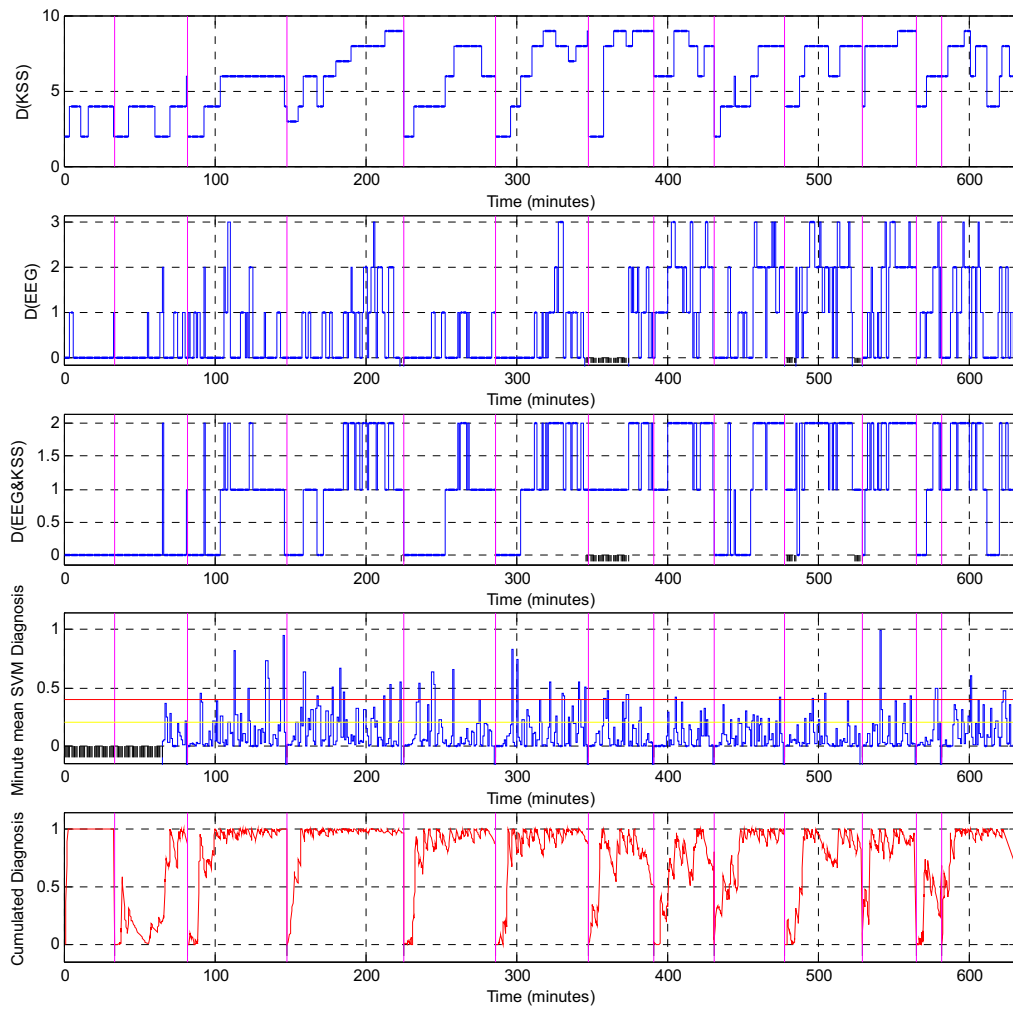


Figure IV.37 : Résultats pour le conducteur C02.

TABLEAU IV.16
Résultats pour le conducteur C02.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	5.747%	3.731%	3.788%	77.78%	82.22%	78.13%	4.255%	10%	8%
Diagnostic cumulé	0%	0%	0%	27.78%	24.44%	16.67%	0%	0%	2.818%

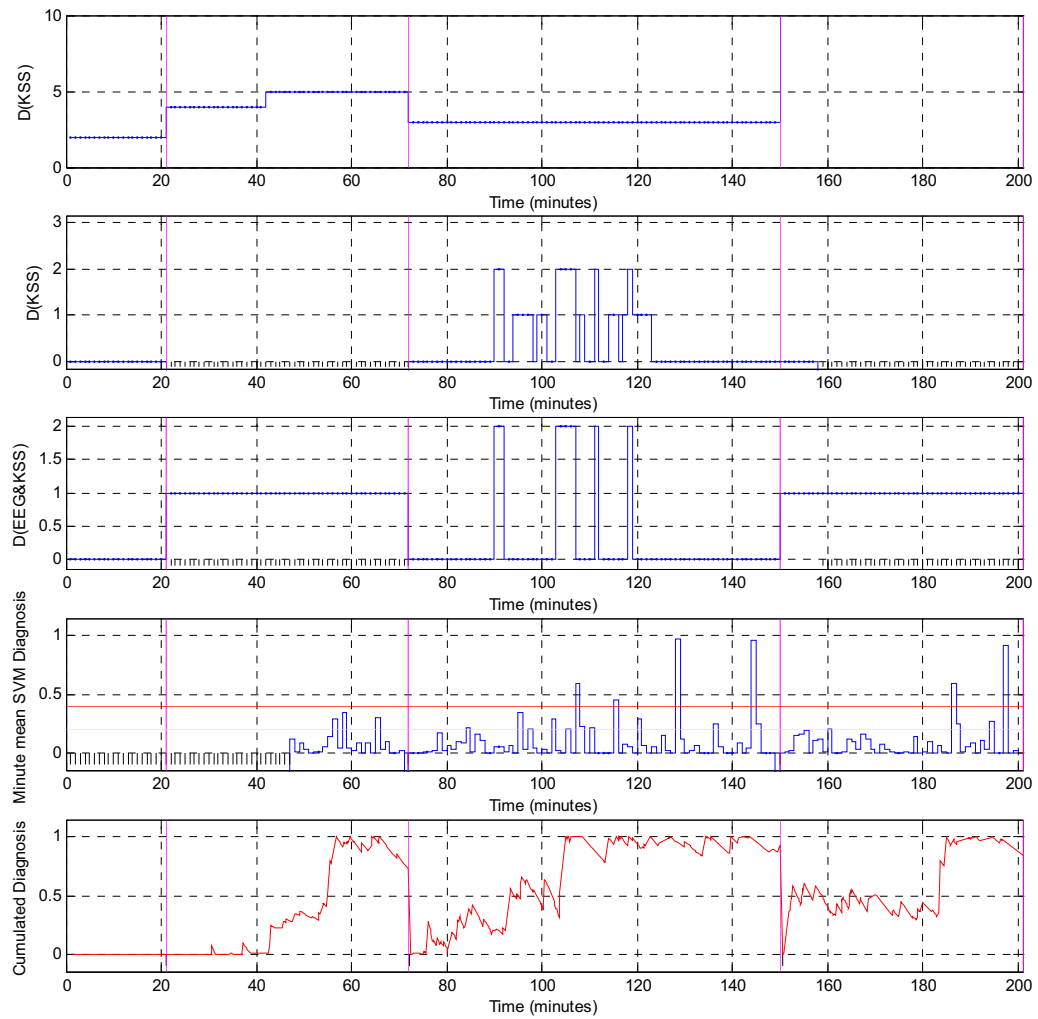


Figure IV.38 : Résultats pour le conducteur C03.

TABLEAU IV.17
Résultats pour le conducteur C03.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	7.576%	9.901%	9.722%	62.5%	88.24%	62.5%	14.29%	10%	5.882%
Diagnostic cumulé	0%	0%	0%	62.5%	0%	0%	0%	0%	0%

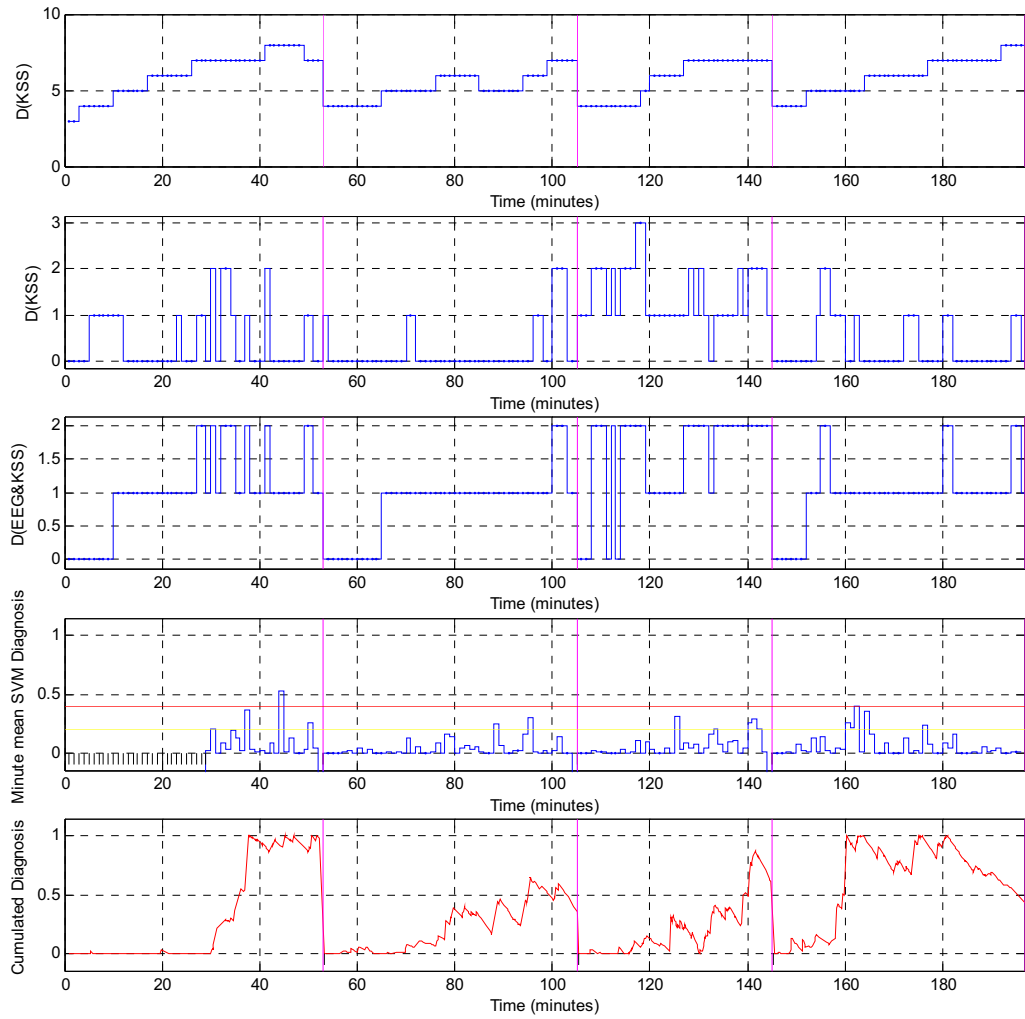


Figure IV.39 : Résultats pour le conducteur C04.

TABLEAU IV.18
Résultats pour le conducteur C04.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	7.547%	0%	0%	92%	83.1%	88.89%	7.407%	8.333%	9.322%
Diagnostic cumulé	1.887%	0%	0%	48%	56.34%	71.43%	1.852%	0%	1.714%

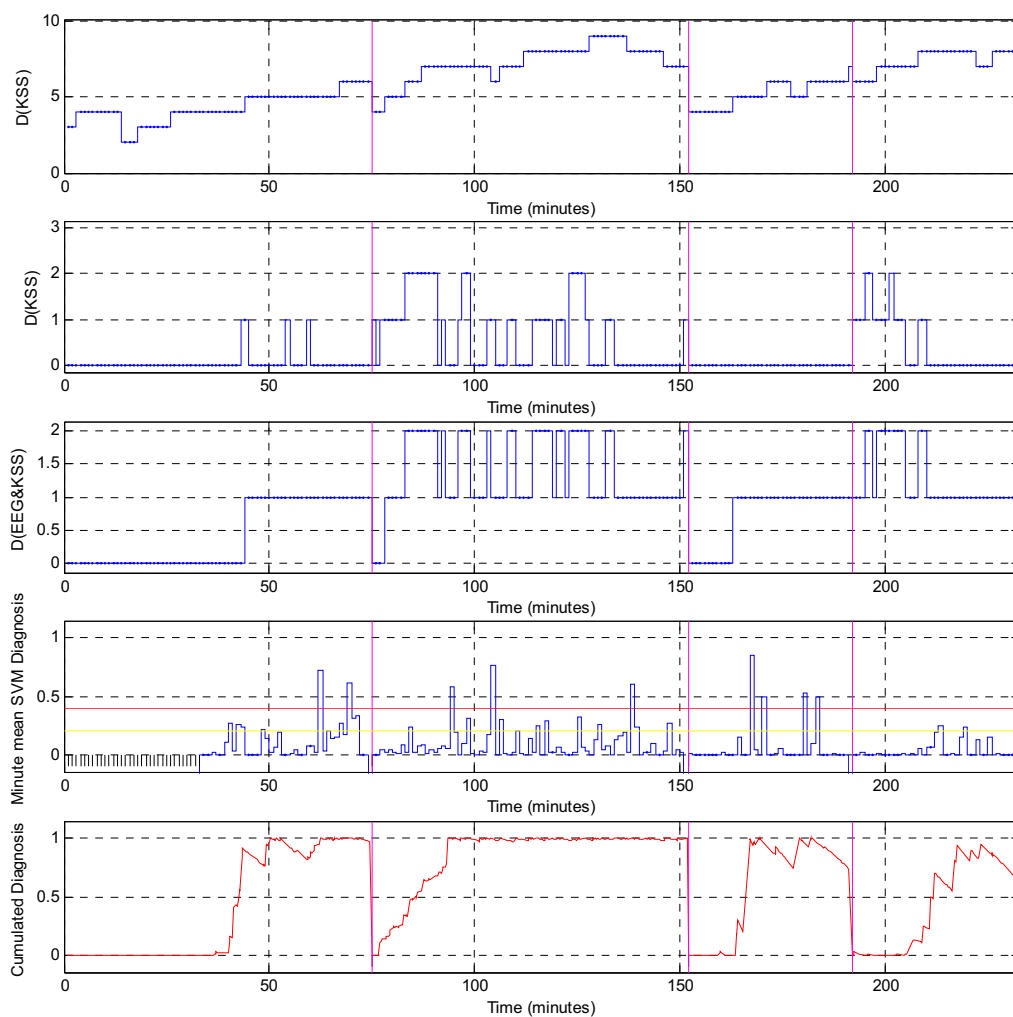


Figure IV.40 : Résultats pour le conducteur C05.

TABLEAU IV.19
Résultats pour le conducteur C05.

[illegible]

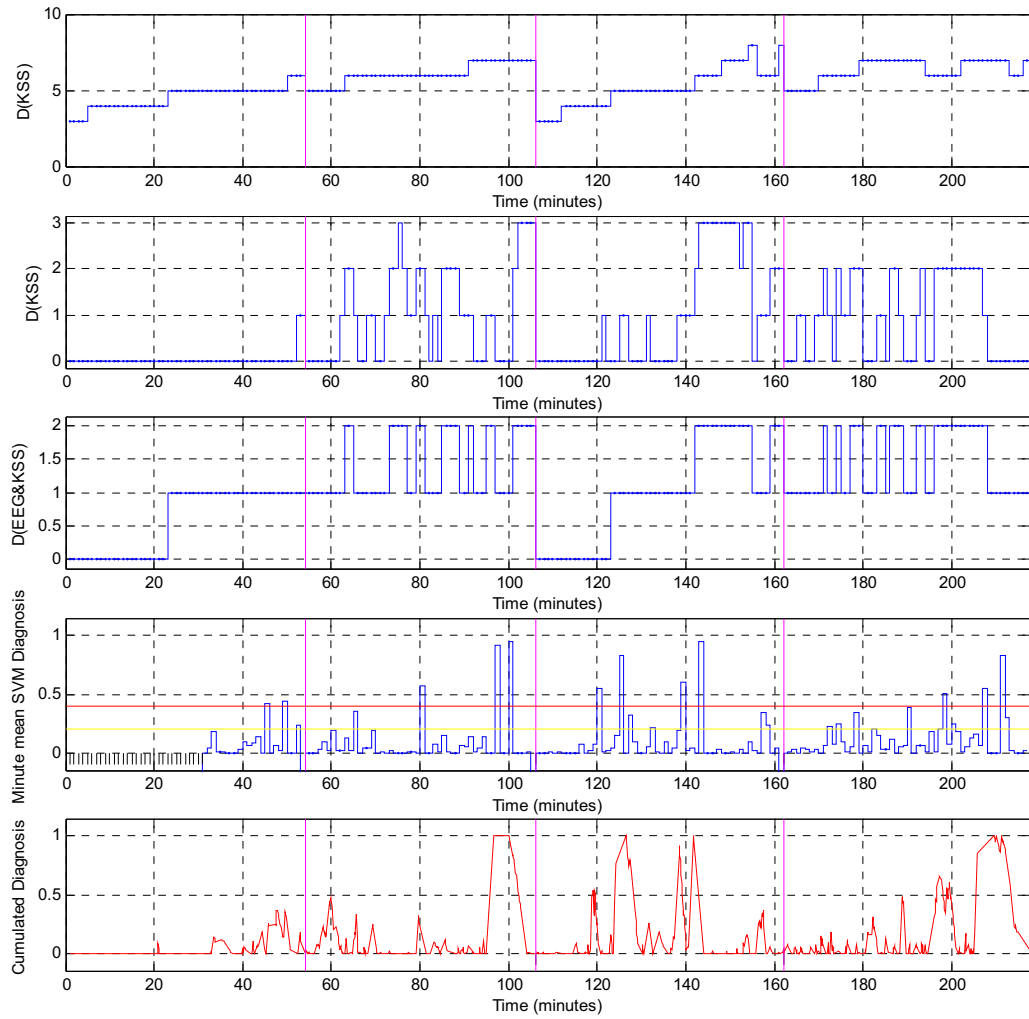


Figure IV.41 : Résultats pour le conducteur C06

TABLEAU IV.20
Résultats pour le conducteur C06.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	5.556%	0%	0%	86.54%	86.54%	86.67 %	7.627%	8.73 %	10.26%
Diagnostic cumulé	0%	0%	0%	46.15%	48.08%	50 %	0%	0%	1.056%

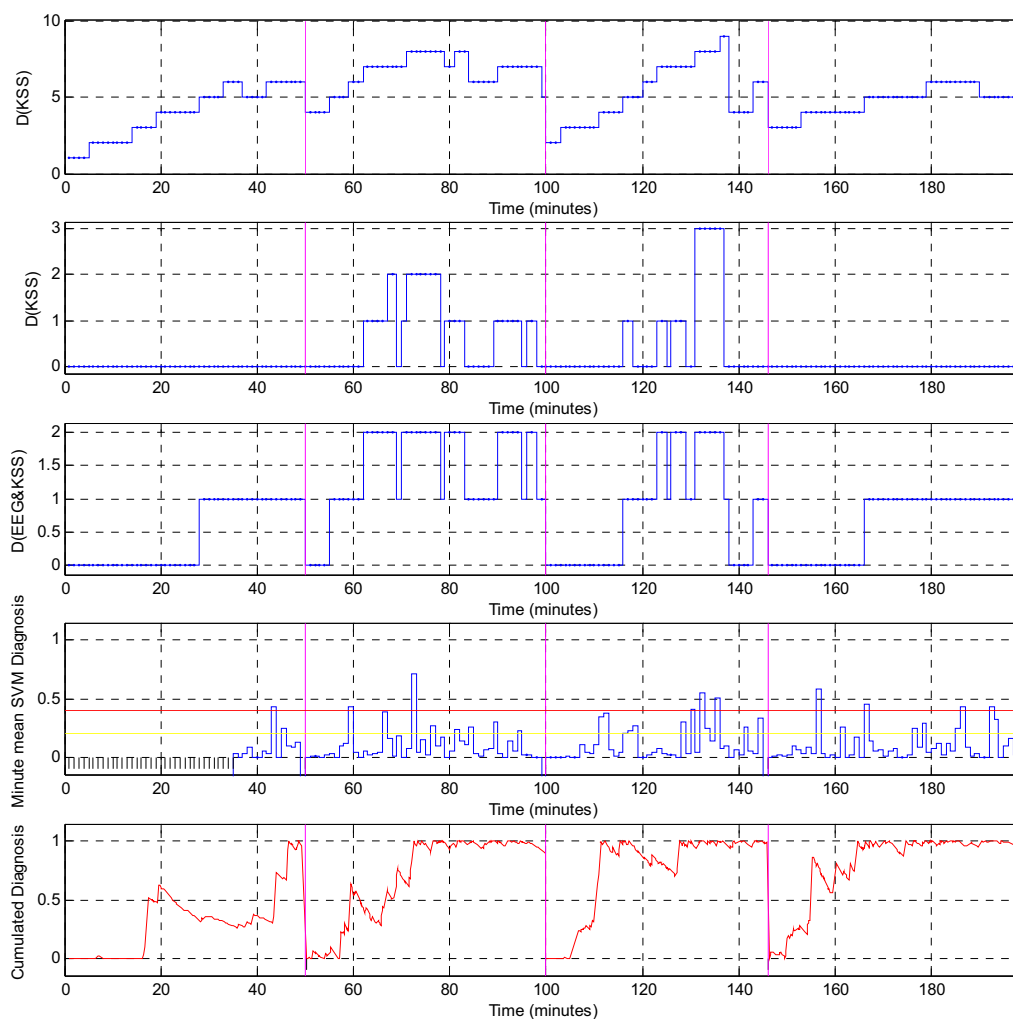


Figure IV.42 : Résultats pour le conducteur C07.

TABLEAU IV.21
Résultats pour le conducteur C07.

[illegible]

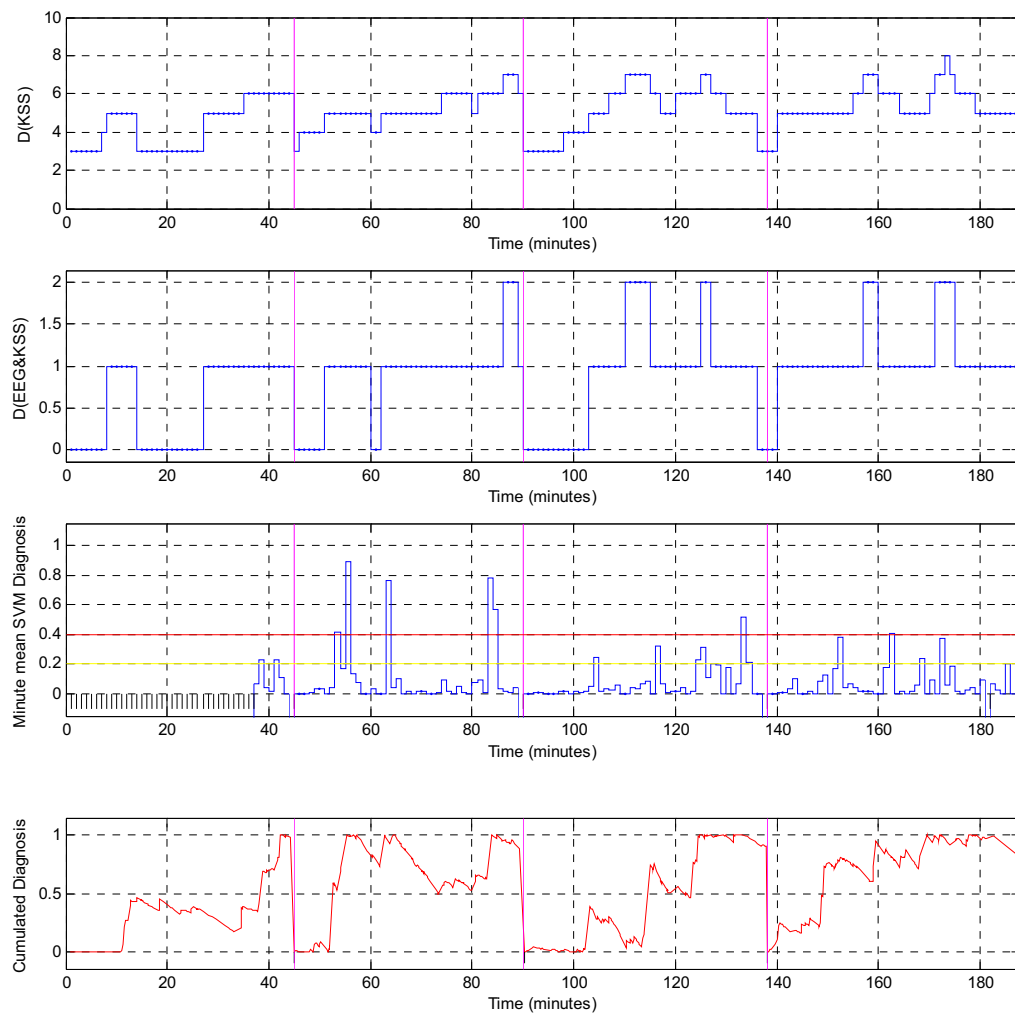


Figure IV.43 : Résultats pour le conducteur C08.

TABLEAU IV.22
Résultats pour le conducteur C08.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	-	0%	0%	-	82.35%	82.35%	-	7.2%	7.2%
Diagnostic cumulé	-	0%	0%	-	0%	0%	-	2.4%	6.597%

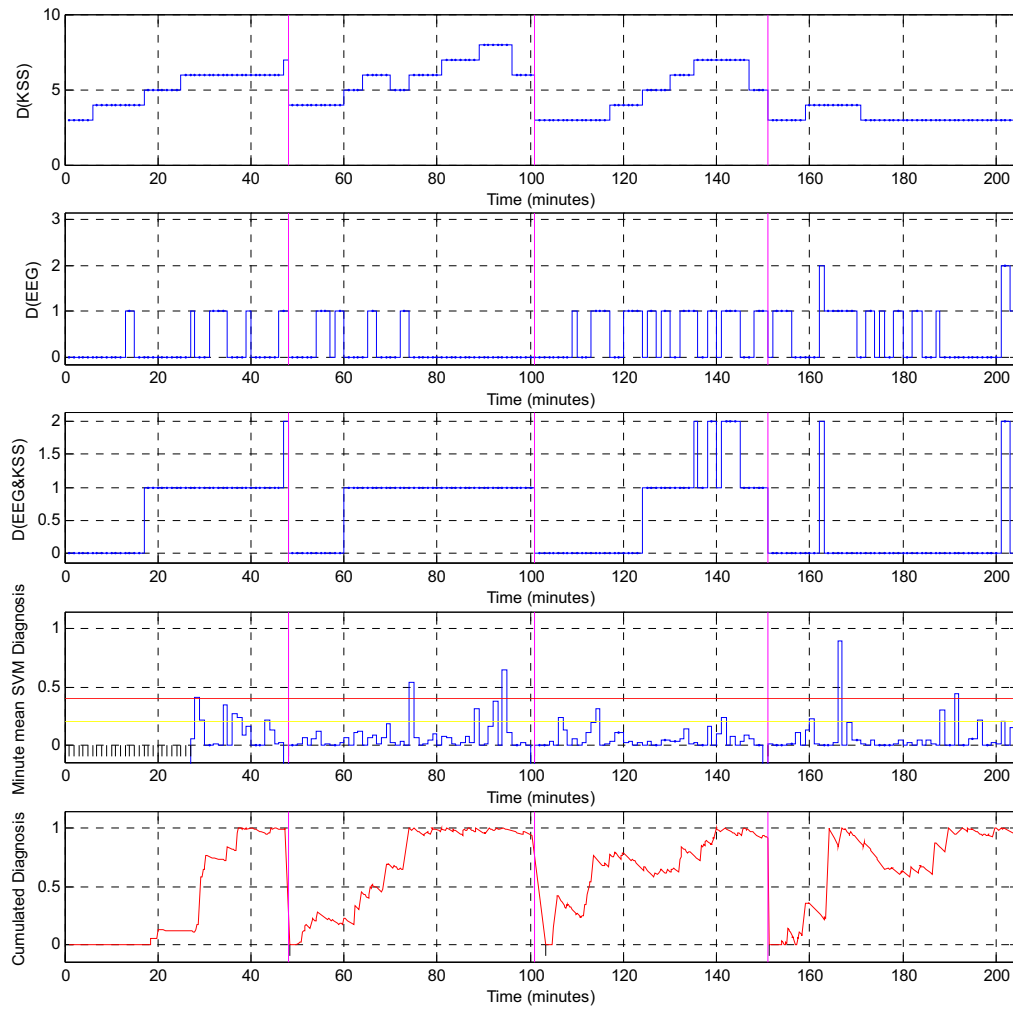


Figure IV.44 : Résultats pour le conducteur C09.

TABLEAU IV.23
Résultats pour le conducteur C09.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	9.244%	6.742%	5.814%	66.67%	85.71%	81.82%	3.03%	7.042%	7.692%
Diagnostic cumulé	7.563%	10.11%	12.5%	0%	0%	0%	0%	0%	4.241%

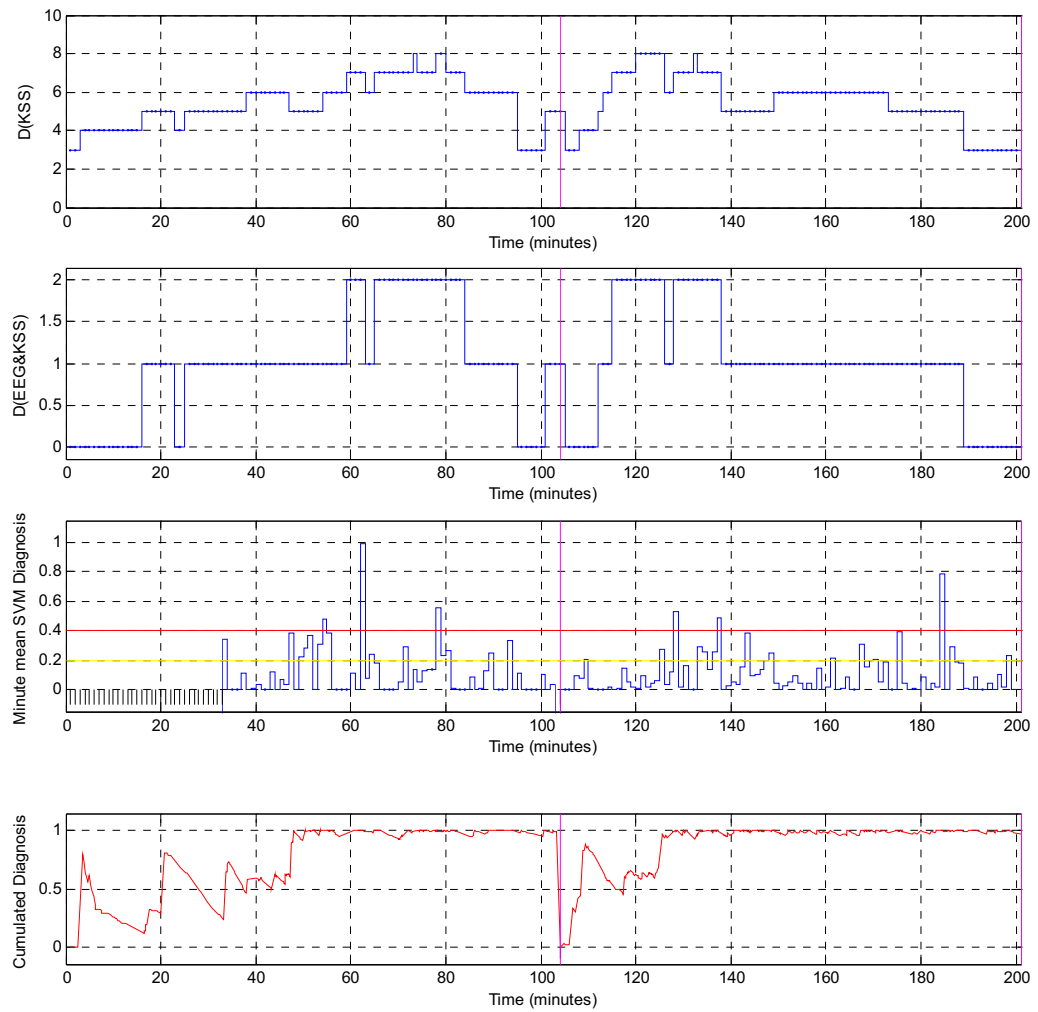


Figure IV.45 : Résultats pour le conducteur C10.

TABLEAU IV.24
Résultats pour le conducteur C10.

	% Fausse Alarmes			% Non Détections			% Zone incertaine		
	EEG	KSS	EK	EEG	KSS	EK	EEG	KSS	EK
Diagnostic moyenné	-	7.407%	7.407%	-	79.55%	79.55%	-	14.91%	14.91%
Diagnostic cumulé	-	25.93%	26.92%	-	0%	0%	-	10.53%	21.71%

IV.8 Conclusion

Dans ce chapitre, nous avons appliquée les concepts sur l'interaction Homme-Machine au problème de la surveillance du conducteur automobile. Nous avons développé une approche méthodologique de supervision temps-réel basée sur les méthodes et outils présentés au cours des chapitres II et III, particulièrement pour la détection des situations à risque (à dynamique lente) et pour le diagnostic de l'hypovigilance (à dynamique lente).

Nous avons présenté les systèmes d'assistance à la conduite, ayant pour but d'aider le conducteur à mieux comprendre son environnement et de l'assister dans sa prise de décision. Ceci nous a amené à étudier un modèle générique de la tâche de conduite pour pouvoir aborder le problème de la perception selon le niveau de contrôle à réaliser. Ensuite, nous avons repéré les facteurs qui interviennent dans la surveillance du conducteur automobile et les types de mesures que l'on peut en réaliser.

Nous avons commencé par une synthèse des approches menées pour le développement de systèmes d'assistance à la conduite automobile, en général. Suite à quoi, nous avons présenté un rappel sur les systèmes d'assistance au conducteur, pour avertir de possibles situations de danger et de la détection de fatigue. Pour cela, nous avons présenté les objectifs des projets AWAKE et PREDIT, dont l'un des objectifs est la conception d'un système de détection temps-réel de l'hypovigilance. Après une analyse des deux architectures, nous avons limité notre démarche au développement de deux modules principaux : un module de diagnostic de situations à risque et un module de diagnostic de l'hypovigilance. Ces efforts sont la suite des projets européens PROMETHEUS et SAVE, du projet national PREDIT phase I et du projet régional Hypovigil, tous accompagnés de quatre thèses LAAS.

Pour le module de diagnostic de situations à risque, nous avons développé une architecture spécifique capable de détecter quand le conducteur est sur le point de sortir de la route et de déclencher une alarme pour l'avertir. Nous avons proposé trois stratégies, dont seulement une prend en compte les actions du conducteur. La comparaison de ces approches à travers les bases de données des essais a mis en évidence la supériorité de notre système basée sur des mesures simples et un système d'inférence floue, capable de représenter les heuristiques liées à ce type de problème de surveillance des situations à risque.

Finalement, nous avons développé notre propre approche pour la mise en œuvre du module de diagnostic de l'hypovigilance à partir de mesures de performance de la conduite. Nous avons présenté une architecture particulière basée sur une analyse en ondelettes des signaux, ce qui définit des intervalles d'analyse où le signal présente une dynamique stable. De cette manière, l'extraction des caractéristiques statistiques et fréquentielles représentant la conduite s'avère plus adéquat, évitant le mélange des informations caractéristiques d'une conduite normale et une anormale. Les machines à vecteurs de support ont montré leur pertinence pour la fusion de ces informations pour en induire une mesure du degré d'anormalité de la conduite. Nous avons testé les performances de ce système sur les bases de données des essais pour évaluer ce système en termes de fausses alarmes et de non détections, à partir d'un ensemble de mesures de référence physiologiques et subjectives.

Conclusion Générale

Nous avons proposé, dans cette thèse, une méthodologie pour aborder la supervision de systèmes Homme–Machine, qui nécessitent une surveillance automatisée globale, notamment pour la détection de défaillances humaines. Cette approche a pour objectif principal la conception de systèmes d’assistance intelligents pouvant fournir des informations pertinentes à l’opérateur sur l’évolution de la performance du système HM entier et l’alerter en cas de risque. Il s’agit d’un objectif difficile qui fait l’objet de plusieurs travaux, particulièrement pour son application dans les transports : l’aéronautique, les transports ferrés, l’automobile, etc.

La complexité de l’interaction HM oblige la prise en compte des facteurs liés aux dynamiques de la machine et à l’environnement pour caractériser la performance de l’opérateur et, donc, son possible état de vigilance. La mise au point d’un système de diagnostic temps–réel implique tout un travail théorique nécessitant une collaboration pluridisciplinaire.

Cette méthodologie est appliquée à la conception d’un système temps–réel pour l’assistance au conducteur dans la conduite automobile. Elle s’appuie sur l’analyse en temps et en fréquence de l’évolution des signaux mécaniques, sur des statistiques et sur des techniques pour la fusion de l’information.

Les travaux ici présentés, ont été réalisés dans le cadre des programmes : Européen AWAKE et le national français PREDIT. Notre responsabilité principale a été celle de développer un module de diagnostic de l’hypovigilance d’un conducteur automobile. Cette démarche pluridisciplinaire, a été réalisée en collaboration de partenaires français et européens issues des différentes disciplines du savoir.

Nous avons commencé, dans le *premier chapitre*, par établir les avantages et les inconvénients de l’impact de l’automatisation dans les systèmes complexes technologiques vis-à-vis de l’interaction de la partie technologique à l’égard de l’opérateur humain. Cela nous a permis d’aborder les aspects de l’interaction homme–machine, notamment leurs différents niveaux de coopération et de partage de tâches. Nous nous sommes centré sur une analyse des défaillances pouvant apparaître dans les systèmes HM, notamment provenant de la partie humaine. Nous avons étudié les principaux facteurs humains impliqués dans l’interaction HM, ce qui nous a permis d’aborder les différentes métriques que nous pouvons exploiter : mesures physiologiques, mesures subjectives et mesures de performance.

A la fin de ce premier chapitre, nous avons introduit les structures de surveillance des systèmes homme–machine, visant la surveillance de la performance du système HM et les actions de l’opérateur humain dans une optique de supervision. Ainsi, nous avons présenté les éléments nécessaires à la construction de systèmes d’assistance intelligents, en commençant par identifier les objectifs du système de surveillance à réaliser, l’identification des informations nécessaires et leur utilisation pour pouvoir inférer une décision (diagnostic ou commande). Dans cette thèse, nous avons restreint l’approche proposée au diagnostic.

Dans le *deuxième chapitre*, nous avons présenté les outils nécessaires au prétraitement des informations provenant des capteurs. Nous avons mis en évidence l’intérêt d’avoir des informations « propres et explicatives » du phénomène à observer. De cette manière, nous pouvons étudier le comportement des signaux d’entrée en temps et en fréquence et en extraire des caractéristiques

statistiques et spectrales. Pour cela, nous avons réalisé une étude des différentes méthodes d'analyse de signaux, dans un cadre unidimensionnel.

Nous avons brièvement abordé l'analyse des signaux continus, pour les systèmes les plus simples, les systèmes LTI, et la transformée de Fourier ; nous avons expliqué l'importance du calcul des convolutions ainsi que leur implémentation à complexité moindre à l'aide de la transformée de Fourier rapide. Le calcul de la convolution rapide permet l'implémentation de filtres avec un temps de calcul réduit, essentiel pour nos besoins temps réel.

Ces concepts nous ont permis d'aborder l'analyse temporelle et fréquentielle des signaux, afin d'examiner leurs composantes transitoires, de durées différentes, à des fréquences différentes. Nous avons étudié les caractéristiques des deux méthodes les plus utilisées dans ce domaine : la transformée de Fourier à fenêtre et la transformée en ondelettes. Ainsi, le calcul d'ondelettes s'avère le plus adéquat pour nos besoins de détection de « singularités » dans les signaux, dépendant des évolutions temporelles et fréquentielles, et pour la « séparation » en différentes bandes de fréquence où les analyses statistiques et fréquentielles apportent des informations intéressantes. A nouveau, l'implémentation temps-réel de la transformée de Fourier à fenêtre, comme celle de la transformée en ondelettes, s'appuie sur les filtres à convolution rapide.

Dans le *troisième chapitre*, nous avons présenté les méthodes de modélisation à partir des données et de l'expertise, à travers des techniques d'apprentissage et de représentation de la connaissance respectivement, ainsi qu'à travers des techniques permettant l'utilisation des deux types d'information. Leur utilisation dans les systèmes temps-réel impose des contraintes de puissance de calcul, particulièrement pour les phases d'apprentissage, afin de conserver une certaine simplicité d'exploitation.

Dans un premier temps, nous avons abordé le cas des techniques d'apprentissage d'un point de vue neuronal. Partant de la théorie de l'apprentissage statistique, nous avons présenté les différentes approches dans un cadre linéaire, pour ensuite passer au cas non-linéaire. Ceci nous a donné les éléments nécessaires pour présenter les machines à vecteurs de support, approche avec des caractéristiques intéressantes mais avec des requêtes importantes de calcul et de stockage mémoire. C'est pour cela que nous avons approfondi nos recherches sur la phase d'apprentissage des SVM pour laquelle il faut résoudre un problème d'optimisation quadratique particulier. Pour ce faire, nous avons proposé des stratégies d'implémentation plus adaptées à l'algorithme d'optimisation quadratique et à l'algorithme de décomposition, pour gérer des problèmes de grande échelle. Nous avons, ensuite, présenté l'efficacité des implémentations sur des bases de données largement connues de la communauté scientifique, pour ne garder qu'une méthode duale d'ensemble actif qui, comparable aux deux méthodes de point intérieur, présente les meilleures performances.

Dans un deuxième temps, nous avons abordé les techniques floues pour la construction de systèmes d'inférence à partir de la connaissance, ceci afin de profiter au maximum de toute l'expertise acquise pour un problème donné. Dans le cas de la surveillance de l'opérateur humain, la connaissance sur une activité, comme celle du conducteur automobile, est vaste et très utile pour former des règles et inférer une décision pour une entrée donnée.

Dans un troisième temps, nous avons abordé les méthodes hybrides, englobant des systèmes où l'on dispose d'observations et d'expertise, à travers deux grandes schémas : les systèmes d'inférence flous conventionnels dotés d'un mécanisme d'apprentissage et les modèles flous représentés sous une architecture en forme de réseau de neurones. Sur ce dernier point, nous avons profité de nos travaux sur les SVM pour la construction de groupes ou clusters représentant des règles d'un système d'inférence flou. Nous avons démontré la pertinence de ce schéma avec son application à l'identification de modèles flous de type Takagi-Sugeno pour des systèmes MIMO. Nous avons comparé cette approche avec deux techniques classiques : les c-moyennes floues et l'algorithme de Gustafson-Kessel, montrant ainsi la supériorité du système hybride SVM-flou sur ces deux approches.

Le *quatrième chapitre* a présenté la conception d'un système de surveillance appliqué à l'assistance à la conduite automobile. Nous avons commencé par une synthèse des approches menées pour le développement de tels systèmes d'assistance. Suite à quoi, nous avons présenté un rappel sur les systèmes

d'assistance au conducteur, pour avertir de possibles situations de danger et la détection de fatigue. Pour cela, nous avons présenté les objectifs des projets AWAKE et PREDIT dont l'un des objectifs est la conception d'un système de détection temps-réel de l'hypovigilance. Après une analyse des deux architectures, nous nous avons limité notre démarche au développement de deux modules principaux : un module de diagnostic de situations à risque et un module de diagnostic de l'hypovigilance

Le développement du module de diagnostic de situations à risque nous a amené à proposer trois stratégies d'implémentation, utilisant différentes variables d'entrée. La comparaison de ces approches, à travers les bases de données des essais, a mis en évidence la pertinence des mesures tel que le temps de réaction, transposé à notre problème comme le temps mis par le conducteur à entreprendre une action corrective. Les systèmes d'inférence floue ont montré leur aptitude pour représenter des heuristiques dans un schéma de surveillance des situations à risque temps-réel.

Nous avons présenté le développement et la mise en œuvre du module de diagnostic de l'hypovigilance à partir de mesures de performance de la conduite. Nous avons présenté une architecture particulière basée sur une phase de prétraitement et une phase d'adaptation/décision. La phase de prétraitement calcule une transformée en ondelettes rapide des signaux, permettant la définition des intervalles d'analyse où le signal présente une dynamique stable. Ainsi, l'extraction des caractéristiques statistiques et fréquentielles de la conduite s'avère plus adéquat, évitant le mélange des informations caractéristiques d'une conduite normale et une anormale. Les machines à vecteurs de support pour la construction d'une frontière de décision ont montré leur pertinence pour la fusion de ces informations et pour en déduire une mesure du degré d'anormalité de la conduite. Nous avons évalué les performances de ce système sur des bases de données des essais pour estimer les taux de fausses alarmes et de non détections, à partir d'un ensemble de mesures de référence physiologiques et subjectives.

Perspectives

L'application de ces techniques de prétraitement de signaux et de fusion de l'information à la surveillance des systèmes en général reste très ouverte, puisqu'il faut satisfaire des contraintes intrinsèques aux systèmes eux-mêmes (dynamiques non linéaires, stochastiques, etc.), de temps-réel et économiques.

Pour le cas de la surveillance du conducteur automobile il faut tenir compte de plusieurs aspects :

Dans la partie de perception, les capteurs actuels doivent être améliorés afin de mieux comprendre l'environnement autour du véhicule, ses dynamiques et le scénario à l'intérieur de l'habitacle. Dans le cas du capteur de position latérale, il doit être complété par d'autres technologies, tels que les radars ou le GPS différentiel, afin d'obtenir des informations sur la route (nombre de voies, largeur de la voie, courbe), géographiques (montagne, montée, descente), ou sur l'état du trafic (densité véhiculaire, travaux).

Dans la partie prétraitement, l'analyse temporelle et fréquentielle donnée par les ondelettes offre encore plusieurs possibilités de recherche. Par exemple, l'évolution des coefficients en ondelettes, associés à chacun des intervalles en fréquence, peut donner des informations utiles sur la transition entre ces intervalles. Ces transitions peuvent être liées à un changement important du comportement du processus.

Dans la partie fusion, la recherche sur les machines d'apprentissage statistique et les systèmes d'inférence flous reste un axe très vaste qu'il faut explorer tant au niveau théorique qu'au niveau de la mise en œuvre.

Finalement, nous avons constaté le besoin d'implémenter un tel système de diagnostic sur des véhicules, pour sa validation par des expérimentations à grande échelle, demandant une intégration système pour une production à niveau semi-industriel (étape de prototype).

Bibliographie

- [1] ABONYI, J., BABUSKA, R. and SZEIFERT, F. (1998). *Modified Gath-Geva Fuzzy Clustering for identification of Takagi-Sugeno Fuzzy Models*. Delft University of Technology, Research Report.
- [2] AIZERMAN, M. A. et al. (1964). *The problem of pattern recognition learning and the method of potential functions*. Autom. Remote control 25, pp. 1175-1193.
- [3] ALEXANDER, L. and DONATH, M. (1999). *Differential GPS Based Control of Heavy Vehicles*. Final Report. Department of Mechanical Engineering University of Minnesota, January.
- [4] ALLOUM, A. (1994). *Modélisation et commande dynamique d'une automobile pour la sécurité de conduite*. Thèse de Doctorat, Université Technologique de Compiègne, UTC.
- [5] ALTMAN, A. and GONDZIO J. (1998). *Regularized Symmetric Indefinite Systems in Interior Point Methods for Linear and Quadratic Optimization*. Optimization Methods and Software 11-12, pp. 275-302
- [6] ANDERSEN, E. D., et al. (1996). *Implementation of interior-point methods for large-scale linear programming*. Technical Report 1996.3, Logilab, HEC Geneva, Section of Management Studies, University of Geneva, Switzerland.
- [7] ANUND, A. and PETERS, B. (2002). *Deliverable 7.1. Preliminary Pilot Plans*. AWAKE Project IST-2000-28062, September.
- [8] ARCOS (2003). *ARCOS 2003 Note de présentation*. Action de Recherche Conduite Sécurisée, Action fédérative PREDIT, édition 201200.
- [9] AssitWare Technologies, Inc (2003). *Interim Report on Road Departure Crash Warning Subsystems*. University of Michigan Transportation Research Institute, Visteon Corporation. September.
- [10] ASTROM, K. et al. (2001). *Control of Complex Systems*. Springer-Verlag. London, Great Britain.
- [11] AZEVEDO-FLHO, A. and SHACHTER, R. D. (1994). *Laplace's Method Approximations for Probabilistic Inference in Belief Networks with Continuous Variables*. In Lopes de Mantara, R and Poole, d (eds). *Uncertainty in Artificial Intelligence*, Morgan-Kaufman, San Francisco.
- [12] BABUSKA, R. and VERBRUGGEN, H. B. (1994). *Applied fuzzy modelling*. In proceedings IFAC Symposium on Artificial Intelligence in Real Time Control, Valencia, Spain, pp. 61-66.
- [13] BABUSKA, R. et al. (1998). *Identification of MIMO systems by input-output TS fuzzy models*. IEEE International Conference on Fuzzy Systems, Anchorage, Alaska, pp. 657-662.
- [14] BABUSKA, R. (2001). *Fuzzy and neural control: DISC course lecture notes*. Delft University of Technology, Control engineering laboratory, Delft, The Netherlands.
- [15] BATAVIA, P. (1999). *Driver-adaptive lane departure warning systems*. Ph.D. dissertation, Robotics Institute, Carnegie-Mellon Univ., Pittsburgh, PA, 1999.
- [16] BAUDAT, G. and ANOUAR, F. (2000). *Generalized discriminant analysis using a kernel approach*. Neural Computation, 12 pp. 2385-2404.
- [17] BEKIARIS, E. et al. (1999). *SAVE project Final Report*. SAVE project TR1047.
- [18] BEKIARIS, E. et al. (2002). *Deliverable 1.1: User needs analysis per category of drivers group*. AWAKE Project IST-2000-28062. May.
- [19] BELLET, T. et al (2002). *"Real time" analysis of the driving situation in order to manage on-board information*. Congrès e- SAFETY, septembre 16-18, Lyon, France, 8 p.

- [20] BELZ, S. M. (2000). *An on-road investigation of self-rating of alertness and temporal separations as indicators of driver fatigue in commercial motor vehicle operators*. PhD Dissertation, Faculty of Virginia Polytechnic Institute and State University. October, 283 p.
- [21] BEN YOUSSEF, C. (1997). *Benchmark of a waste-water treatment process*. LAAS Rep. 97458. Contrat FAMIMO, November.
- [22] BERNUSSOU, J. and TITLI, A. (1982). *Interconnected Dynamical systems: Stability, Decomposition and Decentralisation*. North-Holland Publishing Company. Amsterdam, Netherlands.
- [23] BERSCHANDY, D. (1993). *Comparaison des approches classiques dans l'architecture d'un système d'intelligence artificielle embarqué. Application à la détection temps réel de danger automobile*. Thèse, Université de Paris-Sud (Paris XI) Centre d'Orsay.
- [24] BEZDEK, J. C. (1981). *Pattern Recognition with Fuzzy Objective Function Algorithms*. Plenum Press, New York.
- [25] BILMES, J. A. (1998). *A Gentle Tutorial of the EM Algorithm and its Application to Parameter Estimation for Gaussian Mixture and Hidden Markov Models*. Technical Report TR-97-021, International Computer Science Institute, Berkeley CA, April.
- [26] BISHOP, C. M. and BISHOP, C. (1996). *Neural Networks for Pattern Recognition*. Oxford University Press.
- [27] BISHOP, C. M. and TIPPING, M. E. (2000). *Variational Relevance Vector Machines*. In Proceedings of the 16th Conf. in Uncertainty in Artificial Intelligence, Morgan Kaufmann Publishers, pp. 46-53.
- [28] BITTNER, R., et al. (2000). *Detecting of fatigue states of a car driver*. In Medical data analysis, Lecture notes in computer sci. Brause R. W., Hanisch E. eds. Springer, Frankfurt.
- [29] BOVERIE, S. et al. (2003). *Facteurs Biologiques et Environnementaux de la Dégradation de la Vigilance chez les conducteurs et Système de Détection Automatique des Baisses de Vigilance*. Programme PREDIT III, Conventions de recherche n°02 A 0033-38.
- [30] BOVERIE, S. et al. (2004). *Deliverable 3.2: Optimized HDM for cars and heavy vehicles*. AWAKE N° IST-2000-28062, April 2004, 38p.
- [31] BOSER, B. E., GUYON, I. M. and VAPNIK, V. N. (1992). A training algorithm for optimal margin classifiers. In D. HAUSSLER, editor, Proceedings of the 5th Annual ACM Workshop on Computational Learning Theory. Pittsburgh, USA, July. pp 144-152.
- [32] BOY, G. (1991). *Intelligent Assistant Systems*. Knowledge-Based systems Vol. 6. Academic press. San Diego, California.
- [33] BROWN, M. and HARRIS, C. J. (1994). *Neuro-fuzzy Adaptive Modelling and Control*. Prentice Hall, Series in Systems and Control Engineering, Hemel Hempstead, UK.
- [34] BURGESS, C. J. C. (1998). *A tutorial on Support Vector Machines for Pattern Recognition*. Data Mining and Knowledge Discovery, 2(2), pp. 121-167.
- [35] BURROUGH, P.A., VAN GAANS, P.F.M. and MACMILLAN, R.A., (2000). *High-resolution landform classification using fuzzy k-means*. Fuzzy Sets and Systems, 113(1), pp. 37-52.
- [36] CASTRO, J. L. et al. (2002). *Interpretation of artificial neural networks by means of fuzzy rules*. IEEE Transactions on Neural Networks, 13(1), pp. 101-116.
- [37] COURANT, R. and HILBERT, D. (1953). *Methods of Mathematical Physics*. John Wiley, New York, USA.
- [38] CHEN, M.-S. and WANG, S.-W (1999). *Fuzzy clustering analysis for optimizing fuzzy membership functions*. Fuzzy Sets and Systems, 103(2), pp. 239-254.
- [39] CHIU, S. L. (1994). *Fuzzy model identification based on cluster estimation*. Journal of Intelligent Fuzzy Systems, vol. 2, pp. 267-278.
- [40] De VISSER, W., MARCHAU, V. A.W. J. and Van Der HEIJDEN, R.E.C.M. (1999). *The cost-effectiveness of future driver support systems*. D. Roller (ed.) 32nd International Symposium on Automotive Technology & Automation, Automotive Automation Ltd, Croydon, pp. 421-428.
- [41] De WAARD, D. (1996). *The measurement of drivers' mental workload*. PhD thesis, University of Groningen. Haren, The Netherlands: University of Groningen, Traffic Research Centre.
- [42] DENISON, D., HOLEMS, Chris, MALLICK, B. and SMITH, A. (2002). *Bayesian Methods for Nonlinear Classification and Regression*. John Wiley & Sons, July.

- [43] DiDOMENICO, A. T. (2003). *An investigation on subjective assessments of workload and postural stability under conditions of joint mental and physical demands*. PhD Thesis of the Faculty of the Virginia Polytechnic Institute and State University, Blacksburg, Virginia, July.
- [44] DILLIES-PELTIER, M.-A. and Le FORT-PIAT, N. (2002). *Les systèmes d'aide à la conduite*. Dans La voiture intelligente (Traité IC2, série Systèmes automatisés), GISSINGER Gérard, LE FORT-PIAT Nadine éditeurs, Lavoisier, pp.135-185.
- [45] DINGES, D. F., MALLIS, M. M., MAISLIN, G. and POWELL, J. W. (1998). *Evaluation of techniques for ocular measurement as an index of fatigue and as the basis for alertness management*. Final Report. National Highway Traffic Safety Administration. Report No DOT HS 808 762.
- [46] DINGUS, T. A., HARDEE, L. H. and WIERWILLE, W. W. (1985). *Detection of Drowsy and Intoxicated Drivers Based on Highway Driving Performance Measures*. IEOR Department report No. 8402. Vehicle Simulation Laboratory, Human Factors Group, Virginia Polytechnic Institute and State University, Blacksburg, Virginia.
- [47] DOWDY, S. M. and WEARDEN, S. (1991). *Statistics for Research*. Wiley, New York.
- [48] DUDA, R. O., HART, P. E. and STORK D. G. (2000). *Pattern Classification*. Wiley-Interscience; 2nd edition.
- [49] DUNN, J.C. (1974). *A fuzzy relative of the ISODATA process and its use in detecting compact well-separated clusters*. J. Cybernetics. 3 (3): 32-57, 1974.
- [50] DURSO, F. T., and GRONLUND, S. D. (1999). *Situation awareness*. In F. T. Durso (Ed.), *Handbook of Applied Cognition*. NY: Wiley & Sons. pp. 283-314.
- [51] EMAMI, M. R., TURSKEN, I. B. and GOLDENBERG, A. A. (1998). *Development of A Systematic Methodology of Fuzzy Logic Modeling*. IEEE Transactions on Fuzzy Systems, 6(3), pp. 346-361.
- [52] EUROCONTROL (2002). *Activity 2. Validation framework WP3 Deliverable 5. Human Performance Metrics*. Reference: CARE/ASAS/NLR/02-034, November.
- [53] EUROCONTROL (2003). *Review of workload measurement, analysis and interpretation methods*. Reference: CARE-Integra-TRS-130-02WP2, March.
- [54] European Commission (2001). *WHITE PAPER – European transport policy for 2010: time to decide*. Office for Official Publications of the European Communities. http://europa.eu.int/comm/energy_transport/en/lb_en.html.
- [55] FAIRCLOUGH, S., PLANQUE, S., MARTINEZ, D. and BROOKHUIS, K. A. (1994). *Behavioural responses to a driver impairment monitoring system*. 1st World Congress on Advanced Transport Telematics, Paris (France), November 30–December 3, 8p.
- [56] FLETCHER, Ralph (2000). *Practical Methods of Optimization*. John Wiley & Sons; 2nd edition, May.
- [57] FOULLOY, L., BONOIT E. and MAURIS G. (1993). *Applications of fuzzy sensors*. European Workshop on Industrial Fuzzy Control and Applications. Terrassa, España.
- [58] FREUND, Robert M. and MIZUNO, Shinji (1996). *Interior Point Methods: Current Status and Future Directions*. Optima, Vol. 51, pp. 1-9.
- [59] GATH, I. and GEVA, A. B. (1989). *Unsupervised optimal fuzzy clustering*. IEEE Transactions on Pattern Analysis and Machine Intelligence 7, pp. 773-781.
- [60] GERTZ, Michael and WRIGHT, Stephen J. (2001). *Object-Oriented Software for Quadratic Programming*. Report ANL/MCS-P891-1000, Argonne National Laboratory, Mathematics and Computer Science Division.
- [61] GIRALT, A., JAMMES, B., GONZALEZ-MENDOZA, M. and THOMAS, J. (2003). *A7.2 - Sensors and subsystem technical verification: HDM*. Rapport LAAS No. 04174 (Confidential). Contrat AWAKE N° IST-2000-28062, April 2004, 38p.
- [62] GIRALT, A. et al. (2003). *Deliverable 3.1: Driver hypovigilance criteria, filter and HDM module*. AWAKE Project IST-2000-28062. October 2003, 109p.
- [63] GODTHELP, H., MILGRAM, P. and BLAAUW, G. (1984). *The development of a time related measure to describe driving strategy*. Human Factors, vol. 26, no. 3, pp. 257-268.
- [64] GOLDFARB, D. and IDNANI, A. (1983). *A numerically stable dual method for solving strictly convex quadratic programs*. Mathematical Programming, 27, pp. 1-33.
- [65] GONZALEZ-MENDOZA, Miguel (2000). *Etude du problème d'optimisation dans les Machines à Vecteurs de Support*. Mémoire de DEA, INSAT-LAAS-CNRS, France.

- [66] GONZALEZ-MENDOZA, M. (2002). *Système de Diagnostic par des Machines à Vecteurs de Support : Application à la Détection de l'Hypovigilance du Conducteur Automobile*. 3ème Congrès de Doctorants. Ecole Doctorale Systèmes LAAS-CNRS, May 22-23 2002. Blagnac, France.
- [67] GONZALEZ-MENDOZA, M., TITLI A., HERNADEZ-GRESS, N. (2002). *Boosting Support Vector Machines in Density Estimation Problems*. IEEE Information Processing and Management of Uncertainty, IPMU 2002 Symposium, July 1-5 2002. Annecy, France.
- [68] GONZALEZ-MENDOZA, M., SANTANA-DIAZ, A., TITLI, A., HERNADEZ-GRESS, N. (2002). *Driver Vigilance Monitoring, a New Approach*. IEEE Intelligent Vehicles IV'2002 Symposium, June 20-24 2002. Versailles, France.
- [69] GONZALEZ-MENDOZA, M., SANTANA-DIAZ, A., JAMMES, B., TITLI, A., ESTEVE, D., HERNADEZ-GRESS, N. (2002). *Comparison between Classification and PDF Estimation SVMs for Driver's Vigilance Monitoring*. IBERAMIA 2002 Workshop Minería de datos, November 11-14 2002. Sevilla, Spain. September.
- [70] GONZALEZ-MENDOZA, Miguel, TITLI, André and HERNADEZ-GRESS, Neil (2003). *SVM clustering for identification of Takagi-Sugeno fuzzy models*. 5th IFAC International symposium of Intelligent Components and Instruments for Control Applications, SICICA. July 9-11, Aveiro, Portugal.
- [71] GONZALEZ-MENDOZA, M., JAMMES, B., HERNADEZ-GRESS, N. TITLI, A. and ESTEVE, D. *A Comparison of Road Departure Warning Systems on Real Driving Conditions*. Submitted to IEEE 7th Intelligent Transportation Systems ITS 2004, 3-6 Octobre 2004. Washington, D.C.
- [72] GONZALEZ-MENDOZA, M., TITLI, A. and HERNADEZ-GRESS, N. (2004). *A study of the QP optimization problem in the learning phase of SVM*. IBERAMIA 2004, Puebla, Mexico (submitted).
- [73] GUILLAUME, Serge (2001). *Designing fuzzy inference systems from data: an interpretability-oriented review*. IEEE Transactions on Fuzzy Systems, 9(3), 426-443, June.
- [74] GUSTAFSON, D.E. & KESSEL, W.C. (1979). *Fuzzy clustering with a fuzzy covariance matrix*. Proceedings IEEE CDC. San Diego, CA, USA. 761-766.
- [75] HARRIS, Chris et al. (2002). *Adaptive Modelling, Estimation and Fusion from Data: A Neurofuzzy Approach*. Springer. 2002
- [76] HAYKIN, S. (1999). *Neural Networks: A comprehensive foundation*. 2nd edition. Prentice-Hall, Englewood Cliffs, New Jersey, USA.
- [77] HERNANDEZ, Neil (1998). *Système de Diagnostic par Réseaux de Neurones et Statistiques: Application à la détection d'hypovigilance du conducteur automobile*. Thèse LAAS-ENSEEIH, Toulouse France, Décembre. LAAS Report 98571.
- [78] HERRERA, Francisco, LOZANO, Manuel, VERDEGAY Jose Luis (1994). *Generating fuzzy rules from examples using genetic algorithms*. 4th International Conference on Information Processing and Management of Uncertainty in Knowledge-Bases System, IPMU'94. Paris (France), pp. 675-680.
- [79] HERRERA-CORRAL, J. A. (1995). *Approche Multisensorielle pour la Détection Comportementale de la baisse de la vigilance d'un conducteur automobile*. Doctorat, Institut National des Sciences Appliquées, Toulouse, France. 27 Octobre, 146p. Toulouse. LAAS Report 95358.
- [80] HETTICH, S. and BAY, S. D. (1999). *The UCI KDD Archive*. url: <http://kdd.ics.uci.edu>. Irvine, CA: University of California, Department of Information and Computer Science.
- [81] [53] HLAWATSCH, F. and FLANDRIN, P. (1997). *The Interference Structure of the Wigner Distribution and Related Time-Frequency Signal Representations*. W. Mecklenbrauker and F. Hlawatsch, eds. in : *The Wigner Distribution - Theory and Applications in Signal Processing*, Elsevier, Amsterdam, pp. 59-133.
- [82] HOLSCHNEIDER, M., et al (1989). *A real-time algorithm for signal analysis with the help of the wavelet transform*. In J. M. Combes, A. Grossmann, and Ph. Tchamitchian, editors, *Wavelets, Time-Frequency Methods and Phase Space*, Springer-Verlag. pp. 286-297.
- [83] ICHIHASHI, H., et al. (1996). *Neuro-fuzzy ID3: a method of inducing fuzzy decision trees with linear programming for maximizing entropy and an algebraic method for incremental learning*. Fuzzy Sets and Systems 81, pp. 157-167.
- [84] Intelligent Transportation Society of America and the United States Department of

- Transportation (2002). *National Intelligent Transportation Systems Program Plan: A Ten-Year Vision*, January.
- [85] ISHIBUCHI, H., et al. (1994). *Empirical study on learning in fuzzy systems by rice test analysis*. Fuzzy Sets and Systems, 64, 129-144.
 - [86] JAGER, René (1995). *Fuzzy Logic in Control*. Ph.D. thesis Delft University of Technology, Department of Electrical Engineering, Control laboratory, Delft, The Netherlands.
 - [87] JANG, Jyh-Sing Roger (1993). *ANFIS: Adaptive-network-based fuzzy inference systems*. IEEE Transactions on System, Man and Cybernetics. 23(3), pp. 665-685.
 - [88] JENG, Jin-Tsong, LEE, Tsu-Tain (1999). *Support vector machines for the fuzzy neural networks.*, 1999. Conference Proceedings. 1999 IEEE International Conference on Systems, Man, and Cybernetics, IEEE SMC '99, 12-15 Oct. 1999. pp. 115-120.
 - [89] JIN, Yaochu (2000). *Fuzzy modeling of high-dimensional systems: Complexity reduction and interpretability improvement*. IEEE Transactions on Fuzzy Systems, 8(2), pp. 212-220.
 - [90] JIN, Yaochu (2003). *Advanced Fuzzy Systems Design and Applications*. Series: Studies in Fuzziness and Soft Computing, Vol. 112, Springer-Verlag, Germany
 - [91] JOACHIMS, Thorsten (1998). Making large-scale support vector machine learning practical. In *Advances in Kernel Methods: Support Vector Machines*. B. Schölkopf, C. Burges, A. Smola editors, MIT press, Cambridge, MA, pp. 169-184.
 - [92] JOUFFE, Lionel (1998). *Fuzzy Inference System Learning by Reinforcement Methods*. IEEE Trans. on Systems Man & Cybernetics, PART C: Applications & Reviews, Vol. 28, No. 3, Aug. pp. 338-355.
 - [93] KOWALCZYK, Adam. *Maximal Margin Perceptron*. In A.J. Smola, P.L. Bartlett, B. Schölkopf, and D. Schuurmans, editors, *Advances in Large Margin Classifiers*, pp. 75-114, Cambridge, MA, 2000. MIT Press.
 - [94] KRAMER, A.F. (1991). *Physiological metrics of mental workload: a review of recent progress*. In D.L. Damos (Ed.), *Multiple task performance*. London: Taylor & Francis.
 - [95] LANGARI, R., WANG, L. and YEN, J. (1997). Radial Basis Functions Networks, Expectation Maximisation Algorithm. IEEE Transactions on Systems, Man and Cybernetics, 27(5), pp. 613-623.
 - [96] LAPRIE, J. C. et al. (1995). *Guide de la Sûreté de Fonctionnement*. Cépaduès – Editions, Toulouse, France.
 - [97] LeCUN, R. et al. (1995). *Comparison of learning algorithms for handwritten digit recognition*. Proceedings ICANN'95 – International conference on artificial Neural Networks, volume II, pp53-60. Nanterre, France.
 - [98] MALLAT, Sthéphane (2000). *Une exploration des signaux en ondelettes*. Editions de l'école Polytechnique.
 - [99] MAMDANI, E. H. (1976). *Application of fuzzy logic to approximate reasoning using linguistic synthesis*. Proceedings of the sixth international symposium on Multiple-valued logic table of contents. Logan, Utah, United States, pp. 196-202.
 - [100] MICHON, J. A. (1985). *A critical view of driver behaviour models: what do we know what should we do?* L. Evans and R. C. Schwing Eds. Human behaviour and traffic safety. Plenum Press. New York.
 - [101] MILLEMANN, S. (2001). *Vigilance du conducteur. Eléments clés dans la vision du client du constructeur automobile*. Colloque Vigilance des conducteurs, pilotes et opérateurs: aspects médicaux, développements technologiques, amélioration de la sécurité. Toulouse, 15-16 novembre.
 - [102] MINKA, Thomas P. (2001). *A family of algorithms for approximate Bayesian inference*. PhD thesis of Department of Electrical Engineering and Computer Sciences Department, Massachusetts Institute of Technology.
 - [103] MINSKY, Marvin L. and PAPERT, Seymour A. (1988). *Perceptrons: An Introduction to Computational Geometry*. The MIT Press, Cambridge, MA, expanded edition.
 - [104] MORE, Jorge J. and WRIGHT, Stephen J. (1993). *Optimization Software Guide*. SIAM Publications.
 - [105] MUZET, A. (2001). *Conditions de vie et vigilance*. Colloque Vigilance des conducteurs, pilotes et opérateurs: aspects médicaux, développements technologiques, amélioration de la sécurité. Toulouse, 15-16 novembre.

- [106] NEAL, Radford M. (1993). *Probabilistic Inference using Markov Chain Monte Carlo Methods*. Technical report of the Department of computer Sciences. University of Toronto.
- [107] NICOLA, L., LAGOUTTE, A. et OJEDA, L. (2002). *Spécifications ergonomiques et conception de l'IHM d'un système de surveillance de trajectoire latérale*. Congrès e- SAFETY, septembre 16-18, Lyon, France, 9 p.
- [108] OLARU, Cristina, WEHENKEL, Louis (2003). *A complete fuzzy decision tree technique*. Fuzzy Sets and Systems, 138, pp. 221-254.
- [109] OLIVE, X. (2003). *Approche intégrée à base de modèles pour le diagnostic hors ligne et la conception : application au domaine de l'automobile*. Doctorat de l'Université Paul Sabatier. Toulouse, Décembre. Rapport LAAS No03611.
- [110] OPPENHEIM, Alan V., SCHAFER, Ronald W. and BUCK, John R. (1989). *Discrete-Time Signal Processing*. Prentice-Hall, Englewood Cliffs, NJ.
- [111] OSUNA, Edgar E. et all (1997). *Support Vector Machines: Training and Applications*. MIT Artificial Intelligence Laboratory, March.
- [112] PAPOULIS, A. and PILLAI, S. U. (2002). *Probability, random variables and stochastic processes*. Mc Graw Hill.
- [113] PARK, Minkee, et al. (1999). *A new approach to the identification of a fuzzy model*. Fuzzy Sets and Systems, Vol. 104, Issue 2 pp. 169-181.
- [114] PARKES, A. M. (1991). *Data Capture techniques for RTI usability evaluation*. Dans Commission of the European Communities. Drive Conference. Amsterdam.
- [115] PEEBLES, P. Z. Jr. (2001). *Probability, Random Variables and Random Signal Principles*. 4th Edition, McGraw Hill, New York.
- [116] PILUTTI, T. and ULSOY, A. G. (1999). *Identification of driver state for lane keeping tasks*. IEEE Transactions on Systems Man and Cybernetics–Part A: Systems and Humans, vol. 29(5), September 1999, pp. 486-502.
- [117] PILUTTI, T. and ULSOY, A. G. (2003). *Fuzzy-logic-based virtual rumble strip for road departure warning systems*. IEEE Transactions on Intelligent Transportation Systems, vol 4(1), March 2003, pp. 1-12.
- [118] PLATT, J. (1998). *Sequential Minimal Optimization: A Fast Algorithm for Training Support Vector Machines*. In *Advances in Kernel Methods: Support Vector Machines*. B. Schölkopf, C. Burges, A. Smola editors, MIT press, Cambridge, MA, pp.185-208.
- [119] POPIEUL, J. C. (2001). *Etude des modifications du comportement du conducteur automobile au cours d'un trajet autoroutier de longue durée*. Colloque Vigilance des conducteurs, pilotes et opérateurs: aspects médicaux, développements technologiques, amélioration de la sécurité. Toulouse, 15-16 novembre.
- [120] POWELL, M. J. D. (1985). *ZQPCVX A fortran subroutine for convex quadratic programming*. Technical Report DAMTP/83/NA17, University of Cambridge, UK.
- [121] POWELL, M. J. D. (1987). *Radial basis functions for multivariate interpolation: A review*. In J. C. Mason and M. G. Cox editors. *Algorithms for Approximation of Functions and Data*. Oxford University Press.
- [122] PREECE, J., ROGERS, Y. and SHARP, H. (2002). *Interaction Design: Beyond Human-Computer Interaction*. New York, NY: John Wiley & Sons.
- [123] PRINZEL, Lawrence J., et al. (2001). *Empirical Analysis of EEG and ERPs for Psychophysiological Adaptive Task Allocation*. The NASA STI Program Office, NASA Center for AeroSpace Information, TM-2001-211016, June.
- [124] RAUDYS, Sarunas (2001). *Statistical and Neural Classifiers: An Integrated Approach to Design (Advances in Pattern Recognition)*. Springer-Verlag UK.
- [125] REASON, James T. (1990). *Human error*. Cambridge University Press.
- [126] REECE, Peter (1987). *Perceptrons and Neural Nets*. AI Expert, Volume 2.
- [127] REDMILL, K. A. and OZGUNER, U. (1998). *The Ohio State University Automated Highway System Demonstration Vehicle*. 1998 SAE International Congress and Exposition, vol. 1332, February, pp 117-125.
- [128] RIPLEY, Bryan D. (1996). *Pattern Recognition and Neural Networks*. Cambridge University Press.
- [129] ROGE, J., PEBAYLE, T., and MUZET A. (2001). *Variations of the level of vigilance and of behavioral activities during simulated automobile driving*. Accident Analysis and Prevention, 33. pp. 181-186.

- [130] ROSS, T. J. (1995). *Fuzzy Logic with Engineering Applications*. McGraw-Hill Inc, USA.
- [131] RUMELHART, D. E., HINTON, G. E. and WILLIAMS R. J. (1986). *Learning internal representations by error propagation*. D. E. Rumelhart and J. L. McClelland (Eds). *Parallel Distributed processing: Explorations in the Microstructure of cognition*, volume I. MIT Press.
- [132] RUSSELL, S. et NORVIG, P. (2003). *Artificial Intelligence. A Modern Approach*. Pearson Education Inc. Prentice Hall, New Jersey. Second Edition.
- [133] SANTANA-DIAZ, A., JAMMES, B., GONZALEZ-MENDOZA, M. and D.ESTEVE. *Driver hypovigilance diagnosis using wavelets and statistical learning*. IEEE 5th Intelligent Transportation Systems ITS 2002, Singapore. September 2002.
- [134] SANTANA-DIAZ, A. (2003). *Conception d'un système de détection de la baisse de vigilance du conducteur automobile par l'utilisation des ondelettes et l'apprentissage statistique*. Doctorat, Université Paul Sabatier, Toulouse, France. 6 Janvier 2003, Rapport LAAS No03051. 192p.
- [135] SCHOLKOPF, B. (1997). *Support Vector Learning*. PhD Thesis of Technical University of Berlin, Germany. September.
- [136] SCHOLKOPF, B. et al. (1999). *Estimating the Support of a High-Dimensional Distribution*. Technical Report Microsoft Research. 27 November.
- [137] SCHUERMANN, J. (1996) *Pattern Classification: A unified view of statistical and neural approaches*. Wiley-Interscience, New York.
- [138] SCOTT, D. W. (1992). *Multivariate Density Estimation*. Wiley Inter-Science, New York.
- [139] SHEPHERD, A. J. (1997). *Second-Order Methods for Neural Networks: Fast and Reliable Training Methods for Multi-Layer Perceptrons*. Perspectives in Neural Computing Series, Springer Verlag, June.
- [140] SHLADOVER, S. et al. (1991). *Automated vehicle control developments in the PATH program*. IEEE Transactions on Veh. Technol., vol. 40, pp. 114-120.
- [141] SILVERMAN, Bernard Walter (1986). *Density Estimation for Statistics and Data Analysis*. Chapman & All / CRC Press.
- [142] SMOLA, A. J. (1998). *Learning with Kernels*. PhD thesis, Technische Universität Berlin.
- [143] SMOLA, A. J. et SCHOLKOPF, B. (1998). *A tutorial on support vector regression*. Statistics and Computing. Invited paper.
- [144] SUGENO M. and YASUKAWA T. (1993). *A fuzzy-logic-based approach to qualitative modeling*. IEEE transactions on Fuzzy Systems, 1(1), pp. 7-31.
- [145] TAKAGI, T. and SUGENO, M. (1985). *Fuzzy identification of systems and its application to modelling and control*. IEEE Transactions on Systems, Man and Cybernetics 15(1), pp. 116-132.
- [146] TIJERINA, L. et al. (1999). *A Preliminary of Algorithms for Drowsy and Inattentive Driver Detection on the Road*. National Highway Traffic Safety Administration. Report No DOT HS 808 (TBD).
- [147] TIKHONOV, A. T. and ARSENIN, V. Y (1977). *Solutions of ill-posed problems*. W. H. Winston, Washington D. C.
- [148] TITLI, A. (2001). *Logique floue et systèmes flous*. Notes de cours de l'Institut National des Sciences Appliquées, Toulouse, France.
- [149] TRAVE-MASSUYES, L. (1998). *Raisonnement qualitatif et automatique symbolique*. Rapport LAAS No98589. Habilitation, Université Paul Sabatier, Toulouse.
- [150] TRICOT, N., SONNERAT, D., POPIEUL, J. C. (2002). *Driving Styles and Traffic Density Diagnosis in Simulated driving Conditions*. IEEE Intelligent Vehicles IV'2002 Symposium, June 20-24 2002. Versailles, France.
- [151] VANDERBEI, Robert J. (1997). *Linear programming: Foundations and Extensions*. Kluwer academic Publishers, Hingham, MA.
- [152] VANDERBEI, Robert J. (1999). *LOQO: An interior point code for quadratic programming*. Optimization Methods and Software, Vol. 11. pp. 451-484.
- [153] VAPNIK, V. N. (1977). *Estimating of values of regression at the point of interest*. In Method of Pattern Recognition. Sotvetskoe radio (in Russian).
- [154] VAPNIK, V. N. (1982). *Estimation of Dependences Based on Empirical Data*. Springer-Verlag. New York Inc.

- [155] VAPNIK, V. N. (1995). *The Nature of Statistical Learning Theory*. Springer-Verlag. New York Inc.
- [156] VAPNIK, V. N. (1998). *Computational Learning Theory*. John Wiley & Sons. New York.
- [157] VAPNIK, V. N. (2000). *Support Vector Method for Multivariate Density Estimation*. In Sara Solla, Todd Leen, Klaus-Robert Muller (eds.), *Advances in Neural Information Processing Systems*, MIT Press 25, pp. 659-665.
- [158] VERWEY, W. and ZAIDEL, D. (1999). *Predicting drowsiness accidents from personal attributes, eye blinks, and ongoing driving behaviour*. *Personality and Individual Differences* No: 28, pp. 123-142.
- [159] VOELCKER, J. (2004). *Top 10 Tech cars*. *IEEE Spectrum*, March, pp. 20-27.
- [160] WAHBA, Grace (1990), *Splines Models for Observational Data*. Series in Applied Mathematics, Vol. 59, SIAM, Philadelphia, USA.
- [161] WANG, L. X. and MENDEL, J. M. (1992). *Generating fuzzy rules by learning from examples*. *IEEE Transactions on Systems, Man, and Cybernetics*, 22, pp. 1414-1427.
- [162] WARWICK, K. (1995). *New ideas in fuzzy clustering and fuzzy automata*. Edited by Steele, N. C. *Proceedings of the international ICSC Symposium on Fuzzy logic*. Canada/Switzerland. ICSC Academic Press, pp. XXI-XXVII.
- [163] WESTON, Jason A. E. (1999). *Extensions to the Support Vector Method*. PhD Thesis of Royal Holloway University of London, England. October.
- [164] WIERWILLE, W. W. (1999). *Historical perspective on slow eyelid closure: Whence PERCLOS?* Technical Conference on Ocular Measures of Driver Alertness, Herndon.
- [165] WINSUM, V., BROOHHUIS, K.A. and De WAARD, D. (2000). *A comparison of different ways to approximate time to line crossing (TLC) during car driving*. *Accident Analysis and prevention* (32) 47-56, 2000.
- [166] WOODS, David D. (1996). *Decomposing automation: Apparent simplicity, real complexity*. In R.Parasuraman & M.Mouloua (Eds.), *Automation and human performance: Theory and applications*. Mahwah, NJ: Lawrence Erlbaum Assoc, pp. 3-18.
- [167] WYLIE, C. D., et al. (1996). *Commercial Motor Vehicle Driver Fatigue and Alertness Study: Technical Summary*. Essex Corporation, TP report number 12876E, November.
- [168] XIE, X.L. and BENI, G. (1991). *A Validity measure for Fuzzy Clustering*. *IEEE Transactions on Pattern Analysis and machine Intelligence*, Vol. 13, No4, August pp. 841-847.
- [169] YAOCHU, Jin (2003). *Advanced Fuzzy Systems Design and Applications*. Series: Studies in Fuzziness and Soft Computing, Vol. 112, Springer-Verlag, Germany.
- [170] ZADEH, L. A. (1965). *Fuzzy sets*. *Information and Control*, volume 8, pp. 338-353.
- [171] ZHANG, T. (2001). *Text categorization based on regularized linear classification methods*. *Information Retrieval*, volume 4, April, pp. 5-31.
- [172] ZWINGELSTEIN, G. (1995). *Diagnostic des Défaillances*. Editions Hermès, Paris.
- [173] *Ford Drowsiness study*. <http://www.4x4fords.com/news/drowsy.htm>