



Représentation et reconnaissance des signaux acoustiques sous-marins

Samir Ouelha

► To cite this version:

Samir Ouelha. Représentation et reconnaissance des signaux acoustiques sous-marins. Autre. Université de Toulon, 2014. Français. NNT : 2014TOUL0012 . tel-01136660

HAL Id: tel-01136660

<https://theses.hal.science/tel-01136660>

Submitted on 27 Mar 2015

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.

ÉCOLE DOCTORALE 548
IM2NP-Equipe Signaux Systèmes

THÈSE présentée par :

Samir OUELHA

soutenue le : **11 décembre 2014**

pour obtenir le grade de Docteur en
Sciences pour l'ingénieur : Mécanique, Physique, Micro et Nanoélectronique
Spécialité : Traitement du signal

Représentation et reconnaissance des signaux acoustiques sous-marins

THÈSE dirigée par :

M. COURMONTAGNE Philippe

Enseignant-Chercheur (HDR), ISEN-Toulon

JURY :

M. ADAM Olivier

M. GERVAISE Cédric

M. GARELLO René

M. JAUFFRET Claude

M. COURMONTAGNE Philippe

M. CHAILLAN Fabien

Professeur, Institut Jean le Rond d'Alembert

Chair chorus (HDR), GIPSA-Lab

Professeur, Télécom Bretagne

Professeur, Université du Sud Toulon Var

Enseignant-Chercheur (HDR), ISEN-Toulon

Docteur-Ingénieur, DCNS

Remerciements

Je tiens tout d'abord à remercier Olivier Adam et Cédric Gervaise qui ont utilisé toutes leurs connaissances afin de rapporter mon travail de thèse. De même, je remercie René Garelo et Claude Jauffret pour accepter de participer au jury de ma thèse.

Par la même occasion, je remercie Laurent Silhouette et Gilles Sague pour m'avoir accueilli dans le département de détection Sous-Marine (DSM), au sein du groupe Reconnaissance Acoustique (RAC), sans eux cette thèse n'aurait pas pu exister. Je remercie, de plus, tous les collègues du département DSM que j'ai eu l'occasion de rencontrer. Je remercie particulièrement Jean-Rémi Mesquida pour son aide constante durant ces trois années.

Je voudrais témoigner ma reconnaissance à Philippe Courmontagne, directeur de thèse, pour m'avoir donné goût au traitement du signal et m'avoir orienté dans la direction dans laquelle je suis, notre rencontre a été déterminante pour le déroulement de mes études. Philippe m'a fait confiance et m'a laissé de la liberté dans mon travail de recherche tout en m'incitant à me rattacher à des objectifs précis.

Je tiens à remercier Fabien Chaillan, encadrant industriel, pour m'avoir encadré parfaitement tout au long de ma thèse. Cette fin de thèse met fin à une collaboration de 4 ans de travail, dans une excellente ambiance, durant lesquelles j'ai progressé grâce aux conseils pertinents qu'il a pu me donner.

Je tiens à remercier, mes parents pour tout et ma famille pour leur soutien. Enfin je remercie ma femme Meriam qui a été à mes côtés tout au long de cette thèse.

Table des matières

Remerciements	3
Table des matières	4
Table des figures	8
Liste des tableaux.....	11
Introduction générale.....	12
Chapitre 1 : Chaîne de traitement des signaux sous-marins	14
I. Introduction.....	15
II. Acquisition des données.....	17
III. Formation de voies	18
IV. Veille panoramique	20
V. Extraction et poursuite sur veille panoramique	21
VI. Segmentation d'une représentation temps-fréquence	23
VII. Reconnaissance et identification des pavés temps-fréquence détectés	24
Chapitre 2: Représentations des signaux acoustiques non-stationnaires	26
I. Représentations temps-fréquence	27
A. Introduction	27
B. Transformée de Fourier à court terme (TFCT)	27
1. Principe	27
2. Avantages et limitations	28
3. Expérimentations et résultats.....	30
C. Transformée de Wigner-Ville et ses dérivées	33
1. Principe	33
2. Pseudo Wigner-Ville lissée.....	34
3. Expériences et interprétations.....	35
D. Réallocation spectrale.....	40
E. Transformée en ondelettes	42
1. Introduction	42
2. Définition et propriétés	42
3. Analyse multi-résolution.....	44
4. Avantages et inconvénients	44
II. Représentation basée sur l'audition humaine	45
A. Introduction	45

B.	L'oreille humaine.....	45
1.	Fonctionnement de l'oreille humaine.....	45
2.	Sensibilité de l'oreille	46
C.	Analyse par bandes d'octave	47
1.	Principe	47
2.	Analyse par tiers d'octave.....	47
D.	Une nouvelle technique de représentation de signaux acoustiques: l' <i>Hearingogram</i>	48
1.	Filtres de Mel	48
2.	Principe	51
3.	Formulation discrète de l' <i>Hearingogram</i>	52
4.	Comparaison entre <i>Hearingogram</i> et ondelettes	53
5.	Résultats.....	55
III.	Réduction du bruit des signaux non-stationnaires	60
A.	Introduction	60
B.	Etat de l'art.....	61
1.	Méthodes d'estimation de RSB <i>a priori</i>	61
2.	Règle d'atténuation.....	67
C.	Réduction du bruit au sein des signaux acoustiques sous-marins à partir de l' <i>Hearingogram</i> : <i>Denoised Hearingogram</i>	70
1.	Principe et analyse	70
2.	Reconstruction du signal utile à partir du <i>Denoised Hearingogram</i>	79
D.	Comparaison des différentes techniques	82
	Chapitre 3: Applications à la détection sous-marine dans le contexte du sonar passif.....	94
I.	Reconnaissance des signaux acoustiques sous-marins.....	95
II.	Principe d'un système de reconnaissance automatique.....	98
III.	Classification non supervisée	99
A.	Principe	99
B.	Etat de l'art.....	99
1.	Classification ascendante hiérarchique	100
2.	K-means	101
3.	Spectral Clustering	102
4.	DBScan	102
5.	Fuzzy K-Means	102
C.	Mesure de performance pour la classification non supervisée.....	102

D.	Temps-rythme.....	104
IV.	Classification supervisée.....	106
A.	Introduction	106
B.	Machines à vecteurs supports (SVM)	106
1.	Le choix des SVM.....	106
2.	Principe et calcul des SVM	106
3.	SVM non-linéaires	111
4.	Choix des paramètres	115
5.	SVM multi-classes	119
	Chapitre 4 : Caractérisation des signaux acoustiques sous-marins.....	124
I.	Descripteurs	125
A.	Représentation de l'information par des vecteurs de descripteurs	125
B.	Normalisation.....	125
C.	Détail des descripteurs utilisés	126
1.	Descripteurs temporels.....	127
2.	Descripteurs spectraux	127
3.	Descripteurs cepstraux	127
4.	Descripteurs perceptuels	128
D.	Discussions	129
II.	Sélection des descripteurs.....	130
A.	Nécessité d'une étape de sélection des descripteurs.....	130
B.	Différentes stratégies de recherches.....	131
1.	Best Individual N (BIN)	131
2.	Sequential (SEQ).....	131
3.	Optimisation des paramètres (PO)	132
C.	Différentes taxonomies d'algorithmes	132
1.	Classement ou sélection	132
2.	Différentes familles d'algorithmes de sélection des descripteurs : les filtres, les enrouleurs et les embarqués	133
D.	Etat de l'art de différents algorithmes.....	134
1.	Critère de Fisher.....	134
2.	Minimum Redundancy Maximum Relevance (MRMR).....	134
3.	Diversité marginale maximale (MMD)	135
E.	Extension du critère MMD sur plusieurs dimensions	136

F.	Test et résultats	137
Chapitre 5: Système automatique de reconnaissance des signaux acoustiques sous-marins		141
I.	Performance d'un modèle	142
A.	Métrique de performance	142
1.	Taux de bonnes classifications.....	142
2.	Matrice de confusion	142
B.	Evaluation des performances	143
1.	Méthode HoldOut	143
2.	Validation croisée	144
3.	Bootstrap	144
II.	Description du système automatique d'identification des signaux acoustiques sous-marins.....	145
A.	Segmentation temps-fréquence	145
B.	Classification manuelle et création des classes	145
C.	Calcul des descripteurs	147
D.	Décomposition du problème en problème binaire	147
E.	Sélection des descripteurs	147
F.	Paramétrage du classifieur SVM	148
G.	Création des frontières	148
H.	Mesure de performance	149
Conclusion générale		150
ANNEXE A : Reconstruction du signal temporel à partir de la transformée de Fourier à court terme		152
A.	Signal.....	152
B.	Fenêtre d'observation du signal	152
C.	Observation fenêtrée du signal	154
D.	Analyse harmonique du signal.....	154
E.	Expression algébrique de l'analyse harmonique fenêtrée du signal	155
F.	Expression matricielle des opérateurs.....	155
G.	Expression algébrique de l'analyse et de la synthèse harmonique fenêtrée	157
H.	Expression détaillée	162
Bibliographie.....		164

Table des figures

Figure 1 : Principe du SONAR actif [4]	16
Figure 2 : Principe du SONAR passif [4]	16
Figure 3 : chaîne de traitement de signaux des sous-marins	17
Figure 4 : Illustration du principe de formation de voies [4].....	19
Figure 5: Exemple de veille panoramique	23
Figure 6 : Illustration du principe de segmentation du plan temps-fréquence	24
Figure 7: Spectrogramme d'un chirp linéaire, analyse avec une fenêtre de Hamming de 16 s.....	31
Figure 8: Spectrogramme d'un signal composé de deux chirps linéaires et d'un signal monochromatique, avec une fenêtre de Hamming de 63 échantillons.	32
Figure 9: Spectrogramme d'un signal composé de deux chirps linéaires et d'un signal monochromatique, avec une fenêtre de Hamming de 17 s	32
Figure 10: Plan temps-fréquence d'un chirp linéaire obtenu par transformée de Wigner-Ville	36
Figure 11: Wigner-Ville d'un signal constitué d'un signal monochromatique et de deux chirps linéaires localisé en temps	36
Figure 12: Plan temps-fréquence obtenu par Wigner-Ville d'un chirp en présence de bruit blanc Gaussien (RSB=-3dB)	37
Figure 13: Plan temps-fréquence obtenu par Wigner Ville de l'observation bruitée en en tenant pas compte des interférences entre le signal et le bruit	38
Figure 14: Pseudo Wigner-Viller d'un signal constitué d'un signal monochromatique et de deux chirps linéaires localisé en temps	38
Figure 15: Pseudo Wigner-Ville lissée d'un signal constitué d'un signal monochromatique et de deux chirps localisés en temps.....	39
Figure 16: Spectrogramme réalloué d'un signal composé de deux chirps linéaires et d'un signal monochromatique.....	41
Figure 17: Pavage temps-fréquence classique	42
Figure 18: Pavage temps-fréquence en utilisant les ondelettes	42
Figure 19: L'oreille humaine [31].....	46
Figure 20: Diagramme de Fletcher et Munson (à gauche). Sons audibles de 20 à 20000 Hz (à droite) [31]	46
Figure 21: Banc de filtres de Mel, avec $M=10$ pour $f_{min} = 0 \text{ Hz}$ et $f_{max} = 11025 \text{ Hz}$;.....	50
Figure 22: Réponse impulsionnelle associée à $h5t$	50
Figure 23: Banc de filtres de Mel, avec énergie unitaire et $M=10$	51
Figure 24: Valeurs des fréquences étudiées pour chaque ligne du scalogramme (ligne rouge) et de l'Hearingogram (ligne noire)	54
Figure 25: Schéma fonctionnel de l' <i>Hearingogram</i>	54
Figure 26: Comparaison entre sclaoogramme (à gauche) et <i>Hearingogram</i> (à droite).....	54
Figure 27: Effet du nombre de filtres sur l' <i>Hearingogram</i>	55
Figure 28 : Effet d'un grand nombre de filtres sur l' <i>Hearingogram</i> avec 300 filtres (à gauche) et 800 filtres (à droite).....	56
Figure 29: Comparaison entre le spectrogramme (en haut) et l' <i>Hearingogram</i> (en bas) d'un son de dauphin.....	57

Figure 30: Comparaison entre le spectrogramme (en haut) et l' <i>Hearingogram</i> (en bas) d'écholocations d'orques	58
Figure 31: Comparaison entre le spectrogramme (en haut) et l' <i>Hearingogram</i> (en bas) de vocalises d'orque	59
Figure 32: Courbe de gain G_{Wien}	68
Figure 33: Courbes de gain G_{LSA}	68
Figure 34: Courbes de gain G_{MAP}	69
Figure 35: Courbes de gain G_{JMAP}	70
Figure 36: Schéma fonctionnel du principe de débruitage sur une ligne Zhm	74
Figure 37: Schéma fonctionnel du <i>Denoised Hearingogram</i>	75
Figure 38: Estimation de la densité de probabilité de Z_{hm} correspondant au $M^{ième}$ de Mel, $M=200$	76
Figure 39: <i>Hearingogram</i> (en haut) et <i>Denoised Hearingogram</i> (en bas) d'un son de dauphin Risso .	77
Figure 40: <i>Hearingogram</i> (en haut) et <i>Denoised Hearingogram</i> (en bas) d'écholocations d'orques ...	78
Figure 41: <i>Hearingogram</i> (en haut) et <i>Denoised Hearingogram</i> (en bas) d'un son de dauphin.....	79
Figure 42: Banc de filtre pour la reconstruction du signal (noir: banc de filtres de Mel; rouge: filtres ajoutés afin d'assurer la conservation de l'énergie)	80
Figure 43: Schéma fonctionnel du processus de débruitage proposé.....	81
Figure 44: Spectrogramme d'un signal représentant un morceau de piano échantillonné à 11 kHz	83
Figure 45: Spectrogramme d'un signal représentant un chant de dauphin, où la fréquence d'échantillonnage est égale à 16 KHz	83
Figure 46: Signal test (représentation temporelle et le spectrogramme associé).....	89
Figure 47: Densité de probabilité du bruit	89
Figure 48: Observation bruité et son spectrogramme associé	90
Figure 49: Signal test débruité par la méthode faisant intervenir le <i>Denoised Hearingogram</i> (signal temporel et spectrogramme associé)	90
Figure 50: Signal temporel d'écholocations d'orque et son spectrogramme	92
Figure 51 : Signal d'echolocations d'orque débruité par la méthode faisant intervenir le <i>Denoised Hearingogram</i> (signal temporel et spectrogramme associé)	93
Figure 52: Représentation des multi-trajets pour une onde sonore	96
Figure 53: Divergence sphérique.....	96
Figure 54: Amortissement du son dans l'eau de mer en fonction de la fréquence [46].....	97
Figure 55: Description des étapes de la classification supervisée	98
Figure 56: Exemple de dendrogramme portant sur cinq objets a, b ,c ,d ,e. Les points m, n, p, q sont les nœuds de l'arbre. Le trait horizontal mixte indique un niveau de troncature définissant une partition en trois classes.....	100
Figure 57 : Représentation temps-rythme (en bas) avec le modèle associé (en haut) sur un signal simulé contenant 3 trains de clics différents	105
Figure 58: Illustration du principe des SVM	107
Figure 59: Représentation graphique de l'exemple du XOR	112
Figure 60: Exemple de données non linéaires. Problème de discrimination binaire avec en vert les individus appartenant à la classe 1 et en bleu les individus appartenant à la classe 2.	113
Figure 61: Problème de l'échiquier avec représentation de la fonction de décision idéale. Les individus appartenant à la classe 1 sont en bleu et les individus appartenant à la classe 2 sont en rouge.	117

Table des figures

Figure 62: Surfaces de décision obtenues par apprentissage SVM à noyau gaussien RBF avec $\sigma=0.1$	117
Figure 63: Surfaces de décision obtenues par apprentissage SVM à noyau gaussien RBF avec $\sigma =0.3$	118
Figure 64: Surfaces de décision obtenues par apprentissage SVM à noyau gaussien RBF avec $\sigma =3$	118
Figure 65: Illustration du principe un contre tous, en gris se trouve la zone d'indétermination	120
Figure 66: Illustration du principe un contre un, la zone d'indétermination est hachurée au centre	121
Figure 67: Exemple de graphe de décision, pour un problème à 4 classes	122
Figure 68 : Procédure générale d'un algorithme de sélection des descripteurs	131
Figure 69: Schéma de la sélection de descripteurs de type filtre	133
Figure 70: Schéma principe de la sélection de descripteurs de type enrouleur	133
Figure 71: Schéma de principe de la sélection de descripteurs de type embarqué	134
Figure 72: Système de reconnaissance automatique	146
Figure 73 : Analyse temporelle fenêtrée du signal, avec $N = 11, N_h = 6, N_r = 2, N_o = 14, L = 3$	154

Liste des tableaux

Tableau 1: Bandes d'octave normalisées	47
Tableau 2 : Performances des différents algorithmes de réduction du bruit	85
Tableau 3: Mesures de performances des méthodes de débruitage utilisées, en vert les meilleurs performances et en rouge les performances les moins bonnes.	86
Tableau 4: Résultats moyens des méthodes de débruitage testées.....	87
Tableau 5: Résultat sur les vocalisations.....	91
Tableau 6: Résultat sur les signaux impulsifs	91
Tableau 7: Résultat sur le choc et sa trainée.....	91
Tableau 8: Résultat sur le signal à bande large	91
Tableau 9: Résultat sur signaux impulsif, autre type	91
Tableau 10: Résultats sur vocalisations discontinues	91
Tableau 11: Résultat sur le bruit	91
Tableau 12 : Table de vérité du Ou exclusif	112
Tableau 13 : Table de vérité du Ou exclusif après application d'une transformation	113
Tableau 14 : Caractéristiques des bases employées pour l'évaluation.....	138
Tableau 15: Taux de bonnes classifications avec le classificateur naïves de Bayes et les SVM avec et sans sélection des descripteurs.....	138
Tableau 16: Taux de bonnes classifications avec les SVM, comparaison entre différents algorithmes de sélection des descripteurs.....	139

Introduction générale

Dans le cadre des études et développements menés dans le domaine de la détection et de la reconnaissance des signaux acoustiques sous-marins, cette thèse a pour but de définir et concevoir de nouvelles techniques de représentation des signaux. L'objectif est de faire en sorte que ces techniques augmentent la capacité d'un système SONAR passif à reconnaître et interpréter les signaux reçus par l'antenne placée en amont du système. La finalité de cette démarche n'est pas de substituer la machine à l'humain, dont l'expérience et la finesse d'ouïe le rendent indispensable, mais au contraire de le soulager en lui proposant l'aide à la décision la plus pertinente possible.

Historiquement, le monde de la lutte sous-marine s'est intéressé aux signaux dits stationnaires, c'est à dire présentant des caractéristiques statistiques "relativement" stables dans le temps d'observation. Plus récemment, les signaux non stationnaires ont fait l'objet d'un intérêt particulier pour leur caractère classifiant et énergétique. Une première définition pourrait être de considérer ces signaux comme le complémentaire des signaux stationnaires dans l'ensemble des signaux d'énergie finie. Une seconde définition, plus "physique", serait de considérer un signal non stationnaire comme tout signal de support temporel limité ou bien présentant une variabilité spectrale substantielle dans le temps d'observation. Traditionnellement, ces signaux sont analysés par l'intermédiaire d'une représentation temps-fréquence (RTF) effectuée en première intention par transformée de Fourier à court terme (TFCT), puis en fonction du besoin par les distributions de Wigner-Ville ou encore par transformées en ondelettes. Quel que soit le type choisi de RTF, un compromis entre résolution temporelle et résolution fréquentielle est nécessaire. Chacune de ces représentations a ses avantages et ses inconvénients. Dans notre cas, la RTF sert de donnée d'entrée à un système permettant l'identification des signaux, ainsi ce compromis d'analyse conditionne directement les performances du système, dans le sens où plus la RTF est adaptée à un type de signaux, meilleure sera l'identification de ce dernier.

Les travaux de cette thèse puisent leur inspiration dans ce qui se fait de mieux dans le domaine de la reconnaissance acoustique : l'être humain. Plus spécifiquement, dans le domaine de l'acoustique-sous-marine militaire, cette tâche de reconnaissance est affectée à des experts de l'analyse des bruits sous-marins, appelés « Oreilles d'or ». Ainsi, afin de restituer au mieux le contenu d'un signal audio, une approche consiste à tenter de bio-mimer la capacité de l'être humain à reconnaître des sons, fort de l'efficacité de son acuité auditive. En effet, l'oreille humaine a le pouvoir de différencier deux sons distincts à travers des critères perceptuels psycho-acoustiques tels que le timbre, la hauteur, l'intensité. Etant donné ce contexte applicatif, nous avons conçu une représentation qui se rapproche de la physiologie de l'oreille humaine, autrement dit, de la façon dont l'homme perçoit les fréquences. Pour construire cet espace de représentation, nous utilisons un algorithme original, que nous avons appelé *Hearingogram* et le *Denoised Hearingogram*. En ce qui concerne ce dernier, il correspond au couplage de l'*Hearingogram* avec une technique de réduction du niveau de bruit, spécifiquement dédiée à cette représentation. Ces deux représentations ont pour finalité d'être placées en amont d'un système de reconnaissance.

Il est légitime de se demander comment l'être humain arrive à reconnaître les sons. Apporter une réponse formelle à cette question est une chose très difficile. En effet, si une personne est capable d'identifier un son c'est bien parce que cette action, comme la plupart des processus cognitifs, échappe à la nécessité d'une définition formelle. En réalité, elle repose sur l'apprentissage empirique de très nombreux exemples associés à une ou plusieurs classes, qui nous permet de reconnaître celles-ci en présence d'exemples inconnus. On peut donc dire que le cerveau est alors plus une machine associative qu'une machine logique. Nous utilisons ainsi le même principe pour l'identification automatique, qui vise à affecter des catégories ou des classes à des objets. Dans notre cas, les objets sont des signaux acoustiques sous-marins, dont il faut renseigner préalablement la classe. Cette étape est effectuée avec l'aide d'un expert.

Le manuscrit est organisé de la façon suivante :

- La première partie traite de la présentation de la chaîne de traitement SONAR passif. En particulier l'étude de la genèse du signal en entrée de la chaîne, ainsi que les différents traitements qui sont appliqués au signal.
- La seconde partie décrit la problématique de la représentation du signal acoustique sous-marin, plus précisément il sera fait état des représentations temps-fréquences et temps-échelle, ainsi que de leur mise en pratique informatique. Cette partie décrit également l'*Hearingogram*, représentation basée sur les filtres de Mel, qui caractérisent le comportement de l'oreille humaine. Nous nous sommes ensuite intéressés à la réduction du bruit, et une comparaison des différentes méthodes de ce domaine a été réalisée.
- La troisième partie, traite de la méthodologie de reconnaissance acoustique des signaux-sous-marins, à l'aide du classifieur SVM¹, fondé sur le principe de la séparation à vaste marge sous contrainte d'une bonnes classifications des exemples d'apprentissage. Ainsi, le paramétrage de cet algorithme est décrit afin de définir le jeu de paramètres conduisant aux meilleurs résultats possibles.
- Ensuite, la quatrième partie traite de la question de la représentation des signaux acoustiques par des descripteurs, et de l'étape de sélection de ces derniers.
- Enfin avant de conclure ces travaux de recherche, la dernière partie traite de l'aspect pratique de la reconnaissance des signaux acoustiques sous-marins. Dans ce cadre, un système de reconnaissance automatique est proposé.

¹ Support vectors machine

Chapitre 1 : Chaîne de traitement des signaux sous-marins

I. Introduction

Le milieu marin a toujours été un lieu d'échanges, vital pour l'économie de certains pays, stratégique pour les intérêts d'autres pays. Pour ne s'en tenir qu'à des aspects géopolitiques, les mers et océans sont des lieux de transports de chargements vitaux, à la fois pour les pays pauvres (exportations de matières premières) et pour les pays riches (importations de ces mêmes matières premières).

Depuis l'apparition des sous-marins à la fin du 19^{ème} siècle, les systèmes de reconnaissance sous-marine s'appuient sur le principe de la propagation des ondes acoustiques, car ce sont celles qui se propagent le mieux dans l'eau. Détecter les cibles en toute discrétion, c'est-à-dire en conservant l'avantage acoustique, tel est le défi des sous-marins. Pour y parvenir, ils disposent de systèmes SONAR² qui tentent de détecter et d'analyser les ondes acoustiques captées par les antennes. Le SONAR est composé de modules de traitement complexes dont le but est de transformer le signal acoustique perçu par les antennes en éléments informatifs. Ces derniers nous permettent de réaliser trois étapes essentielles en détection sous-marine, à savoir :

- La détection : extraction du signal utile au sein de l'observation.
- La classification : identification du signal utile extrait.
- La localisation : trajectographie de la cible.

On distingue deux types de SONAR, les SONAR passifs [1] [2] et les SONAR actifs [1] [3]. Un SONAR actif est composé d'une partie émission qui émet une onde sonore et une partie réception qui traite l'onde réfléchi. Le SONAR passif, quant à lui, n'est composé que de la partie réception et traite donc les signaux sonores perçus dans l'eau. La partie réception d'un SONAR (actif ou passif) est composée de 2 grands blocs :

- un ensemble de capteurs acoustiques appelés hydrophones regroupés pour former une antenne que l'on qualifiera de linéaire, cylindrique ou sphérique selon la répartition géométrique des capteurs ;
- un système informatique situé à bord du sous-marin et composé de différents modules prenant en charge une fonction particulière (détection, classification, trajectographie, ...).

Nous allons détailler chacun de ces modules dans les paragraphes suivants.

² Sound Navigation And Ranging

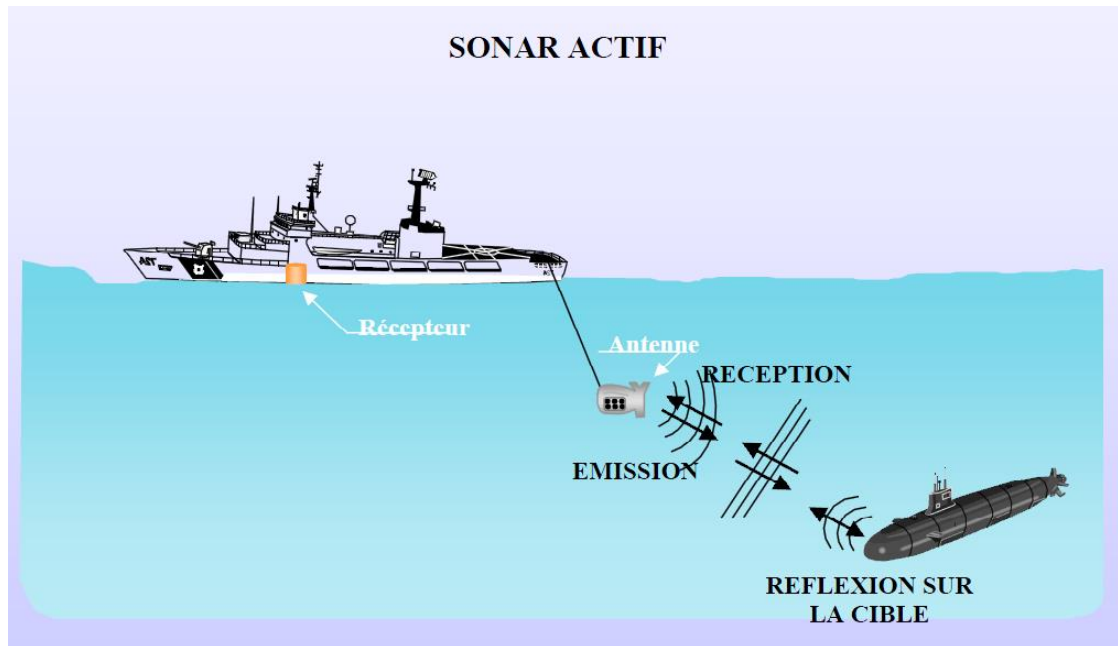


Figure 1 : Principe du SONAR actif [4]

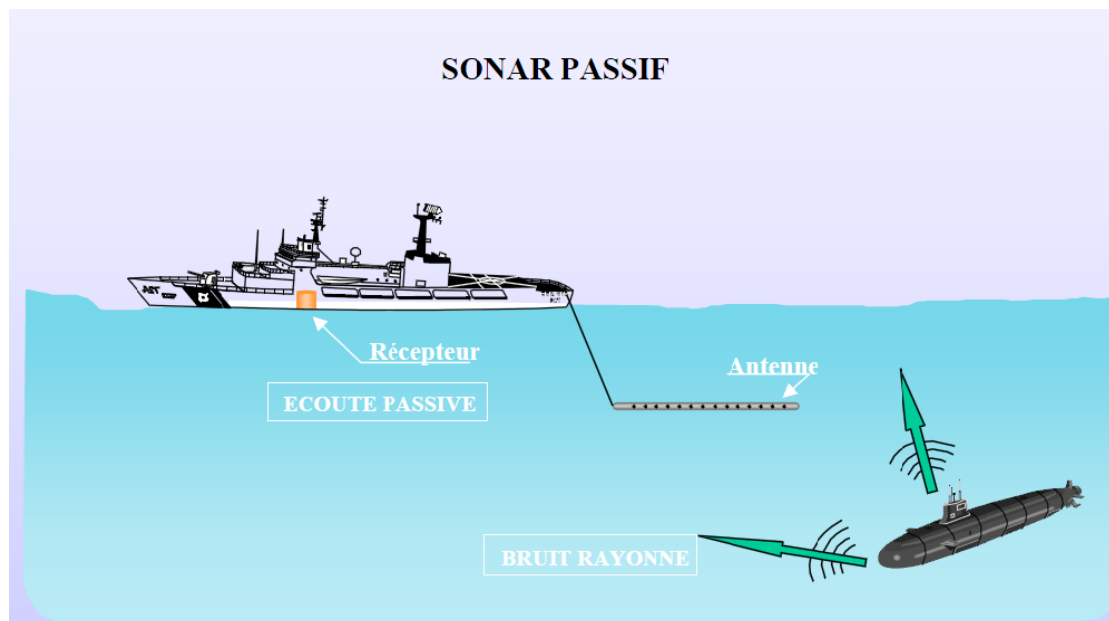


Figure 2 : Principe du SONAR passif [4]

Les travaux développés dans le cadre de cette thèse s'inscrivent dans le contexte d'un SONAR passif et s'intéressent plus particulièrement à la représentation et à la reconnaissance des signaux sous-marins.

La chaîne de traitement pour les signaux sous-marins peut se résumer par les étapes suivantes :

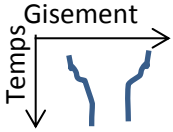
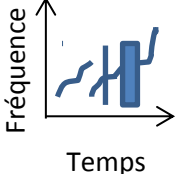
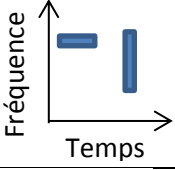
Acquisition des données	
Formation de voies	$f(t, \theta) = \frac{1}{K} \sum_{k=0}^{K-1} g_k h_k(t - \tau_k(\theta))$
Veille panoramique	
Détection	$Z > \text{Seuil}$
Représentation temps-fréquence	
Réduction du bruit	$\hat{S} = Z - B$
Segmentation temps-fréquence	
Reconnaissance	

Figure 3 : chaîne de traitement de signaux des sous-marins

Nous allons expliquer ces différentes étapes dans les paragraphes suivants.

II. Acquisition des données

Les données en entrée d'un récepteur passif sont recueillies par des hydrophones regroupés en antenne. Il existe plusieurs types d'antennes dédiées à l'exploitation et à la reconnaissance de signaux ayant des caractéristiques différentes :

- Antennes linéaires immergées, remorquées par un bâtiment de surface,

- Antennes linéaires fixées au fond de la mer qui écoutent le trafic maritime dans une zone géographique donnée,
- Antennes sphériques, cylindriques etc...

En agissant sur l'espacement entre capteurs, sur le nombre de capteurs, sur la longueur d'une antenne linéaire, on modifie ses caractéristiques. Ainsi la distance inter-capteurs détermine la gamme des ondes perçues par l'antenne. La longueur de l'antenne et la distance inter-capteurs caractérisent sa résolution angulaire. La fonction de directivité de l'antenne, qui dépend du nombre de capteurs, fait apparaître la résolution angulaire en fonction (par exemple) de la voie formée ou de la plage de fréquences traitée. Les caractéristiques des hydrophones modifient également les caractéristiques de l'antenne.

A ce stade de traitement nous disposons d'un signal analogique, nous devons donc le filtrer au niveau de chaque capteur, l'amplifier et l'enregistrer, pour plus de renseignements sur ces dernières étapes le lecteur pourra se référer à [5]. L'étape d'échantillonnage discrétise le signal dans l'espace des temps alors que la quantification introduit un codage en valeur, c'est une discrétisation dans l'espace des amplitudes. Ainsi un nombre binaire obtenu lors de la numérisation représente une discrétisation en niveau et en temps du signal. L'information est donc devenue un signal numérique pour chaque capteur que nous devons remettre en phase.

III. Formation de voies

Tout d'abord avant d'entamer ce chapitre nous allons simplement rappeler deux notions qui sont le gisement et l'azimut :

- L'azimut est l'angle formé entre la direction d'une source et le nord.
- Le gisement est l'angle formé entre l'axe longitudinal d'un navire et la direction de la source.

La formation de voies a pour objectif la description d'une situation physique dépendant du temps et de l'espace au moyen d'une antenne. La formation de voies pourra être interprétée comme un filtrage spatial qui permettra donc d'augmenter les performances en détection, de résoudre de multiples sources et de mesurer les directions de ces sources. On peut la considérer comme un filtrage spatial car si du bruit est présent dans l'observation, il s'additionnera de manière incohérente, alors que le signal qui vient d'une direction θ est pris en compte de manière cohérente.

En toute généralité, la formation de voies consiste à extraire d'un signal les composantes se propageant dans une direction particulière θ et à une vitesse donnée (les composantes pouvant être de fréquences différentes). L'estimation de ce signal se fait grâce à une mesure faite par une antenne réseau de K capteurs. Plusieurs techniques de formation de voies existent [6] tel que :

- Temporelle que nous expliquerons ci-dessous ;
- Fréquentielle, elle est réalisée après analyse spectrale de chaque voie hydrophonique ;
- En deux couches, site puis gisement ou voies grossières puis voies fines ;

- Adaptative, l'objectif de cette dernière technique est la réjection des perturbations sur l'antenne ;
- Avec soustraction de bruit, où nous avons une voie de référence soustraite aux voies pointées (ex : bruit du porteur) ;
- Focalisée, où nous prenons en compte la courbure du front d'onde ;
- A ouverture synthétique, où nous prenons en compte le déplacement de l'antenne.

Nous nous concentrerons uniquement sur la formation de voie temporelle. Le principe est le suivant, le signal temporel $h_k(t)$ reçu par l'un des capteurs k de ce réseau peut être amplifié par un gain g_k et retardé d'un temps τ_k . Le signal formé est la moyenne de ces différents signaux:

$$f(t, \theta) = \frac{1}{K} \sum_{k=0}^{K-1} g_k h_k(t - \tau_k(\theta))$$

III-1

L'expression de $\tau_k(\theta)$ est liée à la géométrie de l'antenne et aux propriétés du milieu. Le choix approprié des gains et des retards permet d'«orienter» le réseau de capteurs dans une direction donnée. Sur la Figure 4 nous voyons une illustration du principe de la formation de voies temporelle.

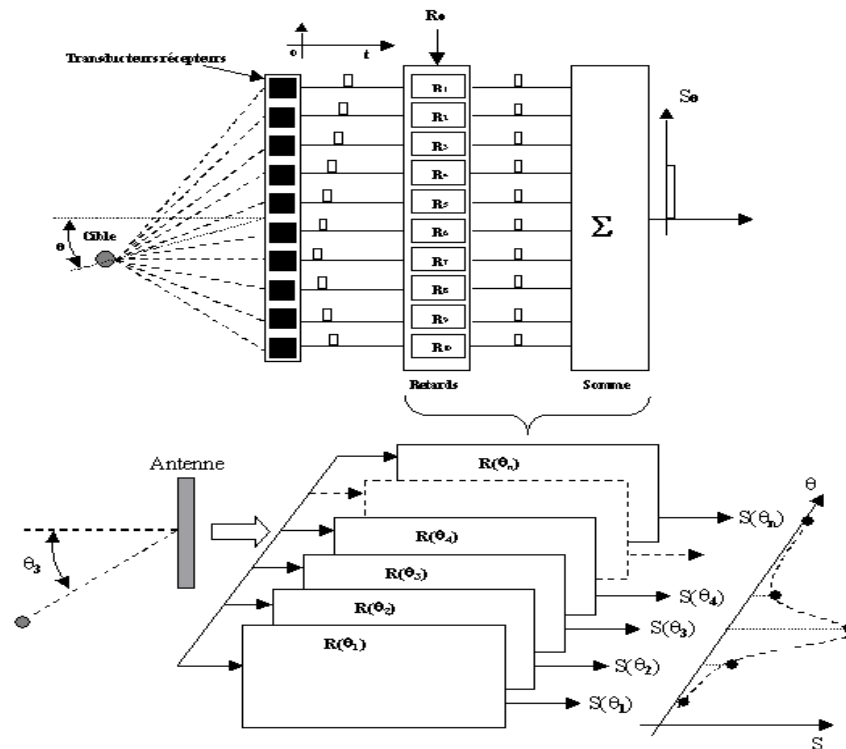


Figure 4 : Illustration du principe de formation de voies [4]

IV. Veille panoramique

Chaque voie formée fait l'objet d'une analyse spectrale. Nous obtenons alors des triplets gisement ou azimuth-fréquence-temps auxquels on associe la valeur de l'énergie, notée $a(\theta, f, \tau)$. Plus précisément :

- Chaque voie traduit le signal perçu dans une direction privilégiée : elle est donnée pour un gisement (ou azimuth) particulier ce qui réalise ainsi une discrétisation de l'espace angulaire. Si le secteur angulaire balayé est tel que :

$$\theta \in [\theta_{min}; \theta_{max}]$$

Et que cet intervalle est divisé en N_θ secteurs angulaires, alors :

$$\Delta\theta = \frac{(\theta_{max} - \theta_{min})}{N_\theta} \quad \text{IV-1}$$

est la valeur du pas angulaire.

Ainsi chaque secteur angulaire est repéré de sorte que $\forall j = 0 \dots N_\theta - 1$:

$$\theta_j = \theta_{min} + j\Delta\theta \quad \text{IV-2}$$

- La fréquence est elle aussi discrétisée, en effet, l'application de la transformée de Fourier discrète à un vecteur temporel composé de N_h échantillons, conduit à un vecteur fréquence de même longueur. Nous pouvons néanmoins utiliser le « zéro-padding » [7], qui est une technique de sur-échantillonnage en fréquence ce qui revient à considérer $N_{FFT} \geq N_h$, où N_{FFT} est le nombre d'échantillons fréquentiels atteints.

On a N_{FFT} canaux fréquentiels avec une finesse d'analyse définie par :

$$\Delta f = \frac{F_e}{N_{FFT}} \quad \text{IV-3}$$

ainsi les canaux analysés ont pour valeur $\forall k = 0 \dots N_{FFT} - 1$:

$$f_k = k\Delta f \quad \text{IV-4}$$

- Pour un secteur angulaire donné, les transformées de Fourier sont calculées à partir des échantillons situés dans une fenêtre temporelle de longueur $T_h = N_h/F_e$, éventuellement élargie jusqu'à $\frac{N_{FFT}}{F_e}$ lors du zéro-padding. En revanche, le temps de récurrence du traitement peut être bien plus court car il s'agit de l'intervalle de temps séparant deux fenêtres glissantes consécutives et potentiellement en recouvrement.

- Ainsi, les échantillons temporels sont définis de sorte que $\forall l = 0 \dots N_{recc} - 1$:

$$\tau_l = t_0 + lT_h \quad \text{IV-5}$$

où t_0 est la date initiale de la veille panoramique de l'acquisition des données et N_{recc} le nombre de récurrences étudiées.

Le maillage de l'espace gisement-fréquence-récurrance est donc effectué de sorte que

$$\{(\theta_j ; f_k ; \tau_l)\}_{\substack{\forall j=0 \dots N_\theta-1 \\ \forall k=0 \dots N_{FFT}-1 \\ \forall l=0 \dots N_{recc}-1}} \quad \text{IV-6}$$

Permettant ainsi d'avoir à disposition l'ensemble de valeurs

$$a(\theta_j ; f_k ; \tau_l) \quad \text{IV-7}$$

Dans un premier temps, nous calculons l'énergie contenue dans une direction θ_j pour la récurrence τ_l selon les fréquences :

$$A(\theta_j ; \tau_l) = \sum_{k=0}^{N_{FFT}-1} a(\theta_j ; f_k ; \tau_l) \quad \text{IV-8}$$

En réalisant ce calcul $\forall j = 0 \dots N_\theta - 1$ et $\forall l = 0 \dots N_{recc} - 1$ nous obtenons une veille panoramique. Cependant, à cette étape, la résolution angulaire de cette dernière n'est pas encore acceptable. Nous réalisons donc une série de traitements qui ont pour but d'augmenter la résolution angulaire [8].

A l'issue de ces traitements nous obtenons une image temps-gisement appelée veille panoramique. Elle représente la variation de l'énergie totale, éventuellement restreinte à une bande de fréquence spécifique, dans les différentes voies au cours du temps.

En pratique, la phase de mise au point d'un tel système est longue et fastidieuse. Obtenir des résultats à la mer sur signaux réels a demandé plusieurs décennies de travail, et demeure un problème ouvert.

L'étape suivante de la chaîne consiste à extraire de la veille panoramique les événements considérés être du signal utile. Les différents traitements sont décrits dans le paragraphe suivant.

V. Extraction et poursuite sur veille panoramique

L'opération d'extraction-poursuite fait la liaison entre le traitement du signal et le traitement de l'information. En pratique elle peut être manuelle ou automatique. Elle consiste à relier entre eux de manière cohérente les événements détectés³. Ainsi une extraction peut être vue comme :

$$\check{A}(\theta_j; \tau_l) = A(\theta_j; \tau_l) \varepsilon(\theta_j; \tau_l) \quad \text{V-1}$$

où $\varepsilon(\theta_j; \tau_l)$ peut être vu comme la sortie d'un organe de détection à seuil. Soit $\xi > 0$ ce seuil, alors l'opération de seuillage est définie par :

$$\varepsilon(\theta_j; \tau_l) = \begin{cases} 0 & \text{si } A(\theta_j; \tau_l) < \xi \\ 1 & \text{si } A(\theta_j; \tau_l) \geq \xi \end{cases} \quad \text{V-2}$$

Une fois ces étapes réalisées, deux éventualités peuvent être envisagées :

- Soit l'évènement détecté se trouve au sein d'une voie (donc une direction) et on extrait le signal audio correspondant à cette direction. Dans ce cas le signal est sur une voie pointée :

$$\theta(\tau) = \theta_{j_0} \quad \text{V-3}$$

Où j_0 est le secteur angulaire dans lequel le signal utile se trouve.

- Soit l'évènement détecté change de voie au cours du temps, dans ce cas nous sommes asservis à pister cet événement et donc la voie est dépendante du temps. Dans ce cas, on a :

$$\theta(\tau) = \phi(\tau) \quad \text{V-4}$$

Où $\phi(.)$ est une fonction quelconque et suffisamment régulière pour modéliser la trajectoire du signal utile observé sur la veille panoramique.

Enfin, nous construisons le signal audio correspondant à la piste détectée.

L'ensemble des étapes conduisant à la synthèse de ce signal n'est pas détaillé, mais regroupées dans l'opérateur $\mathcal{H}$. Ainsi on a $\forall l = 0 \dots N_{recc} - 1$:

$$Z = \mathcal{H}[\check{A}(\phi(\tau_l); \tau_l)]. \quad \text{V-5}$$

Ainsi le vecteur Z contient les N échantillons constitutifs du signal audio à la fréquence d'échantillonnage F_e . Le signal audio ainsi constitué peut être vu comme une observation de la forme :

$$Z = S + B \quad \text{V-6}$$

où S représente le signal utile et B le bruit. Généralement un tel signal audio s'étudie à l'aide d'une représentation temps-fréquence. Sur la Figure 5 on peut voir un exemple de veille panoramique :

³ Cette fonction s'inspire de la faculté d'intégration visuelle de l'opérateur

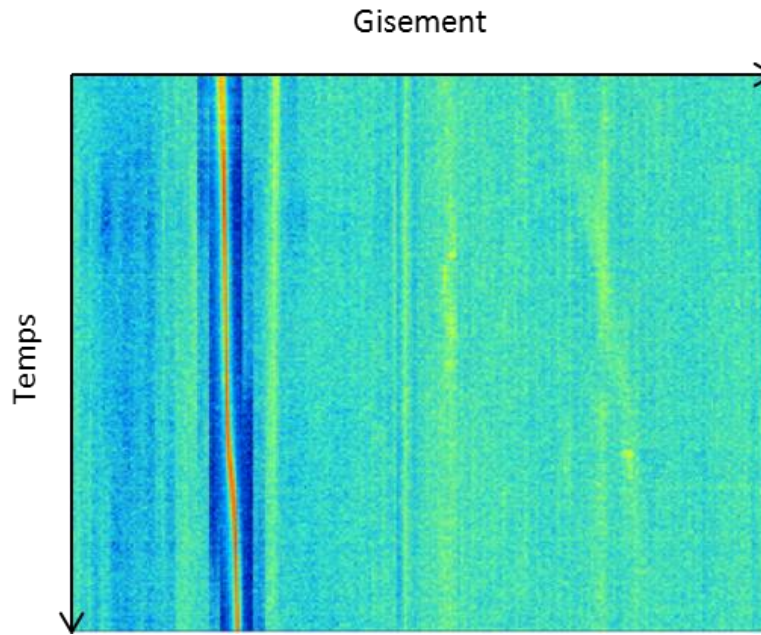


Figure 5: Exemple de veille panoramique

L'image est de type temps-gisement, nous remarquons en rouge une source de bruit qui change de gisement au cours du temps.

VI. Segmentation d'une représentation temps-fréquence

A cette étape, aucune hypothèse n'est faite sur la stationnarité du signal utile S . Il peut présenter un caractère non-stationnaire propre aux signaux réels rencontrés. Dans ces conditions, l'analyse temps-fréquence est la technique qui semble la plus appropriée. Dans notre chaîne de traitement, le but est d'extraire le signal utile qui se trouve dans l'observation afin de soumettre ce dernier à un organe de reconnaissance. Cette tâche se déroule en trois étapes :

- représentation temps-fréquence,
- réduction du niveau de bruit,
- segmentation des motifs représentant le signal utile.

L'opération de segmentation peut être vue comme la donnée d'un sous-ensemble du plan temps-fréquence appelé pavé. Ce dernier est défini par le temps de début et de fin et la fréquence minimale et maximale du signal utile considéré. Sur signaux réels plusieurs problèmes peuvent intervenir, comme :

- la superposition des pavés au sein de l'observation,
- l'atténuation du signal utile perçue par les antennes due à l'absorption du milieu marin,
- la déformation du signal utile due à la variabilité du milieu marin (champ sonore complexe additionné de phénomènes de réverbération et de réflexion),
- la perte d'information due à la quantification du signal.

Les détails de cette étape de segmentation du signal utile sur un plan temps-fréquence sortent du cadre de cette étude, nous considérons donc cette dernière comme un système de type boîte noire prenant en entrée une représentation temps-fréquence et fournissant en sortie un ensemble de pavés. La figure suivante illustre la finalisation de cette étape :

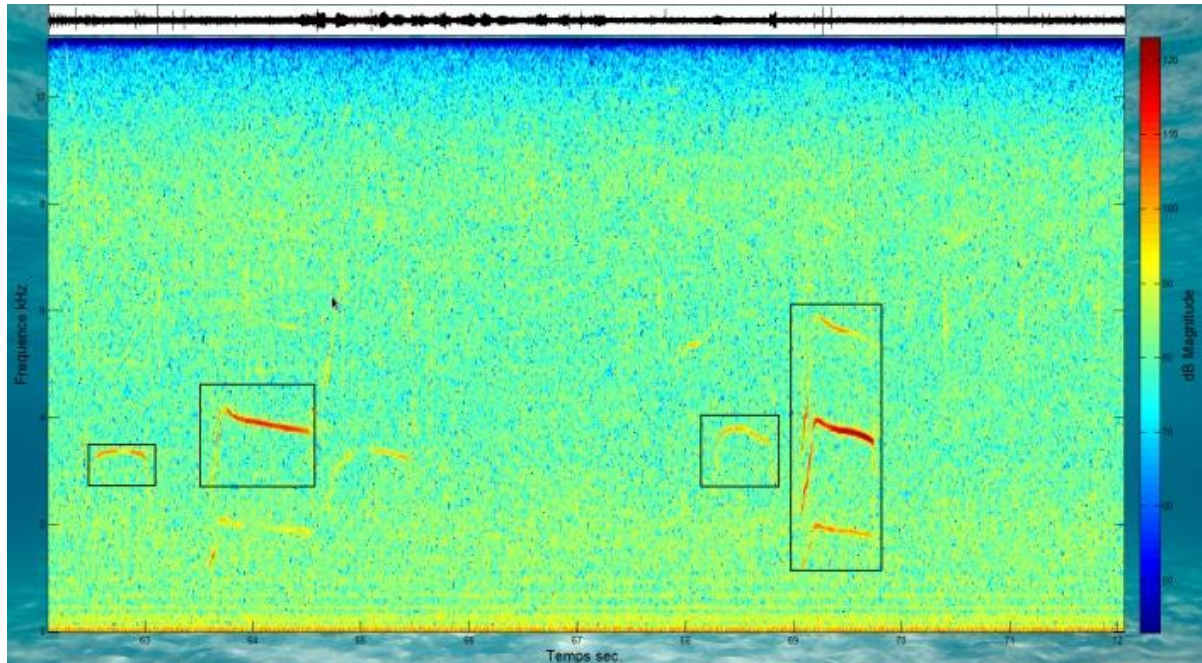


Figure 6 : Illustration du principe de segmentation du plan temps-fréquence

VII. Reconnaissance et identification des pavés temps-fréquence détectés

L'objectif de cette étape est d'interpréter et reconnaître de façon automatique les signaux extraits. En pratique, cette étape est réalisée par des experts en reconnaissance acoustique que l'on appelle « oreille d'or ». A bord d'un sous-marin ou en centre à terre, ces experts écoutent les signaux afin d'en extraire les sons d'intérêts. Ils utilisent de manière combinée les informations audiophoniques, les informations visuelles ainsi que les informations de contexte. De tels spécialistes sont rares, hyper-entraînés, si bien qu'il serait illusoire d'espérer égaler la pertinence de leurs décisions à l'aide d'un processus automatisé. Les bruits écoutés peuvent être d'origine biologique, sismique ou anthropique tels que des bateaux de commerce ou de guerre, des sous-marins, des drones, des travaux sous-marins. L'être humain a la capacité de déterminer la classe d'un signal sonore sur la base de son acuité auditive. En effet, l'oreille humaine a le pouvoir de différencier deux sons distincts à travers des critères perceptuels psycho-acoustiques tels que le timbre, la hauteur, l'intensité. Pour soulager l'utilisateur de la grande quantité d'information qu'il doit traiter, l'idée est d'automatiser

certaines traitements de l'information constituant les tâches les plus fastidieuses, afin de lui proposer une aide à la décision. C'est sur la base de ce principe que nous allons orienter nos travaux.

Ainsi, partant d'un signal audio, les étapes conduisant à la reconnaissance des pavés temps-fréquences détectés dans l'observation sont les suivantes :

- représentation temps-fréquence,
- réduction du niveau de bruit,
- recherche d'indices, puis calcul des éléments descriptifs du signal,
- décision quant à l'appartenance à une classe.

Le chapitre suivant présente les différents modes de représentation temps-fréquence des signaux. Un état de l'art des techniques les plus utilisées en acoustique sous-marine est présenté. De plus, une méthode originale de représentation basée sur l'oreille humaine, appelée l'*Hearingogram*, sera exposée. Enfin cette partie traitera également de la problématique de la réduction du niveau de bruit.

Chapitre 2: Représentations des signaux acoustiques non-stationnaires

I. Représentations temps-fréquence

A. Introduction

Un signal non stationnaire peut se définir par complémentarité à un signal stationnaire, c'est-à-dire par des propriétés statistiques qui varient au cours du temps ou bien par une durée d'existence limitée par rapport à la durée d'observation.

Une représentation du signal comme fonction du temps uniquement donne peu d'indications sur le comportement en fréquence, tandis que son analyse de Fourier ne fait pas apparaître l'instant d'émission et la durée de chacun des éléments composant le signal. Pour les besoins du traitement du signal on a cherché à représenter simultanément les informations temporelles et fréquentielles contenues dans le signal. Ces transformées sont appelées représentations temps-fréquence (RTF). Il convient de considérer le passage par RTF au domaine temps-fréquence, comme une distribution de l'information contenue dans le signal analysé de façon à en faciliter l'interprétation. Les RTF présentent l'avantage majeur de mettre en évidence les comportements non-stationnaires du signal, en effet un signal réel est de façon globale non stationnaire mais il est localement stationnaire, ainsi nous pouvons utiliser des traitements dédiés aux signaux stationnaires sur les parties du signal qui le sont. La relation existant entre représentation temps-fréquence et représentation fréquentielle d'un signal peut se comparer à la relation entre une suite de notes musicales écrites sur une portée et l'histogramme de ces notes, en effet une partition rend compte à la fois de la fréquence des notes mais également de l'ordre dans lequel elles doivent être jouées, alors que l'histogramme nous donne juste les proportions dans lesquelles elles seront présentes.

Diverses méthodes engendrent des RTF de propriétés et performances variées [9]. Dans ce chapitre, un état de l'art des principales méthodes de représentation temps-fréquence est exposé. Pour chaque technique, une version continue et une version discrète seront présentées, afin de pouvoir implanter chacune d'entre elles sur un support informatique.

B. Transformée de Fourier à court terme (TFCT)

1. Principe

La transformée de Fourier à court terme (TFCT) considère implicitement un signal non stationnaire comme une succession de situations localement stationnaires. Historiquement la TFCT, ainsi que ses variantes telles que la transformée de Gabor [10] ont été les premières méthodes proposées. La TFCT utilise une fonction $h(t)$, appelée fenêtre d'analyse. Cette fenêtre sélectionne une portion du signal et se déplace par translation le long de l'axe temporel, on la qualifie de fenêtre glissante. A chaque déplacement de la fenêtre, qui peut s'effectuer avec ou sans recouvrement, on associe le spectre instantané de la portion de signal analysé par la fenêtre à la position temporelle de la fenêtre. Ainsi est obtenue la définition de TFCT, d'un signal $s(t)$:

$$TFCT(t, f) = \int s(u)h^*(u - t)e^{-2i\pi fu} du \quad \text{I-1}$$

Pour obtenir le spectrogramme nous considérons le module au carré de la TFCT :

$$S_s^h(t, f) = \left| \int s(u) h^*(u - t) e^{-2i\pi f u} du \right|^2 \quad \text{I-2}$$

La transformée de Gabor [10] quant à elle est une TFCT où la fenêtre d'analyse choisie est :

$$h(t) = e^{\frac{-t^2}{2\sigma^2}} \quad \text{I-3}$$

C'est donc une fenêtre de forme gaussienne dont le support est déterminé par le paramètre σ . L'écriture discrète de la TFCT est :

$$\text{TFCT}(s[n]) = \text{TFCT}[k, l] = \sum_{n=-\infty}^{+\infty} s[n] h[n - l] e^{\frac{-2i\pi k n}{N}} \quad \text{I-4}$$

Avec h la fenêtre glissante, k le canal fréquentiel et l'indice de l localisation temporelle correspondant à une récurrence temporelle et N le nombre d'échantillons du signal s . On montre dans [11] que la longueur de la fenêtre est le paramètre prédominant de cette méthode.

2. Avantages et limitations

Lors de l'utilisation de la TFCT on considère que le signal est localement stationnaire au sein de la fenêtre à court terme $h(t)$, la résolution temporelle d'une telle analyse est fixée par la largeur de cette fenêtre, la résolution fréquentielle étant fixée par la finesse d'analyse de sa transformée de Fourier. Or, ces deux grandeurs sont inversement proportionnelles, donc les représentations de type TFCT ne conduisent pas à une localisation idéale à la fois en temps et en fréquence des composantes du signal. Par ailleurs, la localisation parfaite d'un atome consisterait à pouvoir diminuer infiniment son support. Or, toute diminution du support temporel provoque un élargissement du support fréquentiel du signal, et réciproquement. En effet, classiquement pour exprimer la largeur du support temporel d'un signal et sa largeur de bande, on a recours à la moyenne pondérée dans le domaine temporel :

$$\langle t \rangle = \frac{1}{E_s} \int t |s(t)|^2 dt \quad \text{I-5}$$

et à la moyenne pondérée dans le domaine fréquentiel :

$$\langle f \rangle = \frac{1}{E_s} \int f |S(f)|^2 df \quad \text{I-6}$$

Où E_s représente l'énergie du signal définie par :

$$E_s = \int |s(t)|^2 dt \quad \text{I-7}$$

Le signal est alors dit concentré en $(\langle t \rangle, \langle f \rangle)$ dans le domaine temps-fréquence. La détermination de la largeur des distributions autour de $(\langle t \rangle, \langle f \rangle)$ s'obtient à l'aide des écarts-type Δ_t et Δ_f du signal dans le domaine temps-fréquence des distributions, avec:

$$\left\{ \begin{array}{l} \Delta_t = \sqrt{\frac{1}{E_s} \int (t - \langle t \rangle)^2 |s(t)|^2 dt} \\ \Delta_f = \sqrt{\frac{1}{E_s} \int (f - \langle f \rangle)^2 |S(f)|^2 df} \end{array} \right. \quad \text{I-8}$$

Qui permettent de mesurer la dispersion de l'énergie dans le plan temps-fréquence. Le principe d'incertitude temps-fréquence repose sur l'inégalité d'Heisenberg-Gabor [12] qui lie les écarts-types précédemment définis :

$$\Delta_t \Delta_f \geq \frac{1}{4\pi} \quad \text{I-9}$$

L'inégalité nous dit que les supports en temps et en fréquence sont dépendants l'un de l'autre, avec égalité si et seulement si $h(t)$ est une gaussienne.

Cette relation est empruntée au principe d'incertitude, déduit en 1927 par Heisenberg dans le domaine de la mécanique quantique. Pour les signaux, le principe d'incertitude postule qu'il n'est pas possible de construire un signal dont le support est infiniment petit à la fois en temps et en fréquence. Ainsi, on peut donc dire que :

- Pour un signal fortement non-stationnaire, une bonne résolution temporelle est requise, ce qui impose de travailler avec une fenêtre $h(t)$ courte, limitant en retour la résolution fréquentielle ;
- Réciproquement, si une analyse fréquentielle fine est nécessaire, une fenêtre $h(t)$ longue doit être utilisée, ce qui a le double effet de moyenner les contributions fréquentielles sur la durée de la fenêtre et de réduire la résolution temporelle.

Bien que fournissant une représentation admissible d'un signal non- stationnaire, on voit donc que la TFCT ne permet pas une analyse à la fois locale en temps et précise en fréquence.

Cependant la TFCT joue un rôle clef au sein des RTF, car elle possède de nombreuses propriétés utiles pour l'analyse des signaux réels. En particulier, elle préserve les translations temporelles, à un facteur de modulation près, ainsi que les translations fréquentielles. Ses avantages sont nombreux et non négligeables :

- Interprétation facile car il n'y pas d'interférences et les échelles sont linéaires en temps et en fréquence.
- Faible coût en temps de calcul grâce à l'algorithme de la FFT (Fast Fourier Transform), que l'on doit à James Cooley et John Tuckey en 1965 [13].
- Mise en œuvre algorithmique assez simple.

Cette méthode est utilisée dans de nombreuses applications issues de divers domaines, on peut citer notamment les domaines suivants :

- Acoustique

- Biomédical
- Géophysique
- Traitement de la parole

Et bien d'autres applications encore, tout ceci montre l'importance de la TFCT dans le domaine du temps-fréquence.

Bien que le paramètre prédominant de l'algorithme reste la taille de la fenêtre d'analyse, le choix de sa forme n'est pas sans conséquence sur le traitement. Par exemple la fenêtre rectangulaire provoque des distorsions. En effet tout signal réel est à durée limitée, en d'autres termes on peut dire que tout signal $s(t)$ est tel que :

$$s_{reel}(t) = s(t)\Pi_T\left(t - \frac{T}{2}\right) \quad \text{I-10}$$

$$\text{avec } \Pi_T(t) = \begin{cases} 1 & \text{si } \frac{-T}{2} \leq t \leq \frac{T}{2} \\ 0 & \text{ailleurs} \end{cases} \quad \text{et} \quad \text{sinc}(t) = \frac{\sin(\pi t)}{\pi t}$$

On introduit un retard dans l'expression de $\Pi_T(t)$ car tout signal réel est causal. Si on passe dans le domaine de Fourier on obtient donc :

$$S_{reel}(f) = S(f) * Tsinc(fT)e^{-i\pi fT} \quad \text{I-11}$$

Le spectre est donc modifié car il y a apparition d'un sinus cardinal, qui selon les applications telles que l'observation des signaux à bande-étroite peut être nuisible à l'interprétation de l'analyse effectuée. Chaque fenêtre a des caractéristiques qui lui sont propres, elles peuvent être adaptées à certaines applications. Dans [14] un tableau récapitulatif : il décrit quelle fenêtre est recommandée suivant le type de signal auquel nous sommes confrontés. Le but étant de trouver le meilleur compromis entre la largeur du lobe central et la présence de lobes secondaires. Nous allons maintenant étudier la reconstruction d'un signal $s(t)$ à partir de sa TFCT, ce qui sera très utile dans les étapes concernant la réduction du bruit.

En annexe A, une mise en œuvre de la reconstruction du signal à partir de la TFCT est présenté. Cette approche sera utilisée pour reconstruire le temporel à partir d'une TFCT où un traitement aura permis de réduire le bruit présent au sein de l'observation.

3. Expérimentations et résultats

Considérons un signal $s_1(t)$ représentant une fréquence modulée linéairement, ce type de signal est appelé chirp. Nous travaillons sur la fréquence normalisée, nous considérons donc la fréquence d'échantillonnage $F_e = 1 \text{ Hz}$. La fréquence centrale de ce signal est 0.1 Hz et sa largeur de bande est de 0.1 Hz . La durée du signal est de $T = 256 \text{ s}$. Le plan temps-fréquence obtenu après spectrogramme est présenté sur la Figure 7. Ce résultat confirme que la résolution du spectrogramme est mauvaise, il en résulte un étalement du signal utile. Il est à noter que les plans

temps-fréquence présentés dans cette partie ont été réalisés à l'aide de la boîte à outils temps-fréquence développé par François Auger, Olivier Lemoine, Paulo Gonçalves et Patrick Flandrin.

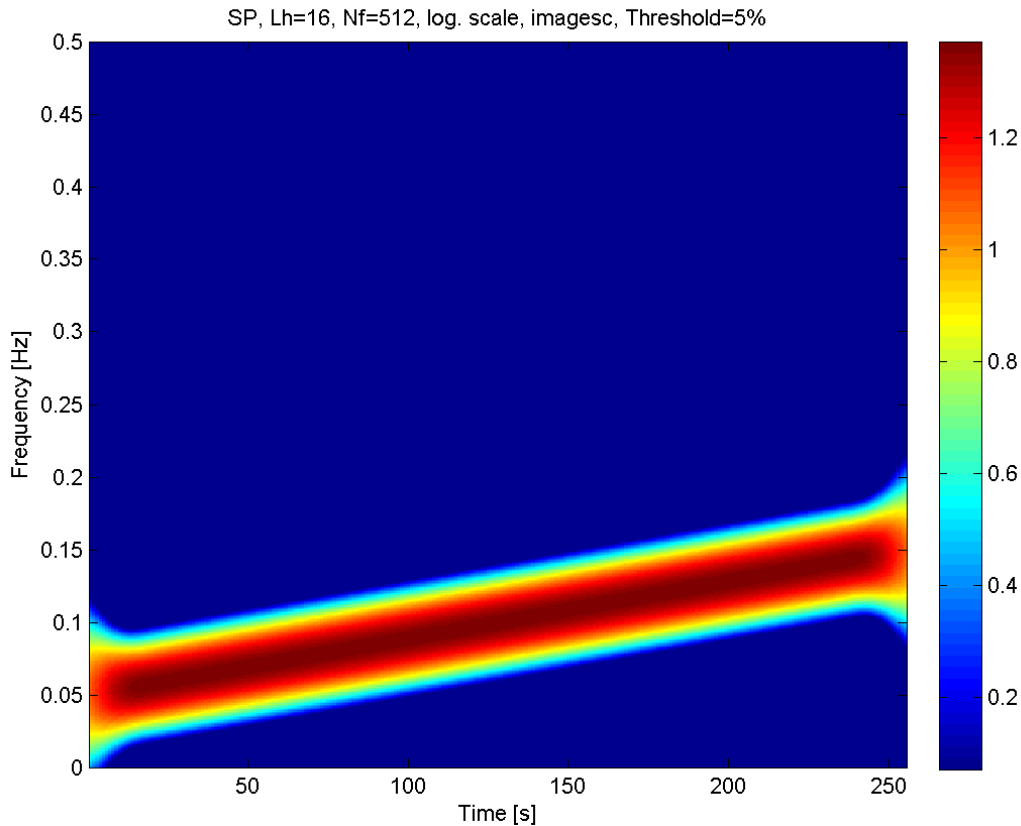


Figure 7: Spectrogramme d'un chirp linéaire, analyse avec une fenêtre de Hamming de 16 s

Considérons maintenant un signal $s(t)$ tel que :

$$s(t) = \begin{cases} s_0(t), & \forall t \in [0; T] \cup [2T; 3T] \\ s_0(t) + s_1(t) + s_2(t), & \forall t \in]T; 2T[\end{cases}$$

où $s_0(t)$ est un signal monochromatique de fréquence 0.22 Hz, $s_1(t)$ correspond au chirp linéaire étudié précédemment et $s_2(t)$ est un chirp linéaire (fréquence centrale 0.35Hz, de largeur de bande 0.1 Hz et de durée T) et où T est égal à 256 s. Le spectrogramme obtenu est présente sur la Figure 8. Nous pouvons noter qu'encore une fois le plan temps-fréquence obtenu n'a pas une résolution optimale, cependant la lisibilité et l'interprétation peuvent être réalisées de façon correcte. Si les composantes était plus proche la faible résolution pourrait entraîner une fusion des différentes composantes de ce signal ce qui deviendrait dérangeant. Nous pouvons voir ce phénomène sur la Figure 9. Cette dernière expérimentation confirme le rôle prépondérant de la taille de la fenêtre de traitement dans le calcul du spectrogramme et le principe d'incertitude d'Heisenberg est vérifié.

La transformation de Wigner-Ville permet de résoudre certains problèmes qu'a le spectrogramme comme nous allons le voir dans le prochain chapitre.

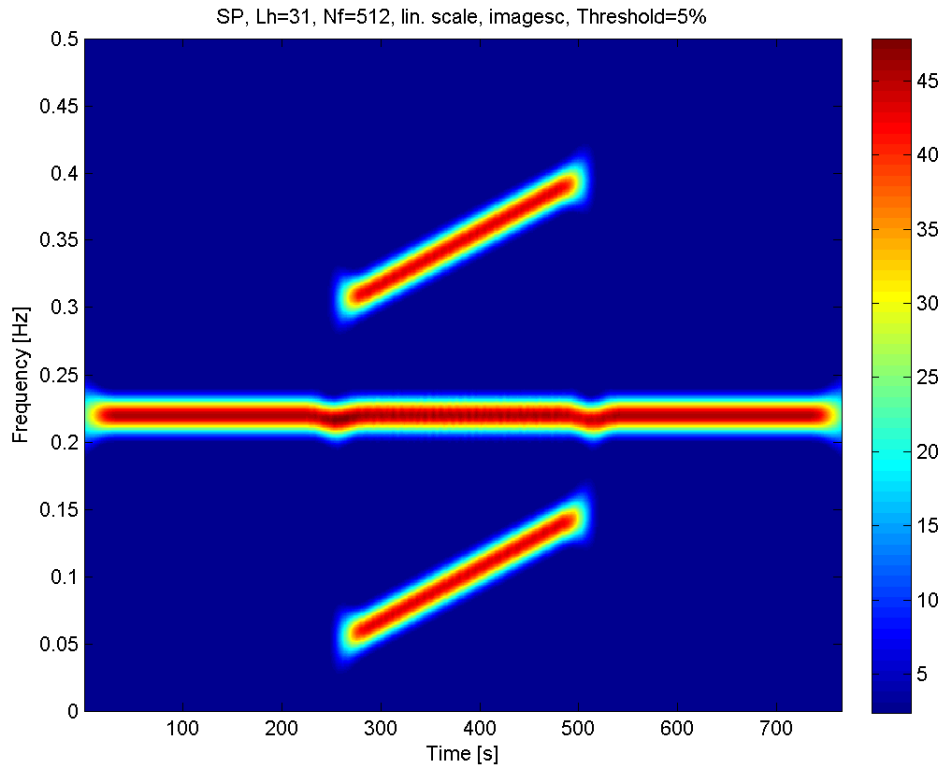


Figure 8: Spectrogramme d'un signal composé de deux chirps linéaires et d'un signal monochromatique, avec une fenêtre de Hamming de 63 échantillons.

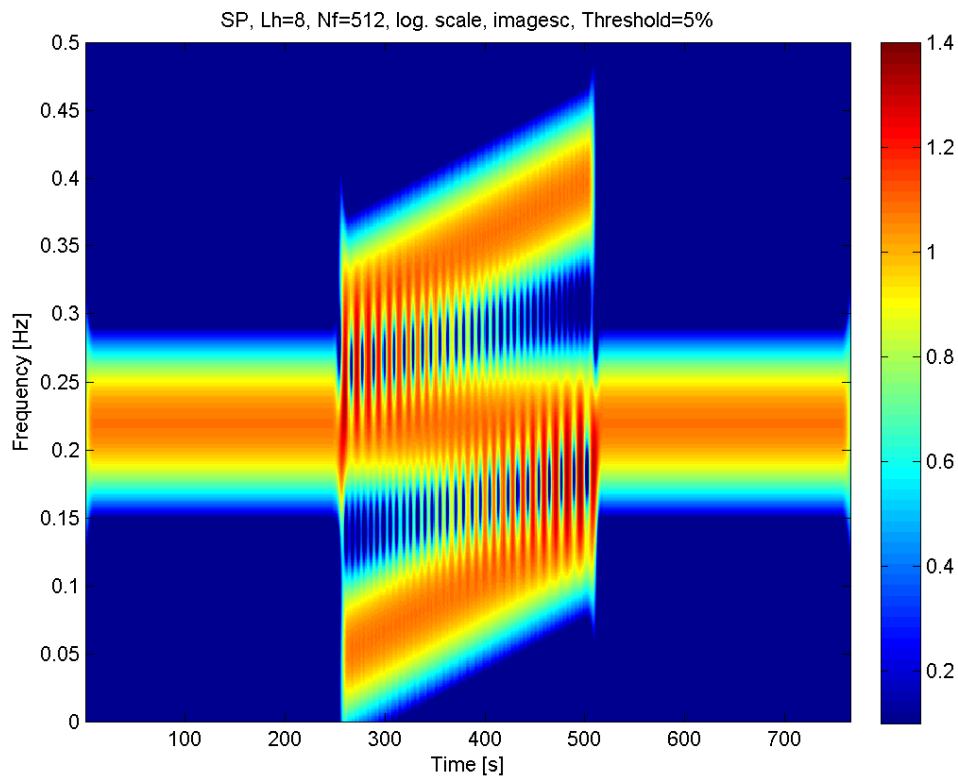


Figure 9: Spectrogramme d'un signal composé de deux chirps linéaires et d'un signal monochromatique, avec une fenêtre de Hamming de 17 s

C. Transformée de Wigner-Ville et ses dérivées

1. Principe

C'est en 1948 que J. Ville [15] propose cette distribution déjà utilisée par Wigner [16] dans le cadre de la mécanique statistique. L'idée générale de cette méthode est de mesurer la symétrie locale d'un signal autour d'un point temps-fréquence. La distribution de Wigner-Ville (WV) est définie par :

$$WV_s(t, f) = \int_{-\infty}^{+\infty} s\left(t + \frac{\tau}{2}\right) s^*\left(t - \frac{\tau}{2}\right) e^{-2i\pi f\tau} d\tau \quad \text{I-12}$$

Elle ne nécessite pas l'introduction d'une fenêtre extrinsèque au signal comme dans le cas du spectrogramme. Toutefois, nous pouvons faire le parallèle entre ces deux RTF, car nous pouvons considérer que la WV est une TFCT dans laquelle la fenêtre est adaptée au signal. De ce fait, les supports temporel et fréquentiel du signal analysé sont conservés, l'énergie du signal l'est aussi.

En discret, cette transformation peut se réécrire :

$$WV_s(t, f) = 2 \sum_{\tau} s_N(t + \tau) s_N^*(t - \tau) e^{-4i\pi f\tau} \quad \text{I-13}$$

où s_N représente le signal à temps discret obtenu par échantillonnage à une période pris comme unité. Dans ces conditions, le plan temps-fréquence est décrit par une fonction périodique de période $\frac{1}{2}$ en f , alors que cette période est unitaire dans le cas du spectre du signal échantillonné, ceci entraîne de fait un non-respect du critère de Shannon-Nyquist, c'est-à-dire un repliement de spectre pour des fréquences supérieures à $\frac{1}{4}$ [17]. Pour remédier à ce problème, soit il faut sur-échantillonner le signal d'un facteur supérieur ou égal à 2 (ce qui engendrera une augmentation de la quantité de calculs à traiter et, de fait, une difficile adéquation avec la tenue du temps réel), soit il faut construire la transformée de Wigner-Ville discrète sur le signal analytique associé au signal réel échantillonné normalement. Dans le deuxième cas, le repliement de spectre des seules fréquences positives n'affectera que les fréquences négatives pour lesquelles la contribution spectrale est nulle car par définition le signal analytique ne présente que des fréquences positives. Notons, que là encore, le calcul de cette transformée implique de prendre en compte deux fois plus d'échantillons que N , c'est-à-dire N échantillons réels et N échantillons imaginaires.

La WV possède de nombreuses propriétés [18], parmi lesquelles :

- elle représente le « spectre instantané » satisfaisant aux contraintes de distribution marginale temporelle et fréquentielle,
- elle satisfait la translation temporelle et fréquentielle, ainsi que le changement d'échelle.
- Aucune hypothèse de stationnarité locale n'est faite, donc pas besoin de fenêtres d'analyse.
- Elle permet d'offrir une bonne localisation des structures énergétiques dans le plan temps-fréquence,
- Elle est parfaitement adaptée aux modulations linéaires dans la mesure où elle concentre l'énergie le long de la fréquence instantanée.

Cependant cette méthode a des inconvénients non négligeables tels que :

- Des valeurs négatives peuvent être présentes ceci est dû à l'incompatibilité entre positivité et bilinéarité assortie de marginales correctes.
- La lisibilité réduite de par la présence de termes d'interférences pouvant être importants.

Lorsque nous sommes en présence de signaux réels, ce facteur est rédhibitoire. En effet, la WV fournit d'excellents résultats pour des signaux mono-composante mais pour les signaux à composantes multiples, elle présente des interférences indésirables. Malheureusement les signaux réels sont souvent à composantes multiples. Pour illustrer ce phénomène, supposons que $s_1(t)$ et $s_2(t)$ soient deux composantes d'un seul signal $s(t)$. La WV est alors :

$$WV_{s_1+s_2}(t, f) = WV_{s_1}(t, f) + WV_{s_2}(t, f) + 2\Re(WV_{s_1, s_2}(t, f)) \quad \text{I-14}$$

Avec :

$$WV_{s_1, s_2}(t, f) = \int s_1\left(t + \frac{\tau}{2}\right) s_2^*\left(t - \frac{\tau}{2}\right) e^{-2i\pi f \tau} d\tau \quad \text{I-15}$$

On peut montrer que les structures interférentielles sont placées à mi-distance des composantes s_1 et s_2 et oscillent suivant l'axe des temps et/ou des fréquences [19]. La règle du point milieu résume la contribution des interférences : deux points du plan interfèrent pour créer une contribution en un troisième, leur milieu géométrique.

Cependant dans des applications réelles nous utilisons la transformée de Wigner-Ville sur le signal analytique $s_A(t)$ associée au signal $s(t)$ [20]. On appelle signal analytique un signal dont la transformée de Fourier est causale, c'est-à-dire un signal dont la transformée de Fourier est nulle pour les fréquences négatives, on a :

$$s_A(t) = s(t) + jH(s(t)) \quad \text{I-16}$$

avec $H(x(t))$ étant la transformée d'Hilbert du signal, elle est définie ainsi :

$$H(s(t)) = \frac{1}{\pi} vp \left\{ \int_{-\infty}^{\infty} \frac{s(\tau)}{t - \tau} d\tau \right\} \quad \text{I-17}$$

vp étant l'abréviation de valeur principale de Cauchy.

Travailler sur $s_A(t)$ nous permet de supprimer les fréquences négatives, ceci réduira une partie des interférences. Il est à noter que même un signal réel mono-composante interférera toujours avec les atomes temps-fréquence décrivant le bruit.

2. Pseudo Wigner-Ville lissée

Précédemment on a vu que la structure de la transformation de Wigner-Ville introduit des termes d'interférences entre composantes. Ces termes possédant une structure oscillante, il est généralement nécessaire de travailler avec une version lissée de la représentation. Une partie des interférences subsistantes peut être supprimée en utilisant une fenêtre d'observation temporelle $h(t)$, à valeur réelles, qui isolera les composantes fréquentielles non simultanées [21]. La nouvelle représentation temps-fréquence ainsi obtenue est la pseudo-distribution de Wigner-Ville (PDWV) [22]. Si elle est associée au signal analytique $s_A(t)$ et non au signal $s(t)$ lui-même, elle a pour expression:

$$PDWV_{s_A}(t, f) = h(t) *_t WV_{s_A}(t, f) \quad \text{I-18}$$

Malgré ce lissage des interférences entre composantes de fréquence positive viennent encore nuire à la lisibilité de la RTF. Un lissage fréquentiel permet de diminuer leur amplitude. La distribution associée est la pseudo-distribution de Wigner-Ville lissée (PDWVL). Son expression est :

$$PDWVL_{s_A}(t, f) = \gamma(t, f) **_{tf} WV_{s_A}(t, f) \quad \text{I-19}$$

Dans cette expression, $\gamma(t, f)$ est un noyau de lissage constant. Plus généralement toute les transformées temps-fréquence peuvent-être décrite par l'équation I-19, comme nous le constatons dans [9].

3. Expériences et interprétations

A titre d'exemple, considérons le chirp linéaire $s_1(t)$ présenté précédemment. Le plan temps-fréquence obtenu par application de la transformée de Wigner-Ville est présenté à la Figure 10. L'analyse de ce résultat montre sans équivoque que cette transformée est parfaitement adaptée à l'étude de signaux modulés linéairement en fréquence. En revanche, nous voyons déjà l'apparition de quelques interférences.

Pour mettre en évidence ce phénomène d'interférence, considérons le signal $s(t)$ décrit précédemment.

Le plan temps-fréquence obtenu est proposé sur la Figure 11. Les motifs temps-fréquence, dénotés par ①, ② et ③, correspondent respectivement aux signaux $s_0(t)$, $s_1(t)$ et $s_2(t)$; les trois autres correspondent à des interférences :

- ④ : interférences entre $s_0(t)$ et $s_2(t)$;
- ⑤ : interférences entre $s_1(t)$ et $s_2(t)$;
- ⑥ : interférences entre $s_0(t)$ et $s_1(t)$.

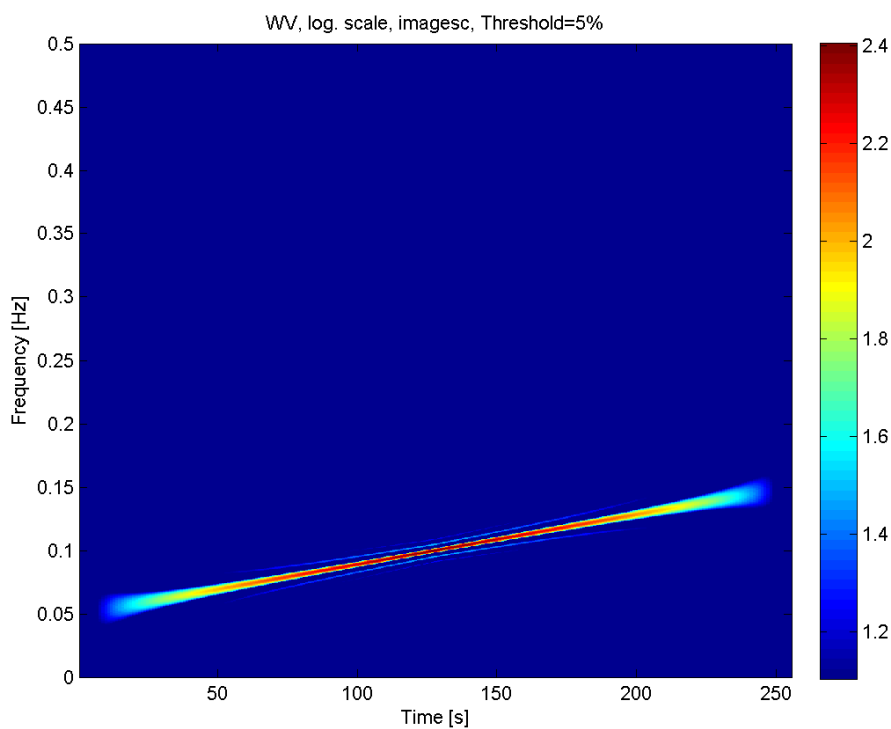


Figure 10: Plan temps-fréquence d'un chirp linéaire obtenu par transformée de Wigner-Ville

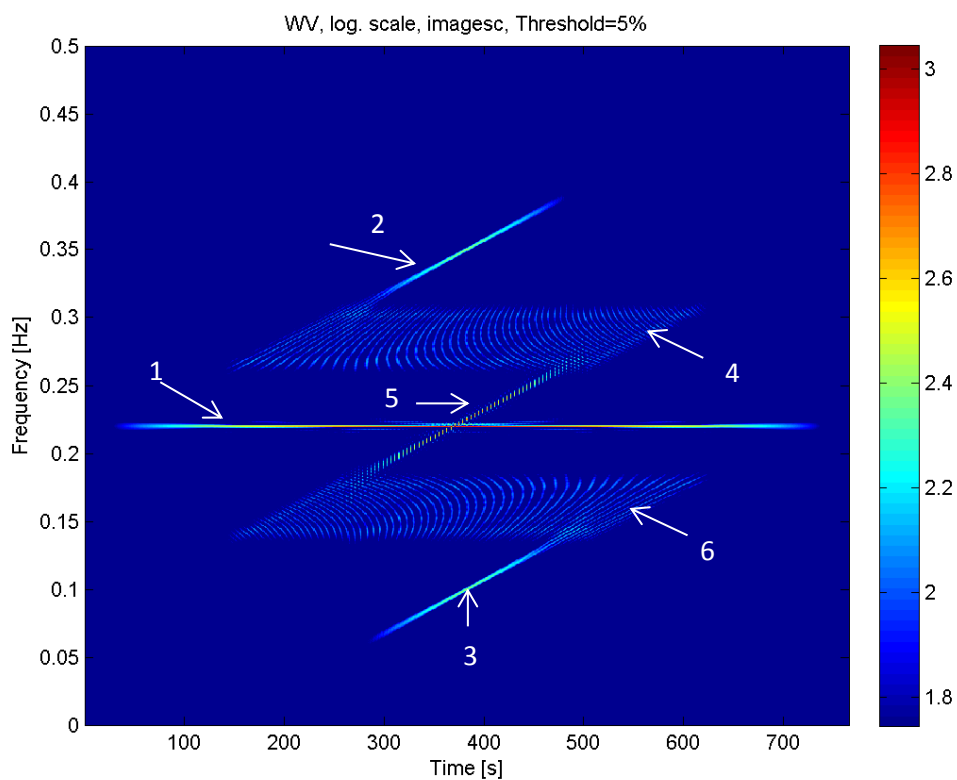


Figure 11: Wigner-Ville d'un signal constitué d'un signal monochromatique et de deux chirps linéaires localisé en temps

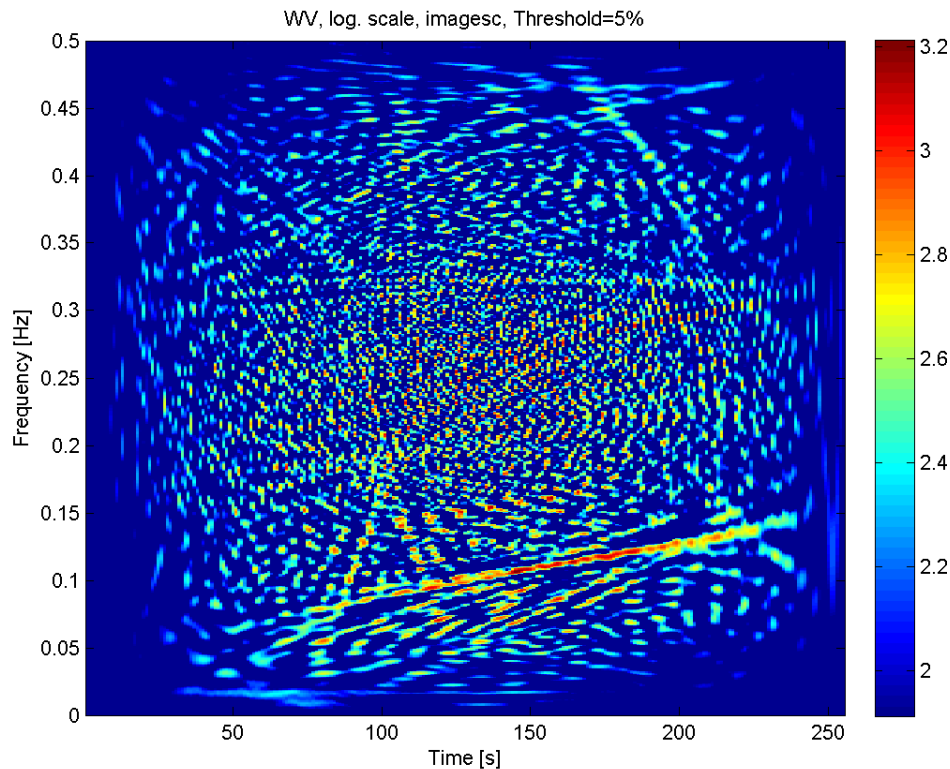


Figure 12: Plan temps-fréquence obtenu par Wigner-Ville d'un chirp en présence de bruit blanc Gaussien (RSB=-3dB)

De plus, lorsque le signal d'intérêt est entaché de bruit, des interférences vont apparaître entre le signal et le bruit, mais également entre les différentes fréquences liées au bruit ; ces dernières peuvent entraîner une destruction du motif temps-fréquence associé au signal d'intérêt et ainsi engendrer une dégradation du rapport signal à bruit dans le plan temps-fréquence. Considérons, par exemple, le cas d'une observation correspondant à la superposition du signal $s_1(t)$ et d'un bruit blanc Gaussien, tel que le rapport signal à bruit (RSB) soit -3dB. Le plan temps-fréquence obtenu est présenté en Figure 12.

Il apparaît qu'il est difficile de percevoir la présence du chirp. Pour prendre conscience de l'influence des interférences liées au bruit, considérons, à présent, le plan temps-fréquence sur la Figure 13 qui correspond à la superposition additive des plans temps-fréquence obtenues par WV du signal d'intérêt seul et du bruit seul.

Sur ce dernier le chirp linéaire apparaît plus clairement que précédemment. Ainsi, une comparaison entre ces deux plans temps-fréquence met en exergue l'influence néfaste des interférences sur le RSB.

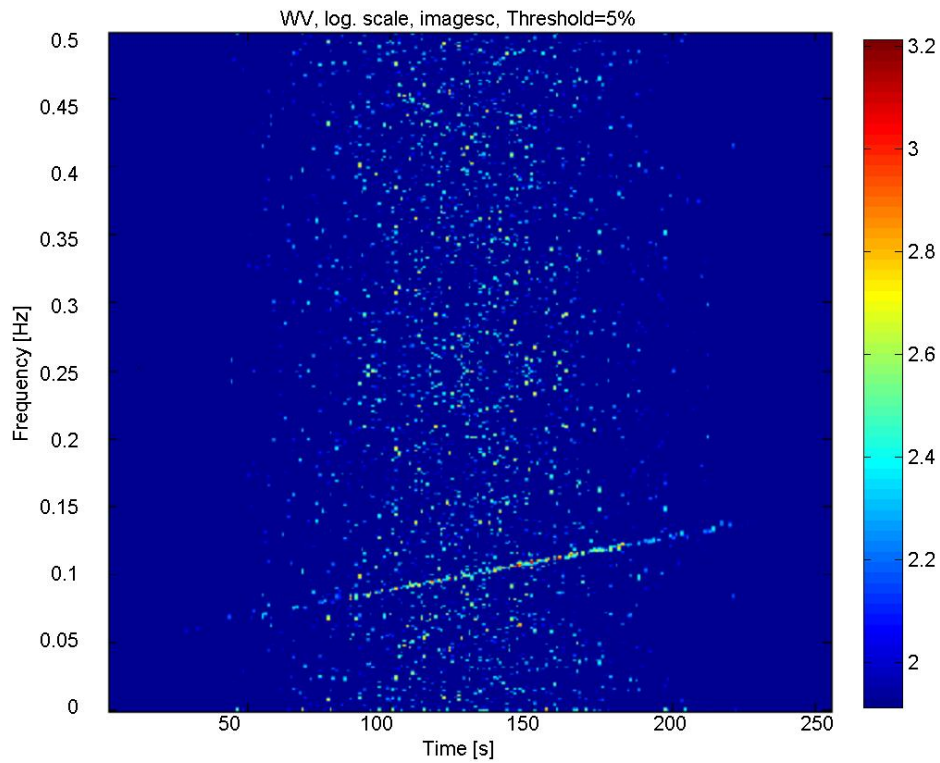


Figure 13: Plan temps-fréquence obtenu par Wigner Ville de l'observation bruitée en en tenant pas compte des interférences entre le signal et le bruit

Enfin, afin de quantifier l'apport du lissage sur la réduction des interférences, considérons, de nouveau, le signal constitué de la superposition de deux chirps et d'un signal monochromatique. Le plan temps-fréquence obtenu est proposé sur la Figure 14.

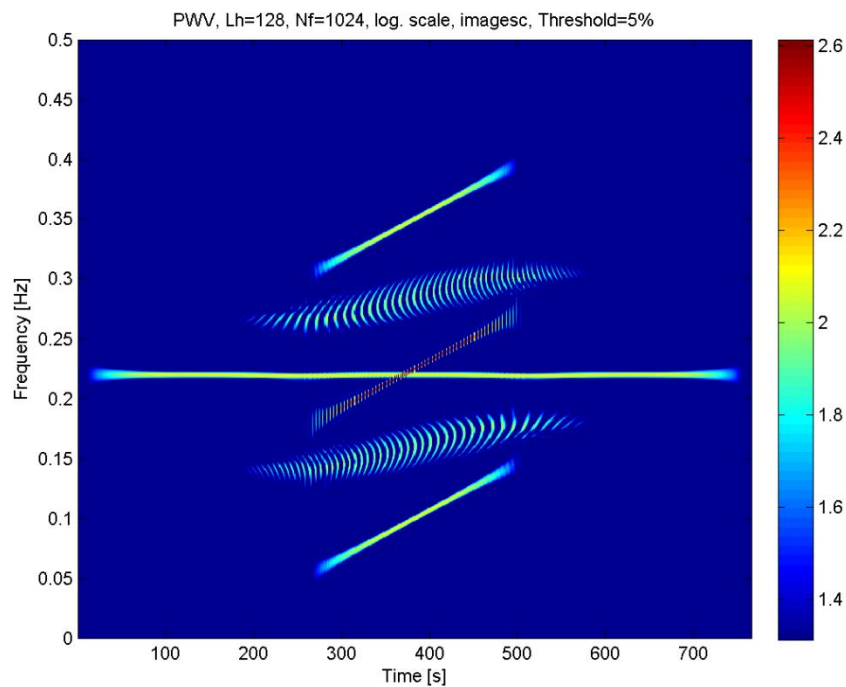


Figure 14: Pseudo Wigner-Viller d'un signal constitué d'un signal monochromatique et de deux chirps linéaires localisé en temps

L'analyse de cette figure révèle que l'utilisation de la pseudo Wigner-Ville permet de réduire l'influence des interférences grâce à un lissage en fréquence, tout en ne les éliminant pas totalement. On notera en particulier que la disparition des interférences de moindre amplitude entre les signatures temps-fréquences réelles (notées précédemment ①, ② et ③) et les interférences principales (dénotées ④, ⑤ et ⑥).

Enfin, afin de quantifier l'apport de la transformée pseudo Wigner-Ville lissée, étudions le plan temps-fréquence obtenu avec les signaux étudiés précédemment. Le plan obtenu dans le cas de la superposition additive de deux chirps et d'un signal monochromatique est présenté ci-après. L'analyse de ce dernier montre clairement que l'objectif est atteint, les interférences ayant totalement disparues, en revanche ceci se fait au détriment de la résolution. En effet, d'une part le lissage temporel affecte la résolution temporelle, c'est-à-dire la capacité à distinguer deux impulsions successives et, d'autre part, le lissage fréquentiel s'effectue aux dépens de la résolution fréquentielle.

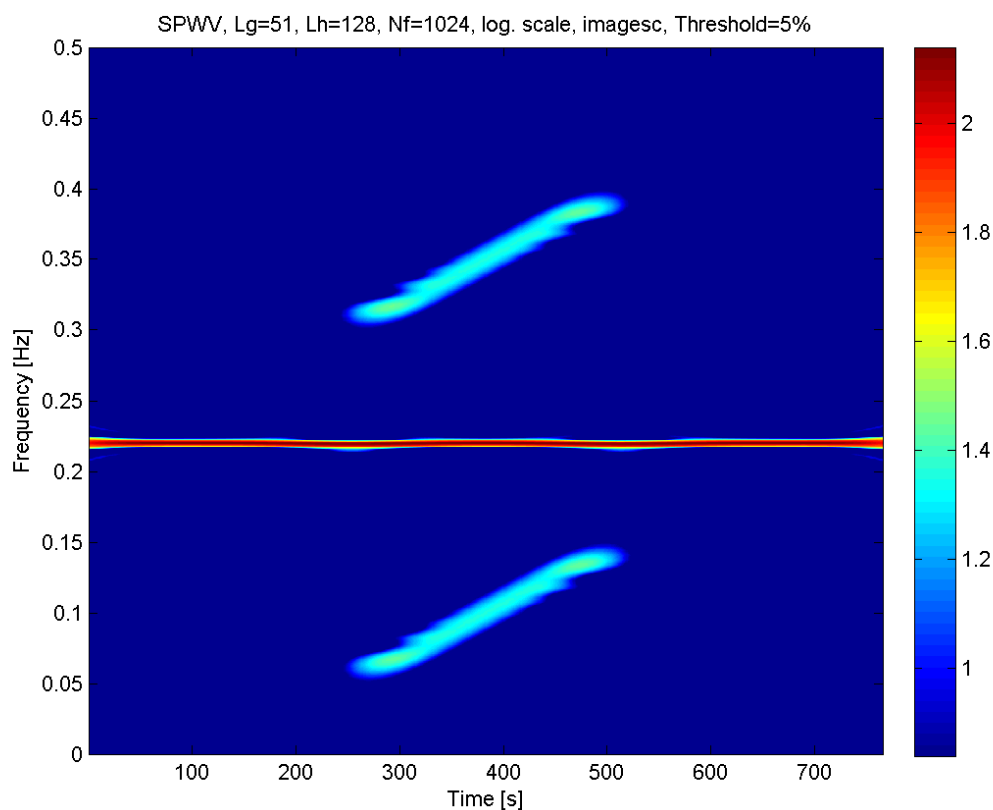


Figure 15: Pseudo Wigner-Ville lissée d'un signal constitué d'un signal monochromatique et de deux chirps localisés en temps

Après avoir vu différentes transformations temps-fréquence, intéressons-nous à un post-traitement qui permet d'améliorer la résolution après traitement.

D. Réallocation spectrale

Nous avons vu que la RTF était devenu un élément incontournable du traitement du signal. Cependant chaque RTF a ses faiblesses et nous nous rendons compte qu'il est impossible d'avoir une représentation « idéale ». Nous devons réaliser un compromis entre précision de la localisation des zones contenant du signal et lisibilité du temps-fréquence. La RTF est généralement un moyen d'analyse préalable à des traitements avals tels que la localisation ou l'identification de motifs temps-fréquence. L'efficacité de ces traitements dépend fortement de la qualité de la représentation, c'est ainsi qu'est née l'idée de réaliser des post-traitements succédant à la phase d'analyse et visant à en affiner la représentation: la réallocation est l'un de ces post-traitements. Elle a été introduite par [23] en 1976 et ce n'est que récemment, dans [24], que les auteurs ont montré la pertinence de cette méthode en tant qu'outil complémentaire pour l'analyse temps-fréquence. Le principe consiste à déplacer les valeurs des représentations temps-fréquence afin d'améliorer la résolution fréquentielle et temporelle de la RTF et donc améliorer la lisibilité du plan temps-fréquence et par conséquent aussi les traitements avals.

L'idée de la réallocation est la suivante : si nous nous plaçons dans le domaine de la mécanique, nous pouvons affecter la masse totale d'un objet au centre géométrique de ce dernier. Ce choix n'est pas forcément le meilleur car dans ce cas on considérerait l'objet comme ayant une distribution uniforme, ce qui n'est pas le cas. Un choix plus judicieux serait d'associer la masse de l'objet à son centre de gravité. C'est exactement l'idée de la réallocation. Elle a pour but de recentrer à leurs vraies places, l'énergie des termes mono-composantes du signal étalés par l'opération de lissage. En effet, en remarquant que le spectrogramme peut s'écrire comme une transformée de Wigner-Ville lissée [25] [9]:

$$S_x^h(t, f) = \iint WV_x(\tau, \nu) WV_h(\tau - t, \nu - f) d\tau d\nu \quad \text{I-20}$$

Où $WV_x(t, f)$ est la transformée de Wigner-Ville du signal, et $WV_h(t, f)$ le noyau de lissage égal à la distribution de Wigner-Ville de la fenêtre $h(t)$. Ce lissage va donc étaler la répartition de l'énergie issue de la distribution de Wigner-Ville. Le principe de la réallocation est de « refocaliser » le spectrogramme. On déplace l'énergie du point (t, f) sur un nouveau point (t', f') , centre de gravité de la distribution de Wigner-Ville du signal dans un voisinage, fonction de la fenêtre d'analyse $WV_h(t, f)$, [26]:

$$t'(t, f) = \frac{1}{S_x^h(t, f)} \iint t W_x(\tau, \nu) W_h(\tau - t, \nu - f) d\tau d\nu \quad \text{I-21}$$

$$f'(t, f) = \frac{1}{S_x^h(t, f)} \iint f W_x(\tau, \nu) W_h(\tau - t, \nu - f) d\tau d\nu \quad \text{I-22}$$

En pratique la réallocation peut être vue comme un processus agissant en deux étapes:

- Un lissage, qui nous permet d'enlever les termes d'interférence, au détriment de la bonne localisation des composantes du signal.
- Une compression des composantes qui subsistent après le lissage.

Représentations temps-fréquence

Cette méthode donne des résultats très satisfaisants, comme l'atteste la Figure 16, en effet elle permet de retrouver une résolution proche de la transformée de Wigner-Ville, en ayant l'avantage de ne plus présenter d'interférences. Cependant cette méthode est sensible au bruit [25], pour l'utiliser sur signaux réels, il faut donc réaliser un traitement permettant d'augmenter le rapport signal à bruit, de plus cette méthode est coûteuse en termes de temps de calcul, même si une méthode de réallocation rapide a été codée [26].

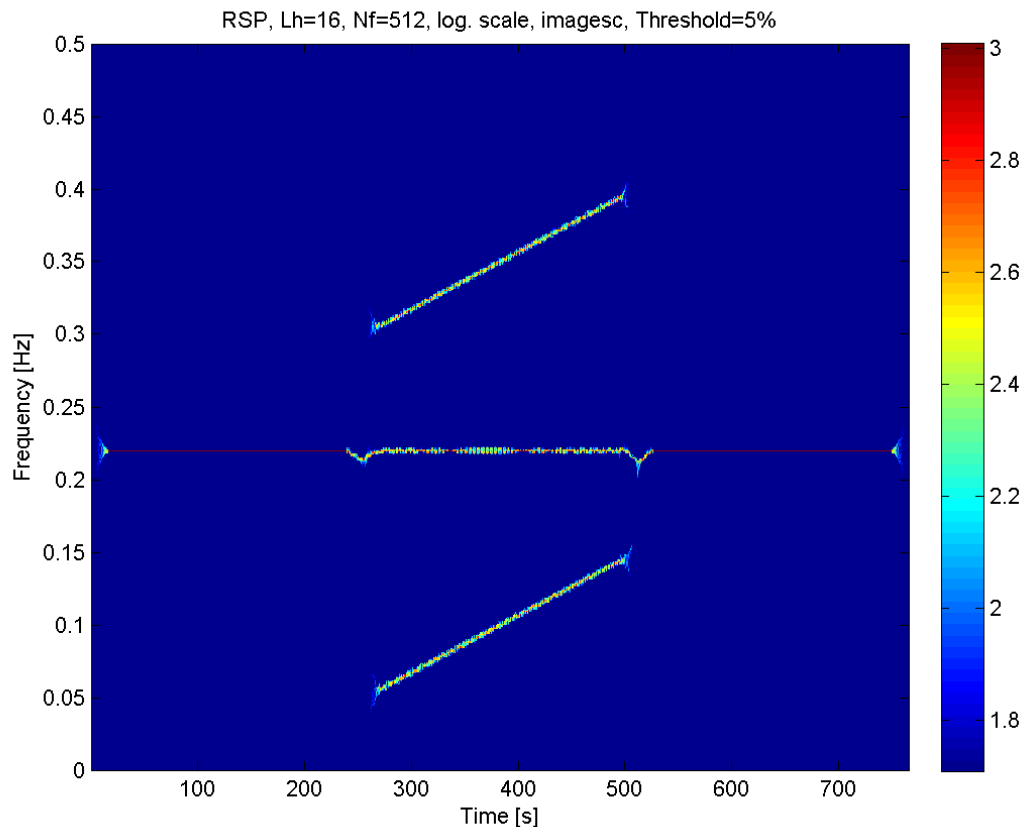


Figure 16: Spectrogramme réalloué d'un signal composé de deux chirps linéaires et d'un signal monochromatique.

Après avoir réalisé un état de l'art des représentations temps-fréquences classiques intéressons-nous aux transformées temps-échelle et plus particulièrement à l'analyse multi-résolution.

E. Transformée en ondelettes

1. Introduction

Les deux transformées étudiées précédemment utilisent une fenêtre de taille fixe couvrant le domaine spatio-temporel (Figure 17). Or dans certains cas on souhaiterait avoir une fenêtre qui s'adapte en fonction des irrégularités du signal, c'est-à-dire les variations brusques dans le signal ou hautes fréquences (Figure 18). C'est pour cette raison qu'a été créée la transformée en ondelettes.

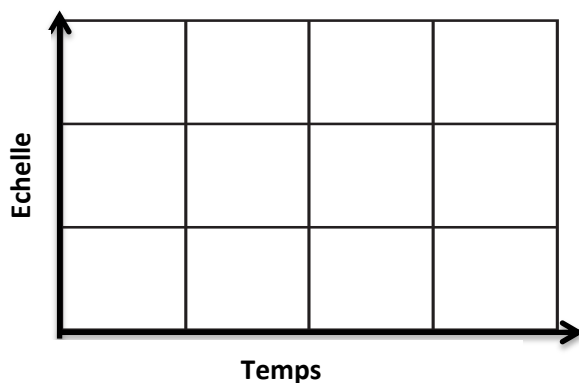


Figure 17: Pavage temps-fréquence classique

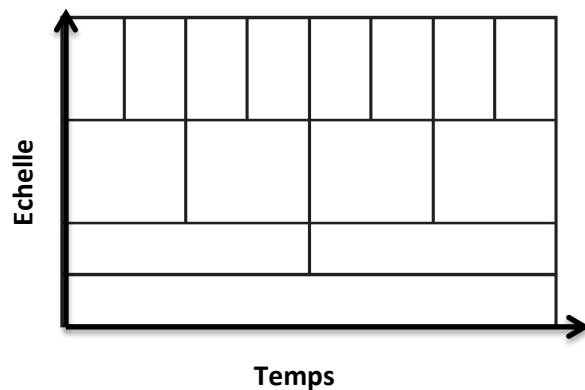


Figure 18: Pavage temps-fréquence en utilisant les ondelettes

Les ondelettes ont été introduites au début des années 1980 par Morlet et Grossmann [27]. Il s'agit d'une décomposition du signal sur des fonctions élémentaires. Le pavage temps-fréquence est maintenant remplacé par le pavage temps-échelle. L'ondelette est à la fois translatée et dilatée ou contractée, c'est pour cela que le pavage temps-échelle n'est pas régulier.

Les ondelettes trouvent une place incontournable dans les problèmes de traitement des signaux et des images. Elles sont utilisées notamment pour l'analyse, la compression et la réduction du niveau de bruit des signaux, on parle alors d'analyse multi-résolution [28].

2. Définition et propriétés

La transformée en ondelettes est une représentation temps-échelle du signal. Elle transcrit un filtrage du signal avec une fenêtre (ondelette) qui ne reste pas constante pour toutes les fréquences comme c'est le cas pour la transformée de Fourier à court-terme. L'ondelette adapte son support en fonction du paramètre d'échelle. C'est une onde qui est localisée dans le temps, en quelque sorte c'est une « petite » onde, d'où le nom d'ondelette. Il s'ensuit que la condition d'admissibilité de l'ondelette est qu'elle soit à moyenne nulle.

La famille d'ondelettes est obtenue à partir d'une ondelette mère $\Psi(t)$ grâce à la relation suivante:

$$\Psi_{a,b}(t) = \frac{1}{\sqrt{a}} \Psi\left(\frac{t-b}{a}\right)$$

I-23

Ainsi nous réalisons des opérations de translation en temps grâce au paramètre b et dilatation (ou contraction) grâce au paramètre d'échelle a . La transformée en ondelettes continue (CWT) a été développée comme une alternative à la transformée de Fourier à court terme pour résoudre le problème posé par les résolutions temporelle et fréquentielle. Par son principe, l'analyse par ondelettes est semblable à l'analyse TFCT, mais elle se distingue de cette dernière sur deux points :

- Contrairement à la transformée de Fourier, il n'y a qu'une seule raie visible pour représenter une sinusoïde, c'est-à-dire que les fréquences négatives ne sont pas considérées;
- La largeur du support temporel de la fenêtre varie tout au long de la transformation.

La CWT est définie de la façon suivante :

I-24

$$CWT_x(a, b) = \langle x, \Psi_x \rangle = \frac{1}{\sqrt{a}} \int_{-\infty}^{+\infty} x(t) \Psi^* \left(\frac{t-b}{a} \right) dt$$

où $\langle . \rangle$ désigne le produit scalaire.

Le signal ainsi transformé est donc une fonction de deux variables, a et b , représentatifs respectivement de l'échelle et de la translation. Le paramètre d'échelle pourrait être comparé à l'échelle utilisée en cartographie. Tout comme cette dernière, une échelle importante correspond à des zones homogènes, présentant peu de détails, elles correspondent aux basses fréquences et ont généralement une durée importante, contrairement à une échelle de faible valeur qui correspond à des zones de détails qui sont les hautes fréquences qui ont généralement une courte durée. En effet, plus le facteur d'échelle a est important, plus l'ondelette est étalée. Pour un facteur d'échelle assez grand, la représentation des coefficients d'ondelettes donne une représentation de la forme générale de la fonction, comme si on observait la fonction de loin. Par contre, un facteur d'échelle faible correspond à une représentation des singularités, comme si on regardait la fonction de près. Cette propriété de «microscope» est très utile pour l'étude de la régularité d'une fonction ou pour l'analyse des structures fractale, où une même structure se retrouve à des échelles différentes.

La transformée en ondelettes continue possède quelques propriétés telles que :

- Conservation de l'énergie du signal analysé
- Linéarité
- Invariance par translation
- Dilatation

Similairement à la TFCT, la transformée en ondelettes discrète (DWT) a pour expression :

I-25

$$DWT(m, n) = a_0^{\frac{m}{2}} \int_{-\infty}^{+\infty} s(t) \Psi(a_0^{-m} t - nb_0) dt$$

pour une discrétisation du plan temps-échelle $a = a_0^m$ et $b = b_0^n$. Le cas particulier de $a_0 = 2$ et $b_0 = 1$ désigne la transformée en ondelettes dyadique :

$$DWT(m, n) = \sqrt{2}^m \int_{-\infty}^{+\infty} s(t) \Psi\left(\frac{t}{2^m} - n\right) dt$$

3. Analyse multi-résolution

Le concept de la multi-résolution est d'utiliser plusieurs fois la transformation en ondelettes afin de décomposer le signal de départ en plusieurs signaux (une échelle de signal) contenant des informations différentes. D'une manière très simpliste, à la première échelle, la décomposition en ondelettes va extraire les détails les plus fins du signal, puis à la seconde échelle, apparaîtront les détails un peu plus grossiers, et ainsi de suite, jusqu'à obtention d'un signal complètement lissé donc basse fréquence.

Ainsi l'analyse multi-résolution est en quelque sorte la formalisation mathématique du phénomène suivant: lorsqu'on observe un objet, suivant la distance à laquelle on se trouve de ce dernier, on percevra plus ou moins de détails, pourtant l'objet n'a pas changé. On peut ainsi dire que l'espace dans lequel il est représenté n'est pas le même suivant la distance d'observation, d'où cette présence ou absence de détails.

Les niveaux grossiers sont décrits par l'espace d'approximation qui est constitué de sous-espaces V_j , imbriqués, et qui sont associés à un facteur d'échelle. Le signal $s(t)$ d'une part et le signal $s\left(\frac{t}{2}\right)$ correspondant à une dilatation de facteur 2 d'autre part, appartiennent respectivement à V_j et V_{j+1} . Le passage de l'espace V_j vers l'espace V_{j+1} correspond à un zoom sur le signal. L'ensemble des sous-espaces $V_{j+1} \subset V_j \subset \dots \subset V_0 \subset V_{-1}$ recouvrent complètement l'espace du signal.

En pratique les coefficients issus de l'analyse multi-résolution sont estimés à l'aide des techniques de filtrage numérique, il y a deux algorithmes qui nous permettent de réaliser l'analyse multi-résolution qui sont l'algorithme de Mallat [28] et l'algorithme à trous [29].

4. Avantages et inconvénients

Les principaux avantages de l'utilisation des ondelettes sont :

- Grand choix possible d'ondelettes (ondelette de Morlet, de Daubechies, de Haar, chapeau mexicain etc...)
- L'analyse multi-résolution : ainsi on peut « zoomer » sur les zones du signal qui nous intéressent.

Les inconvénients majeurs des ondelettes sont :

- Le comportement des résolutions temporelles et fréquentielles,
- L'absence de critère de choix sur le type d'ondelette à utiliser.

II. Représentation basée sur l'audition humaine

A. Introduction

La qualité de représentation des signaux acoustiques est d'une importance capitale dans beaucoup de domaines. Etant donné que tous les signaux réels présentent des non-stationnarités, leur représentation doit refléter le plus fidèlement ces caractéristiques. Les êtres humains ont une remarquable capacité à analyser les sons, parfois même très bruités, ils sont capables de reconnaître très rapidement des signaux et de les assigner à des classes prédéfinies.

De nos jours, des chercheurs du monde entier tentent de développer des outils de reconnaissance aussi performants que l'être l'humain. Cette activité est due au potentiel économique des différentes applications que l'on pourrait faire grâce à la reconnaissance automatique. Nous pouvons citer comme exemple d'application :

- Les outils pour smartphone tel que Shazam® qui est capable de reconnaître une musique sur quelques mesures seulement. Siri® sur l'iPhone® qui est un outil de reconnaissance de mots.
- Robotique,
- Biomédical...

Des travaux ont montré que la perception acoustique humaine est fondée sur la perception de fréquences sonores [30]. Alors partant de l'hypothèse que les humains sont les meilleurs classificateurs de signal sonore, la contribution que nous allons proposer en cette fin de chapitre est de s'inspirer des mécanismes de la perception humaine pour l'identification d'un son, afin d'élaborer des outils de reconnaissance acoustique sous-marine.

B. L'oreille humaine

1. Fonctionnement de l'oreille humaine

On peut voir sur la Figure 19 que le son stimule le tympan (fine membrane), qui se met à vibrer. Le tympan transmet son mouvement aux osselets qui le répercutent à l'oreille interne par l'intermédiaire de la fenêtre ovale. Les vibrations se propagent dans un liquide enfermé dans la cochlée. La cochlée est tapissée de minuscules cils reliés au nerf auditif.

Oreille externe (pavillon -> conduit -> tympan): le pavillon recueille le signal auditif et le guide dans le conduit auditif comme le ferait un réflecteur, tout en favorisant les fréquences élevées (5KHz). Les dimensions et les parois du conduit en font un résonateur pour les fréquences voisines de 2 kHz qui sont justement les fréquences vocales. Le tympan vibre et transmet le mouvement aux organes qui constituent l'oreille moyenne, qui a pour fonction de réaliser une adaptation d'impédance et protection contre les bruits trop forts. Le signal arrive alors dans l'oreille interne, milieu liquidien où la cochlée le transforme en impulsions électriques et chimiques conduites par le nerf auditif, aux zones du cerveau concernées.

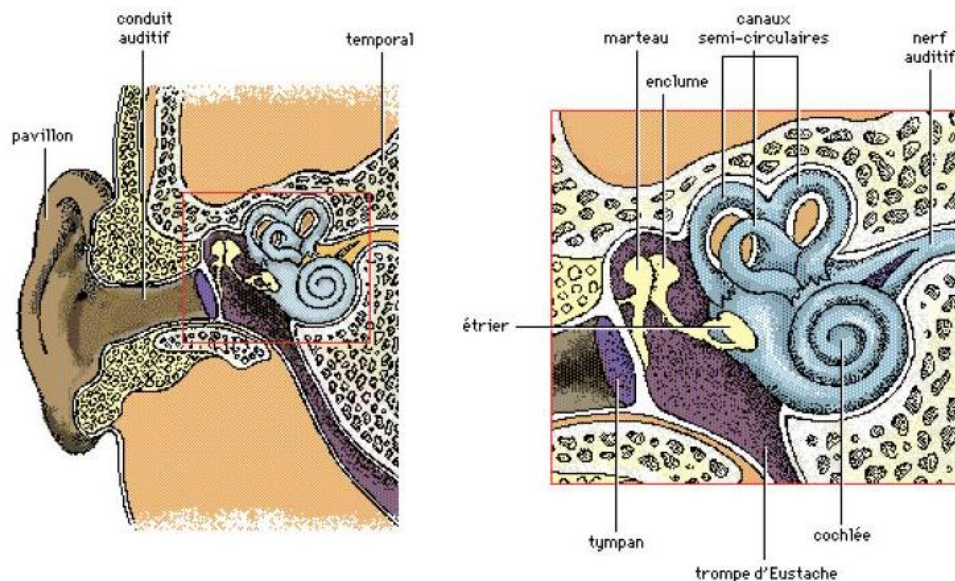


Figure 19: L'oreille humaine [31]

2. Sensibilité de l'oreille

L'oreille transforme les pressions acoustiques en sensation auditive. Elle ne perçoit pas de la même manière toutes les fréquences. L'ensemble des courbes de la Figure 20 est appelée diagramme de Fletcher et Munson. Il représente les courbes d'égale sensation sonore d'une oreille humaine normale en fonction de la fréquence. La zone d'audition normale est comprise entre la limite de la douleur (vers 120 décibels) et le seuil d'audition (0dB à 1000 Hz). Elle est en outre limitée vers 30 Hz pour les fréquences basses et vers 15000 Hz pour les fréquences hautes. Ce diagramme est une moyenne. Il montre que la sensibilité de l'oreille est maximale entre 1000 Hz et 5000 Hz. Les limites évoluent d'un sujet à l'autre et pour un même individu en fonction de l'âge ou des maladies et accidents. Entre 20 Hz et 15000 Hz, la sensation auditive produite par un son pur de niveau L varie en fonction de la fréquence. On a la même sensation pour un son de fréquence 1000 Hz à 40 dB, pour un son de fréquence 100 Hz à 65 dB et pour un son de fréquence 10000 Hz à 50 dB. Ces 3 sons ont un même niveau d'isophonie.

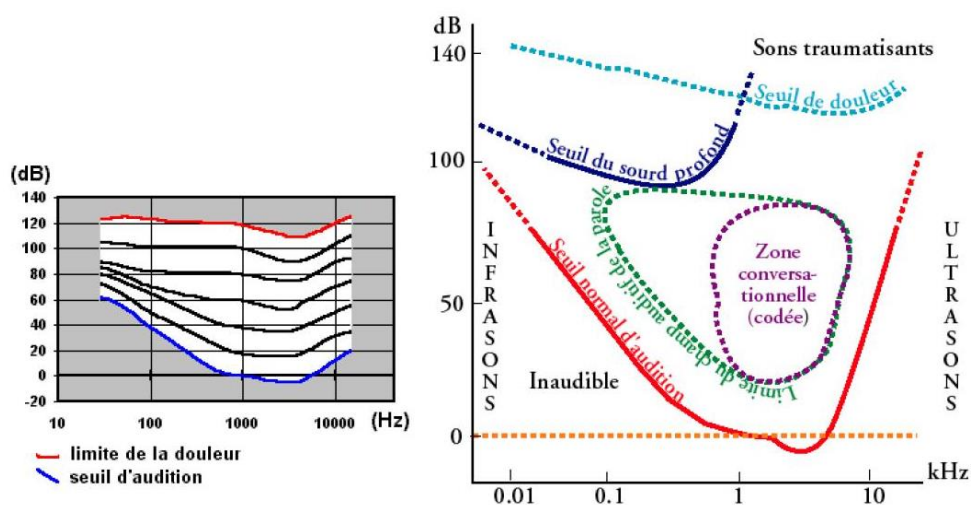


Figure 20: Diagramme de Fletcher et Munson (à gauche). Sons audibles de 20 à 20000 Hz (à droite) [31]

C. Analyse par bandes d'octave

Ce type d'analyse fut beaucoup utilisé par les analystes durant les dernières décennies. Elle permet de représenter l'énergie par bandes de fréquence réparties de façon logarithmique.

1. Principe

Une octave est l'intervalle entre deux fréquences f_1 et f_2 qui sont telles que :

$$f_2 = 2f_1 \quad \text{II-1}$$

La largeur de bande Δ_f n'est pas constante pour les différentes octaves étudiées en effet, on a :

$$\Delta_f = f_2 - f_1 = 2f_1 - f_1 = f_1 \quad \text{II-2}$$

La fréquence médiane de chaque octave est définie comme suit :

$$f_{med} = \sqrt{f_1 f_2} = \sqrt{2f_1 f_1} = \sqrt{2}f_1 \quad \text{II-3}$$

De plus l'analyse par bande d'octave est à $\frac{\Delta_f}{f_{med}}$ constant car :

$$\frac{\Delta_f}{f_{med}} = \frac{f_1}{f_1 \sqrt{2}} = \frac{1}{\sqrt{2}} \quad \text{II-4}$$

Il nous faut donc déterminer une fréquence f_1 pour définir la première octave. En général, la première fréquence est fixée à 20 Hz, car c'est la plus petite fréquence audible par l'oreille humaine.

Le tableau suivant donne les bandes d'octave de l'analyse normalisée, cette analyse est basée sur la fréquence 1000 Hz:

$f_1(\text{Hz})$	20	40	80	160	320	640	1280	2560	5120
$f_2(\text{Hz})$	40	80	160	320	640	1280	2560	5120	10240
$f_{med}(\text{Hz})$	31.5	63	125	250	500	1000	2000	4000	8000

Tableau 1: Bandes d'octave normalisées

2. Analyse par tiers d'octave

Un tiers d'octave correspond à un intervalle entre deux fréquences telles que :

$$f_2 = 2^{\frac{1}{3}}f_1 \quad \text{II-5}$$

Ainsi, au sein d'une octave nous avons trois tiers d'octave.

Prenons deux fréquences f_1 et f_2 situé à un tiers d'octave d'intervalle :

La bande de fréquence est :

$$\Delta_f = f_2 - f_1 = 2^{\frac{1}{3}}f_1 - f_1 = \left(2^{\frac{1}{3}} - 1\right)f_1 \quad \text{II-6}$$

La fréquence médiane est définie :

$$f_{med} = \sqrt{f_1 f_2} = \sqrt{2^{\frac{1}{3}}f_1 f_1} = f_1 \sqrt{2^{\frac{1}{3}}} = f_1 2^{\frac{1}{6}} \quad \text{II-7}$$

On a donc :

$$f_{med} \approx 1,12f_1$$

On peut vérifier que l'analyse par tiers d'octave est bien à $\frac{\Delta_f}{f_{med}}$ constant :

$$\frac{\Delta_f}{f_{med}} = \frac{\left(2^{\frac{1}{3}} - 1\right)}{2^{\frac{1}{6}}} \approx 0,23f_{med} \quad \text{II-8}$$

D. Une nouvelle technique de représentation de signaux acoustiques: l'*Hearingogram*

Dans le but de développer une nouvelle méthode de représentation des signaux acoustiques sous-marins, nous proposons un plan temps-fréquence basé sur la physiologie humaine. Cela signifie que le but est d'obtenir une similarité perceptuelle entre ce que nous entendons et ce que nous voyons sur le plan temps-fréquence. En d'autres termes « on voit ce que l'on entend avec la même sensation ». Nous allons donc décrire une approche permettant d'obtenir un plan temps-fréquence représentatif de l'audition humaine, et pour arriver à cela nous utilisons les filtres de Mel [32].

1. Filtres de Mel

Dans le domaine de la physiologie humaine il a été prouvé que l'oreille agit comme un banc de filtres, qui sont concentrés sur seulement certaines composantes fréquentielles [30]. C'est pourquoi les filtres de Mel sont espacés non-uniformément sur l'axe des fréquences [32], où nous avons beaucoup de filtres dans les régions basses fréquences et peu de filtres dans les régions hautes fréquences (Figure 21 et Figure 23). Plus précisément, les filtres de Mel forment un banc de filtres et chacun des filtres a un gabarit triangulaire. Les sommets des différents triangles sont espacés selon l'échelle de Mel, donnée par la formule suivante fonction de la fréquence f en Hz :

$$mel(f) = 2595 \log_{10} \left(1 + \frac{f}{700}\right) \quad \text{II-9}$$

Ainsi, l'échelle des fréquences de Mel est linéaire en dessous de 1000 Hz et devient logarithmique au-dessus de 1000 Hz , ce phénomène peut être observé sur la Figure 21.

Si nous considérons M filtres de Mel $H_{Mel}(f; m)$, chacun d'eux est centré sur une fréquence f_m , $\forall m = 1, 2, \dots, M$ la largeur $B(m)$ est définie comme suit :

$$B(m) = f_{m+1} - f_{m-1} \quad \text{II-10}$$

$$\forall m = 2, 3, \dots, M - 1$$

La fréquence centrale f_m est calculée à partir de sa fréquence centrale sur l'échelle de Mel en utilisant la formule inverse suivante :

$$f_m = 700 \left(10^{\frac{mel(f_m)}{2595}} - 1 \right) \quad \text{II-11}$$

$$\text{avec } mel(f_m) = \frac{m}{M+1} [mel(f_{max}) - mel(f_{min})]$$

où f_{min} et f_{max} correspondent respectivement à la plus haute et à la plus basse fréquence étudiée du signal, généralement $f_{min}=0$ et $f_{max} = \frac{F_e}{2}$ où F_e est la fréquence d'échantillonnage.

Dans ces conditions la réponse impulsionnelle $h_{Mel}(t; m)$, associée au filtre de Mel $H_{Mel}(f; m)$ est donnée par :

$$h_{Mel}(t; m) = h_{Mel}^+(t; m) + h_{Mel}^-(t; m) \quad \text{II-12}$$

avec $h_{Mel}^-(t; m)$ définie comme :

$$\begin{aligned} h_{Mel}^-(t; m) = & \frac{2}{f_m - f_{m-1}} \left(f_m^2 \text{sinc}(2f_m t) - f_{m-1}^2 \text{sinc}(2f_{m-1} t) \right) \\ & + \frac{2}{f_{m-1} - f_m} \left(f_m \text{sinc}(2f_m t) - f_{m-1} \text{sinc}(2f_{m-1} t) \right) \\ & - \frac{f_m^2 - f_{m-1}^2}{f_m - f_{m-1}} \left(f_m^2 \text{sinc}(2f_m t) - f_{m-1}^2 \text{sinc}(2f_{m-1} t) \right) \end{aligned} \quad \text{II-13}$$

avec $h_{Mel}^+(t; m)$ définie comme :

$$\begin{aligned} h_{Mel}^+(t; m) = & \frac{2}{f_m - f_{m-1}} \left(f_{m+1}^2 \text{sinc}(2f_{m+1} t) - f_m^2 \text{sinc}(2f_m t) \right) \\ & + \frac{2f_{m+1}}{f_{m+1} - f_m} \left(f_{m+1} \text{sinc}(2f_{m+1} t) - f_m \text{sinc}(2f_m t) \right) \\ & - \frac{f_{m+1}^2 - f_m^2}{f_m - f_{m+1}} \left(\text{sinc}((f_m + f_{m+1})t) \text{sinc}((f_{m+1} - f_m)t) \right) \end{aligned} \quad \text{II-14}$$

$$\text{où } \text{sinc}(t) = \frac{\sin \pi t}{\pi t}.$$

La Figure 21 présente un banc de filtre de Mel,

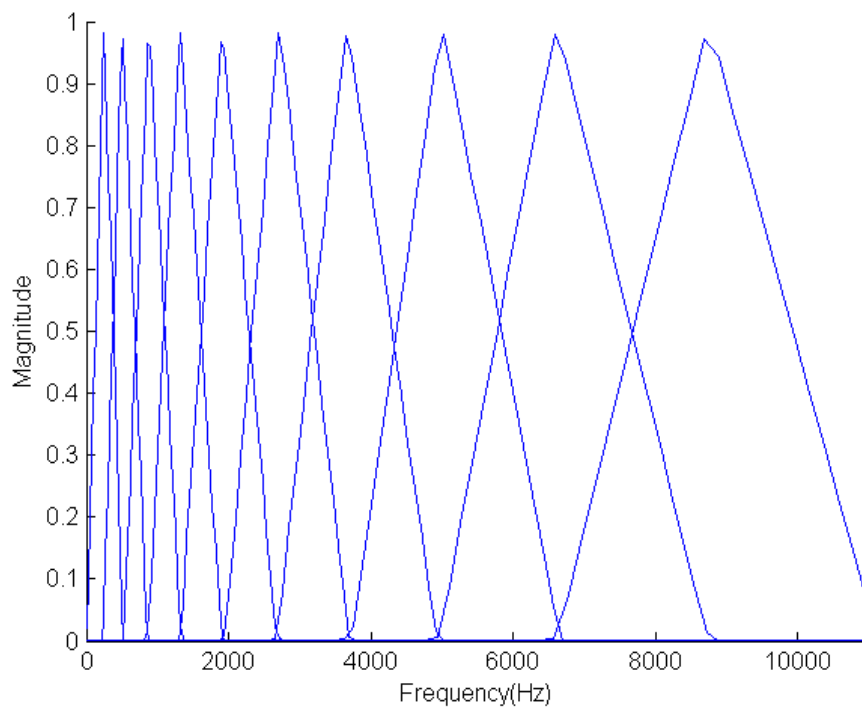


Figure 21: Banc de filtres de Mel, avec $M=10$ pour $f_{min} = 0$ Hz et $f_{max} = 11025$ Hz ;

Sur la Figure 22, nous présentons la réponse impulsionnelle $h_5(t)$.

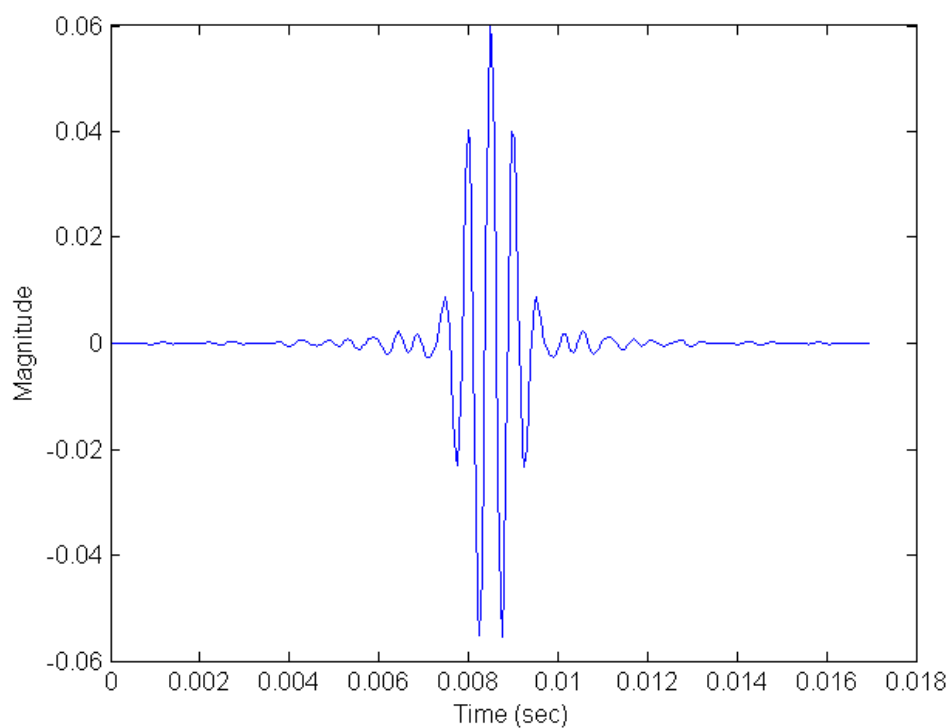


Figure 22: Réponse impulsionnelle associée à $h_5(t)$

Si nous analysons la réponse impulsionnelle d'un de ces filtres, nous remarquons qu'elle a les caractéristiques d'une ondelette mère, c'est à dire localisée en temps et avec une nature oscillante, et semble être similaire à une ondelette de Morlet, mais nous verrons plus tard que notre approche diffère de l'approche des ondelettes.

Il existe dans la littérature, beaucoup de manières d'implémenter un banc de filtres de Mel. Davis [33] nous propose un exemple d'implémentation avec un espacement des fréquences linéaires jusqu'à 1KHz et après un espacement logarithmique, où l'amplitude du filtre est constante (Figure 21). Les filtres sont donc à énergie croissante. Pour notre approche, nous allons utiliser deux implémentations différentes, la première est celle qui vient d'être décrite et dans la seconde (Figure 23) chaque filtre a une énergie unitaire (en d'autres termes l'aire de chaque triangle est unitaire). Cette implémentation est décrite dans [14], mais nous considérons toutes les fréquences contrairement à ce qui est fait dans la référence citée précédemment, où les fréquences en dessous de 133 Hz sont supprimées.

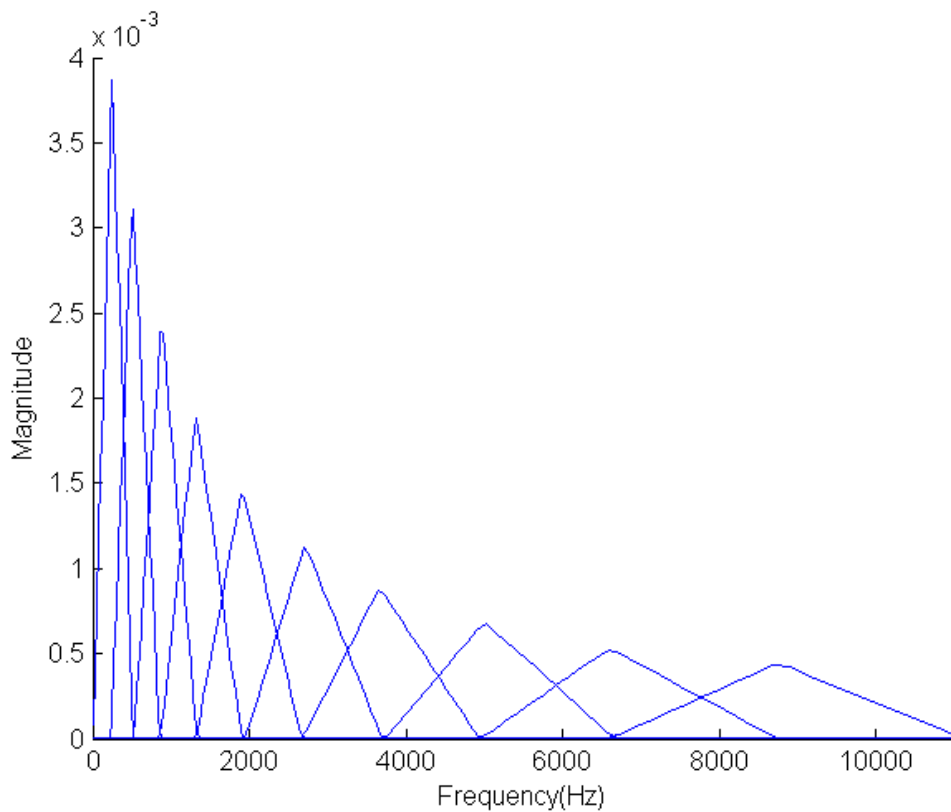


Figure 23: Banc de filtres de Mel, avec énergie unitaire et M=10

2. Principe

Le produit de convolution entre un signal $s(t)$ et la réponse impulsionnelle $h_{Mel}(t; m)$ correspond à un filtrage dans la bande de fréquences $H_{Mel}(f; m)$, ainsi grâce à ce filtrage nous obtenons des informations à propos des composantes spectrales autour de la fréquence f_m . Quand cette étape de filtrage est réalisée $\forall m = 1 \dots M$, il est possible de construire un plan temps-fréquence, où les fréquences qui correspondent aux fréquences centrales f_m de chaque filtre sont espacées selon

l'échelle de Mel. En élevant au carré chaque valeur du plan temps-fréquence, nous obtenons la distribution de la puissance instantanée du signal dans le plan de Mel :

$$\theta_S^{inst}(t, f_m) = |s * h_{Mel}(t; m)|^2 \quad \text{II-15}$$

Le plan ainsi obtenu est appelé « *Instantaneous Hearingogram* ».

Nous remarquons que si le nombre de filtres de Mel tends vers l'infini, la largeur de bande $B(m)$ de chaque filtre de Mel, décrite à la relation II-10, tendra vers 0, tel que chaque filtre de Mel $H_{Mel}(f; m)$ sera comparable à un Dirac centré sur la fréquence f_m . Dans ces conditions, la réponse impulsionnelle $h_{Mel}(t; m)$ tendra vers un signal monochromatique et ainsi l'*Instantaneous Hearingogram* tendra vers le spectrogramme, obtenue par la transformée de Fourier à court terme, avec une fenêtre glissante rectangulaire, nous parlons aussi de sonographe [9].

Dans le but de réduire le niveau de bruit, nous construisons maintenant l'*Hearingogram*, qui représente la distribution énergétique obtenue en intégrant l'*Instantaneous Hearingogram* le long de l'axe temporel :

$$\theta_S(\tau, f_m) = \frac{1}{T} \int_{\tau}^{\tau+T} \theta_S^{inst}(t, f_m) dt \quad \text{II-16}$$

Cette étape est donc une opération de lissage. Plus le temps d'intégration T est important plus le niveau de bruit sera atténué, cependant le signal utile sera affecté par cette opération. Nous avons donc choisi pour l'*Hearingogram*, la durée T correspondant à la plus petite durée des réponses impulsionnelles des différents filtres constituant le banc, c'est-à-dire le triangle avec la plus grande largeur de bande utile, donc le dernier filtre du banc. Un autre moyen pour atténuer le bruit au sein de l'*Hearingogram*, sera vu dans le chapitre concernant la réduction du bruit.

3. Formulation discrète de l'*Hearingogram*

Nous avons vu dans la partie précédente le principe de l'*Hearingogram*. Nous allons maintenant en déduire son écriture discrète, afin d'envisager son implantation de la façon la plus simple possible.

Considérons une observation Z qui contient M_{obs} échantillons, qui est un vecteur de données obtenues par l'enregistrement d'un signal continu échantillonné à la fréquence F_e . Pour construire l'*Hearingogram* nous réalisons donc le produit de convolution entre l'observation Z et la réponse impulsionnelle de chaque filtre de Mel h_m , $\forall m = 1 \dots M$.

Une remarque doit être faite sur le nombre d'échantillons à utiliser pour décrire la réponse impulsionnelle de chacun des filtres. La taille de cette réponse impulsionnelle sera dépendante de la fréquence f_m du filtre, plus particulièrement de la bande $B(m)$. En effet, comme nous l'avons déjà vu, la largeur de bande est plus petite lorsque le filtre est en basse fréquence et inversement. En conséquence, le support temporel de la réponse impulsionnelle est plus petit pour les hautes fréquences que pour les basses fréquences. Ainsi si nous appelons M_{h_m} le nombre d'échantillons utiles pour décrire la réponse impulsionnelle du $m^{ième}$ filtre, nous avons alors la relation suivante :

$$M_{h_p} < M_{h_q}, \quad \forall p > q \quad \text{II-17}$$

Nous réalisons donc l'étape de convolution entre le signal et chaque filtre de Mel du banc de filtres. Après mise en quadrature, on obtient la distribution de la puissance instantanée du signal dans le plan de Mel :

$$\theta_Z^{inst}[k, m] = \left(\sum_{n=1}^{M_{h_m}} Z[k - n] h_m[n] \right)^2 \quad \text{II-18}$$

Après discrétisation, l'étape d'intégration nous permettant de réduire le bruit (tout en dégradant le signal utile comme nous l'avons vu plus haut) est décrite par la relation suivante :

$$\theta_Z[k_0, m] = \frac{1}{K} \sum_{k=k_0}^{k_0+K-1} \theta_Z^{inst}[k, m] \quad \text{II-19}$$

Où K correspond aux nombres d'échantillons utilisés pour l'étape de lissage. Plus ce nombre est grand et moins il y aura de bruit, cependant le signal utile sera dégradé et il y a donc un compromis à trouver. Le nombre d'échantillons utiles pour le lissage sera égal au nombre d'échantillons permettant de décrire la plus petite réponse impulsionnelle du banc de filtre, on a donc :

$$K = M_{h_M} \quad \text{II-20}$$

En faisant cela nous réalisons donc une étape de lissage qui a pour conséquence la diminution du niveau de bruit mais aussi la dégradation de la résolution temporelle. C'est pour cela que dans la partie concernant la réduction du bruit une technique basée sur les ondelettes sera exposée.

L'algorithme peut être résumé grâce au schéma fonctionnel présenté en Figure 25

4. Comparaison entre Hearingogram et ondelettes

Nous avons vu précédemment en observant une réponse impulsionnelle d'un des filtres de Mel que cette dernière semble avoir les caractéristiques d'une ondelette. Il est alors légitime de se demander si l'*Hearingogram* n'est rien d'autre qu'une transformée en ondelettes continue, mais avec une ondelette mère particulière. Dans le but de réfuter une telle idée, nous allons comparer les deux techniques appliquées sur les mêmes signaux.

La résolution temporelle dépend de l'ondelette utilisée pour calculer le scalogramme, mais la résolution fréquentielle est directement liée au principe de la transformée en ondelettes continue, ce qui est vérifié par les tracés de courbe présentés, en Figure 24.

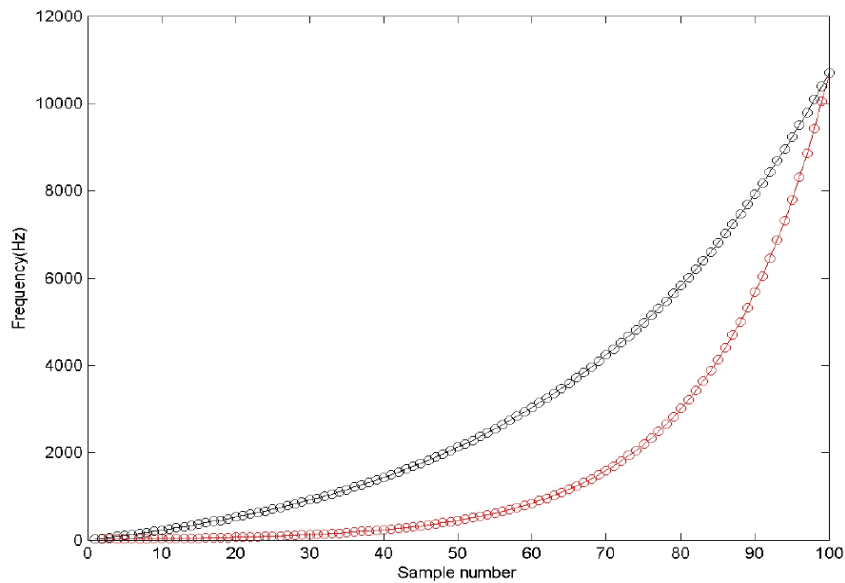


Figure 24: Valeurs des fréquences étudiées pour chaque ligne du scalogramme (ligne rouge) et de l'Hearingogram (ligne noire)

Figure 25: Schéma fonctionnel de l'Hearingogram

Sur la Figure 26 une comparaison est effectuée entre l'Hearingogram obtenu avec 100 filtres et le scalogramme obtenu avec 100 échelles, où la plus grande échelle est calculée à partir de la plus petite fréquence centrale du banc de filtres.

La comparaison entre ces deux plans laisse apparaître clairement que les résolutions temporelle et fréquentielle sont différentes.

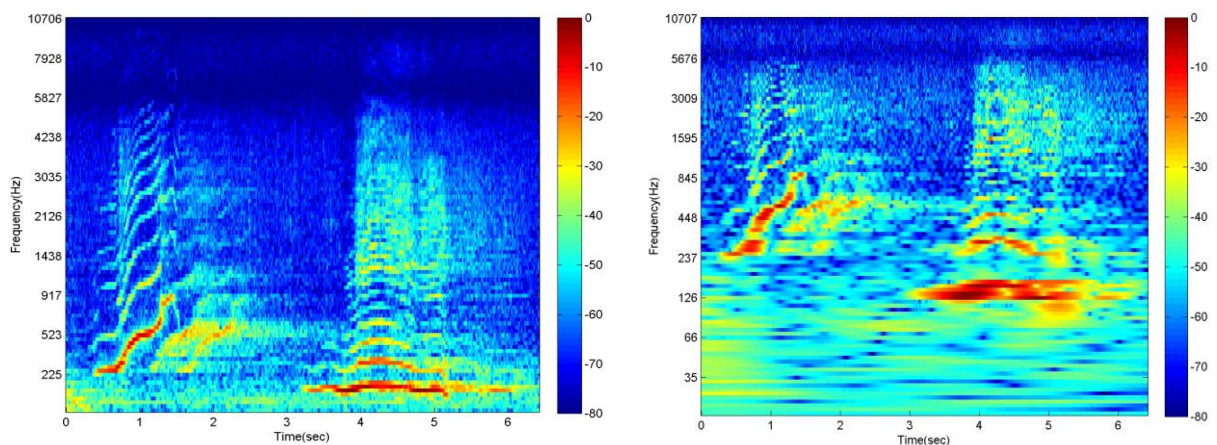


Figure 26: Comparaison entre scaogramme (à gauche) et Hearingogram (à droite)

Ainsi nous pouvons dire que l'Hearingogram est une nouvelle méthode de représentation des signaux, à mi-chemin entre le spectrogramme et le scalogramme. De plus, la transformée en ondelettes s'appuie sur une ondelette mère en la dilatant ou compressant, conservant ainsi le

nombre d'oscillations présentes, alors que l'*Hearingogram* utilise des réponses impulsionnelles où les oscillations dépendent de la bande de fréquence du filtre.

5. Résultats

Dans cette section nous allons présenter trois résultats de l'*Hearingogram* sur des signaux réels de mammifères marins.

Tout d'abord, nous avons testé l'influence du nombre de filtres sur la qualité de la représentation. Sur la Figure 27, nous utilisons respectivement 10, 20, 50 et 100 filtres de Mel, sur un signal représentant une vocalise de baleine.

Ces résultats montrent que même avec un petit nombre de filtres de Mel, une information significative peut être extraite des données, ainsi il est possible d'extraire rapidement quelques caractéristiques du signal, ce qui est très intéressant pour l'identification de signatures présentes dans les données. Pour un plus grand nombre de filtres de Mel, comme la résolution fréquentielle est plus fine plus de détails sont visibles, en particulier les harmoniques présentes dans le signal, qui sont à présent bien visibles.

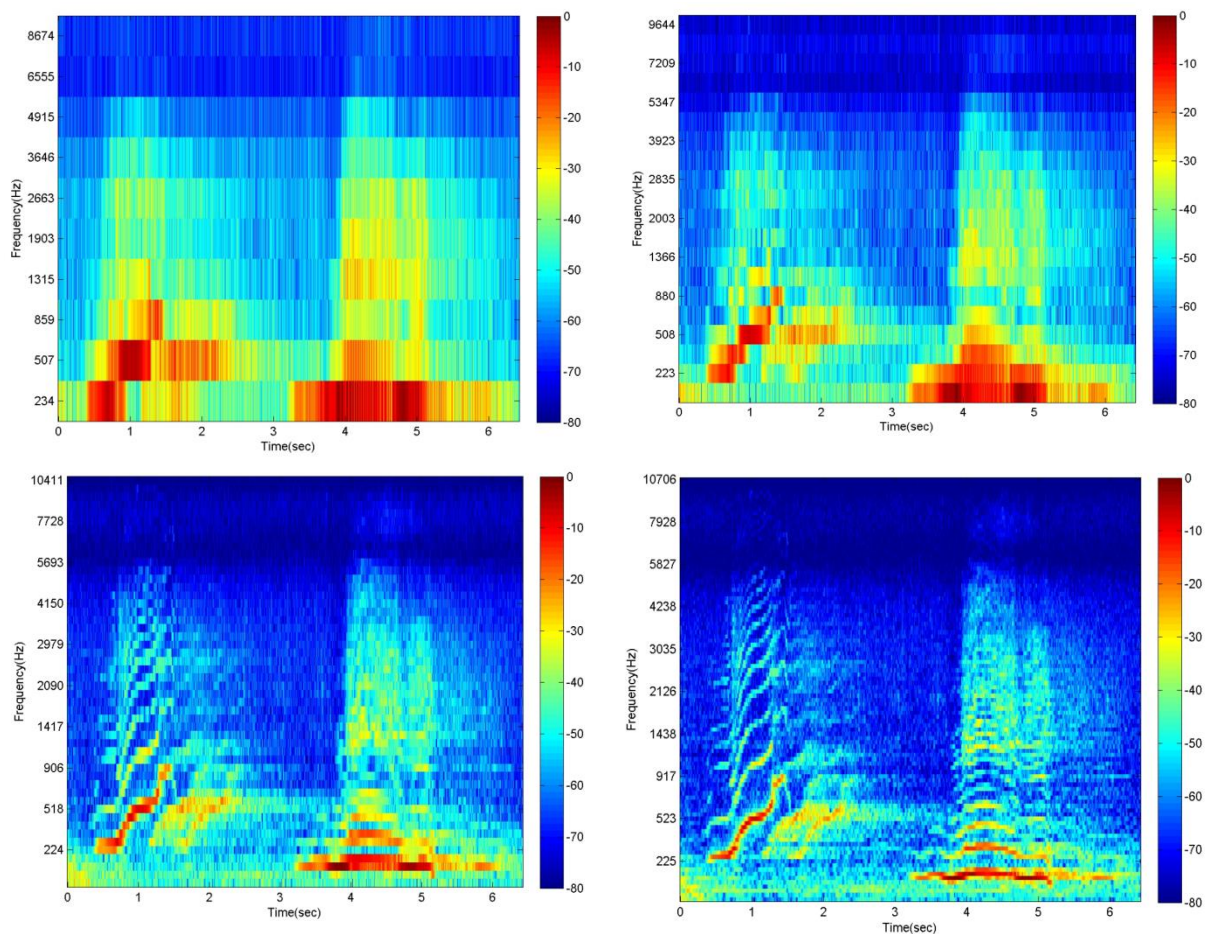


Figure 27: Effet du nombre de filtres sur l'*Hearingogram*

Nous réalisons maintenant l'expérience avec un nombre de filtres plus élevé, c'est-à-dire avec 300 et 800 filtres. L'expérience est réalisée sur le même signal que précédemment. Le résultat est présenté sur la Figure 28.

Nous remarquons, que sur l'*Hearingogram* effectué avec 800 filtres, la résolution temporelle est nettement moins bonne que dans les autres. Ceci s'explique par le fait que plus le nombre de filtres sera élevé plus la base de chaque filtre de Mel sera fine, ce qui signifie que la réponse impulsionnelle aura besoin d'être décrite avec beaucoup plus d'échantillons, entraînant de fait une perte de qualité sur la résolution temporelle. Le compromis entre résolution temporelle et fréquentielle est ainsi présent, le choix de 200 filtres nous paraît approprié compte tenu de notre connaissance des signaux réels.

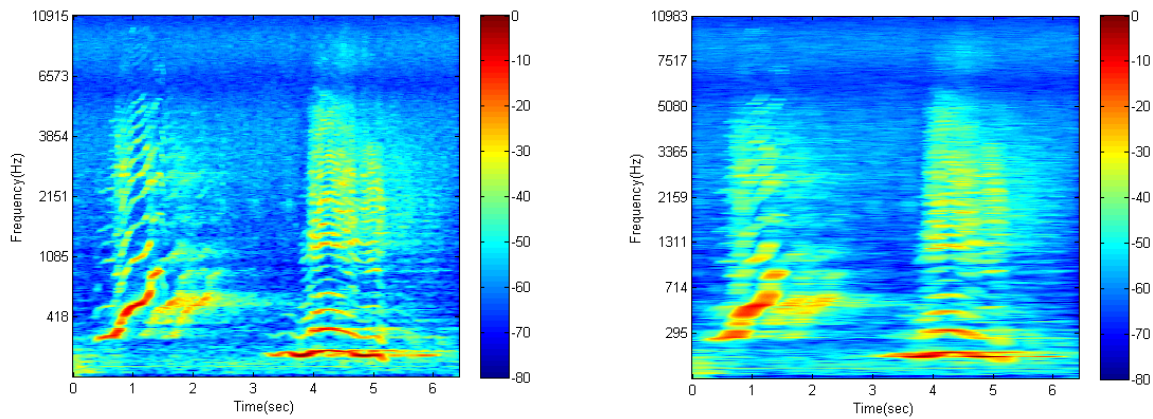


Figure 28 : Effet d'un grand nombre de filtres sur l'*Hearingogram* avec 300 filtres (à gauche) et 800 filtres (à droite)

Nous présentons, à présent, quelques résultats appliqués sur des signatures de mammifères marins. Dans cette expérience une comparaison entre le spectrogramme et l'*Hearingogram* est réalisée. Tous les plans ont été calculés en utilisant 200 filtres de Mel à amplitude constante.

Le premier exemple correspond à un chant de dauphin ($F_e = 22050 \text{ Hz}$). Une analyse des plans temps-fréquence obtenus autour de 2 secondes et de 2 kHz met en exergue la richesse de l'*Hearingogram* comparé au spectrogramme. En effet, dans le spectrogramme, nous remarquons qu'il y a un motif temps-fréquence qui est difficilement interprétable, tandis que nous voyons clairement un enchevêtrement d'harmoniques sur l'*Hearingogram*, ce qui laisse à penser la présence de plusieurs dauphins.

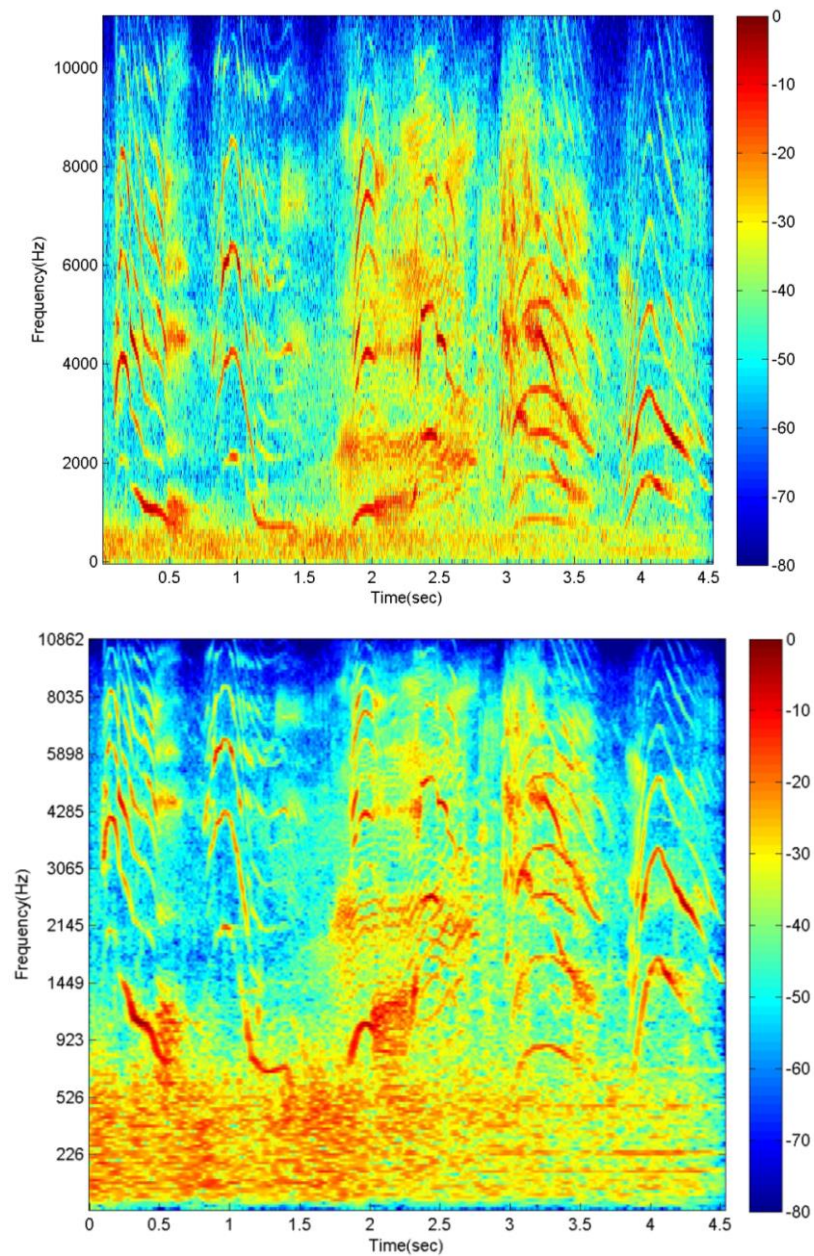


Figure 29: Comparaison entre le spectrogramme (en haut) et l'*Hearingogram* (en bas) d'un son de dauphin

Le signal suivant représente l'écholocation d'un d'orque. Bien que l'écholocation corresponde à plusieurs signaux de durée très courte (signaux impulsifs), chaque clic est clairement visible sur l'*Hearingogram* avec un contraste plus élevée que dans le spectrogramme.

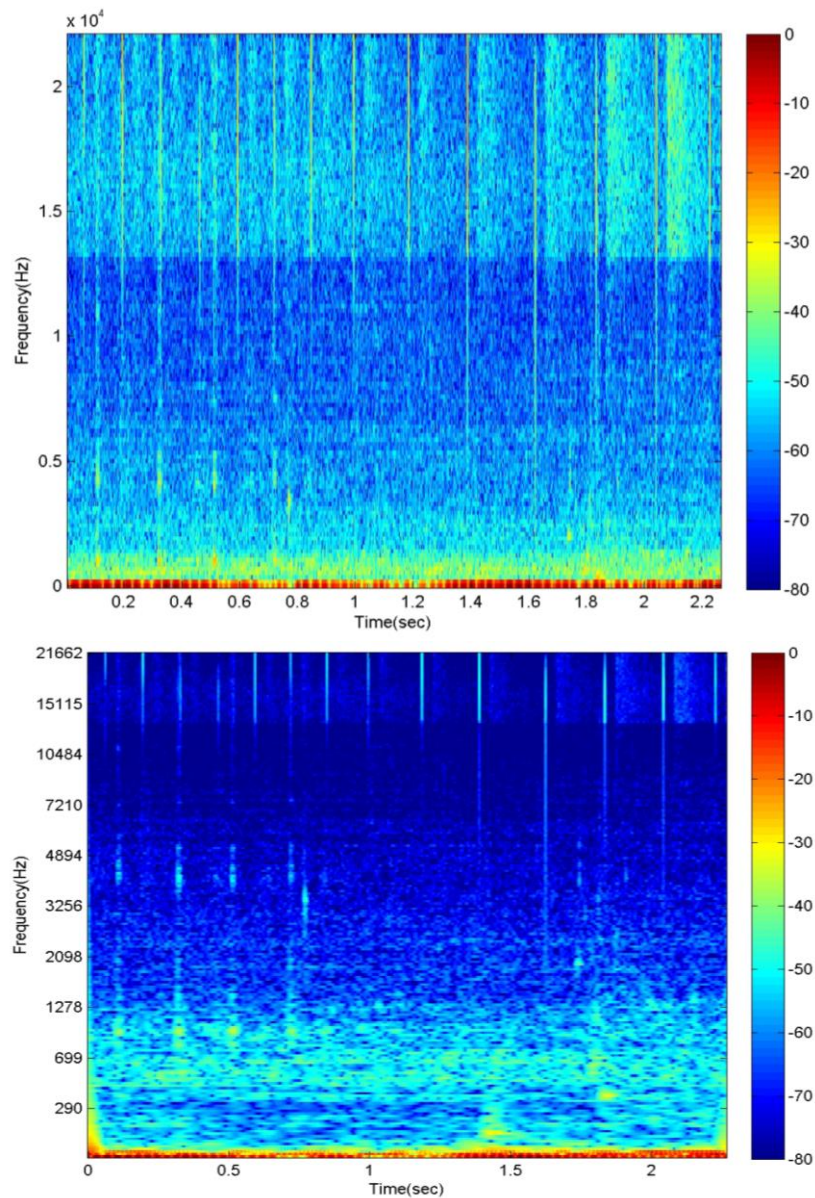


Figure 30: Comparaison entre le spectrogramme (en haut) et l'*Hearingogram* (en bas) d'écholocations d'orques

Le dernier signal étudié contient des vocalises d'orques, le résultat est présenté sur la Figure 31. Encore une fois l'*Hearingogram*, nous permet d'obtenir plus d'informations sur le signal étudié, en particulier autour de t égal à 1 s et t égal à 4 s et entre t égal à 5 et 6 pour les composantes basse-fréquences.

Ces expériences révèlent que l'*Hearingogram* permet d'obtenir une meilleure restitution de plusieurs phénomènes en comparaison avec le spectrogramme. Il est à noter que nous avons réalisé un paramétrage qui permet d'après un critère visuel d'obtenir des performances correctes en termes de résolutions temporelle et fréquentielle. Malgré tout, il est difficile d'obtenir un paramétrage optimal au sens d'un critère bien défini, et qui donne de bons résultats sur tous types de signaux.

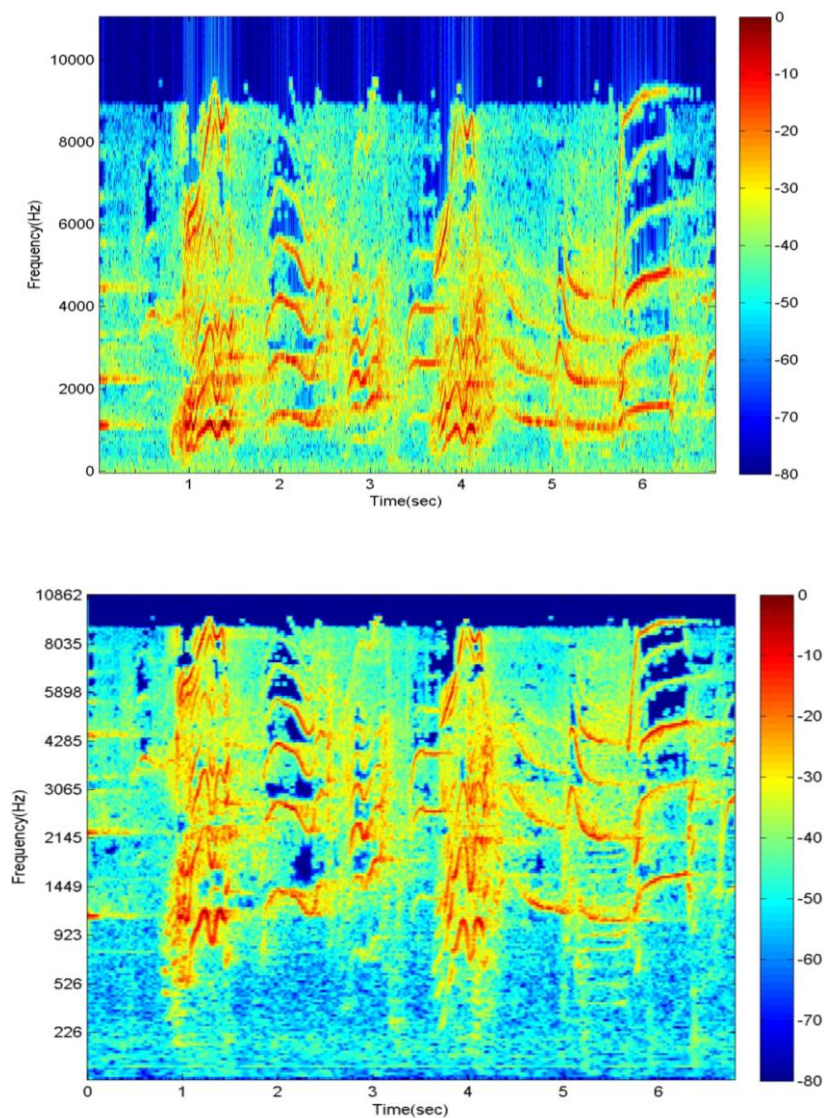


Figure 31: Comparaison entre le spectrogramme (en haut) et l'*Hearingogram* (en bas) de vocalises d'orque

Enfin, chaque spectrogramme a été calculé avec 2048 lignes alors que chaque *Hearingogram* a été calculé avec 200 lignes, ce qui permet d'économiser l'espace. Ainsi l'*Hearingogram* nous permet d'avoir une information plus pertinente que le spectrogramme, tout en occupant moins de place en mémoire et avec un paramétrage beaucoup moins compliqué puisque le seul paramètre de la méthode est le nombre de filtres à choisir. Si le spectrogramme est calculé plus rapidement que l'*Hearingogram* à 200 filtres, l'*Hearingogram* reste parallélisable ce qui laisse un temps de calcul acceptable pour une application temps réel.

Après avoir vu différentes manière de représenter le signal dans le plan temps-fréquence, nous allons nous intéresser à la réduction du bruit qui perturbe le signal utile.

III. Réduction du bruit des signaux non-stationnaires

A. Introduction

Un état de l'art des techniques de réduction de bruit a été réalisé, ce dernier est un préambule à la définition de nouveaux algorithmes plus performants. Des techniques ont été à ce titre développées et comparées. Ceci pour répondre à des problématiques de détection sous-marine rencontrées en reconnaissance acoustique:

- La restitution audio sur veille panoramique : les signaux observés sont des extractions audio de veille panoramique. Il s'agit alors de réaliser des techniques de traitement du signal améliorant la restitution audio, capables de s'affranchir de sources parasites telles que les bruiteurs saturants.
- La visualisation des signaux acoustiques sous-marins : les signaux audio extraits sont analysés à travers leur représentation temps-fréquence afin de mettre en évidence les événements non stationnaires. Le bruit entachant l'information utile nuit à la détection et donc à l'interprétation visuelle de tels événements.
- La reconnaissance des signaux acoustiques sous-marins : le bon fonctionnement de ce système est conditionné par une description exhaustive et robuste des formes observées. Or, en pratique la variabilité du contexte environnemental induit la non-stationnarité du RSB des données observées : ainsi, les pavés temps-fréquence détectés comme étant du signal utile ne se manifestent pas de façon identique dans le plan temps-fréquence si bien que le paramétrage du système peut être désadapté à une forme donnée, ce qui engendre un biais dans la description de cette dernière et donc une possible erreur de classification.

La réponse à ces problèmes spécifiques passe par la formulation du problème de réduction du bruit des signaux acoustiques sous-marins. De façon générale, le modèle d'observation est tel que le signal utile contenu dans les données est perturbé par du bruit de mer. Dans ces conditions, l'opération de réduction du bruit consiste à extraire le signal utile de son observation. Elle doit conjointement éliminer le bruit et préserver les composantes du signal utile.

Les techniques développées doivent permettre de réduire du bruit les signaux issus de systèmes SONAR passifs et par là améliorer la performance des fonctions d'analyse et de reconnaissance qui suivent leur détection. En bout de chaîne, l'analyste ou un organe de décision automatisé, en tant qu'utilisateur final, a besoin d'outils de représentation des signaux lui apportant le confort audio et visuel nécessaire aux opérations de reconnaissance acoustique aboutissant à une identification. En ce sens, augmenter la performance de la chaîne de traitements de visualisation et de restitution audio par réduction de bruit des signaux engendre un gain d'intégration visuelle et audio améliorant la restitution de l'information et donc l'aide à la décision apportée à l'opérateur. Ce gain est à rattacher au gain de traitement (GT) rencontré dans l'équation du SONAR [34].

B. Etat de l'art

Nous allons dresser un état de l'art des techniques de réduction de bruit. Nous allons nous intéresser uniquement aux méthodes qui utilisent la représentation temps-fréquence pour réduire le bruit car elles sont plus performantes que les méthodes qui travaillent uniquement sur le temporel. En effet, les techniques de réduction de bruit dans le domaine temporel s'appliquent généralement à l'aide de traitements par fenêtre glissante dont la durée est fixée compte tenu de la longueur de cohérence du signal. Sur la durée de la fenêtre, le signal utile est ainsi considéré comme stationnaire. Cette hypothèse appliquée pour toute la durée de la fenêtre ne peut être vérifiée, entraînant de fait un fort lissage du signal d'intérêt. On préfère donc s'appuyer sur les plans temps-fréquence afin de caractériser le caractère non-stationnaire du signal.

Toutes les méthodes présentées ont été construites à la base pour réduire le bruit au sein des signaux de la parole. Voici la philosophie générale de chacun de ces algorithmes:

- Calcul de la transformée temps-fréquence.
- Estimation des paramètres statistiques du bruit.
- Estimation de :

$$RSB_{prio}[k, l] = \frac{E\{|S[k, l]|^2\}}{E\{|Z[k, l]|^2\}} \quad \text{III-1}$$

- Estimation du signal débruité grâce à une fonction de gain dépendant des quantités précédentes:

$$\hat{S} = G(R\hat{S}B_{prio}, R\hat{S}B_{post})Z \quad \text{III-2}$$

- Reconstruction du signal débruité

1. Méthodes d'estimation de RSB *a priori*

a. Estimateur Decision directed (DD)

1. Description

Ephraim et Mallah sont les premiers à introduire le concept de RSB *a priori* dont ils développent un estimateur dans [35]. Ils en construisent l'estimateur du maximum de vraisemblance $R\hat{S}B_{prio}[k, l]$ de $RSB_{prio}[k, l]$ qui maximise la quantité suivante :

$$P(Z[k, l] | E\{|B[k]|^2\}, E\{|S[k]|^2\}) \quad \text{III-3}$$

définie par :

$$\hat{S}[k, l] = \max \left(\frac{1}{N_{lissage}} \sum_{n=0}^{N_{lissage}-1} |Z[k, l-n]|^2 - E\{|B[k]|^2\}, 0 \right) \quad \text{III-4}$$

où $N_{lissage}$ correspond au nombre de spectres sur lequel on effectue l'opération de moyennage.

Ils en déduisent un estimateur du RSB *a priori* :

$$R\hat{S}B_{prio}[k, l] = \max \left(\frac{1}{N_{lissage}} \sum_{n=0}^{N_{lissage}-1} RSB_{post}[k, l-n] - 1, 0 \right) \quad \text{III-5}$$

Cet estimateur prend donc la forme d'une moyenne glissante qu'ils choisissent de remplacer par une moyenne réursive car plus rapide à calculer.

L'estimateur devient :

$$R\hat{S}B_{prio}[k, l] = \alpha R\hat{S}B_{prio}[k, l-1] + (1-\alpha) \max(RSB_{post}[k, l] - 1, 0) \quad \text{III-6}$$

avec $0 \leq \alpha < 1$.

En remarquant que $\hat{S}[k, l]$ peut s'écrire sous la forme :

$$\hat{S}[k, l] = G(R\hat{S}B_{prio}[k, l-1], RSB_{post}[k, l-1])Z[k, l] \quad \text{III-7}$$

où G est une fonction de gain (spectral) à valeur dans $[0,1]$ qui supprime la part de bruit des coefficients $Z[k, l]$ en utilisant l'information des deux RSB. L'estimateur du *RSB a priori* devient :

$$\begin{aligned} R\hat{S}B_{prio}^{DD}[k, l] &= \frac{|\hat{S}[k, l]|^2}{E\{|B[k]|^2\}} \\ &= \alpha G^2(R\hat{S}B_{prio}[k, l-1])RSB_{post}[k, l] \\ &\quad + (1-\alpha) \max(RSB_{post}[k, l] - 1, 0) \end{aligned} \quad \text{III-8}$$

Ce qui est l'estimateur Decision-directed développé par Ephraim et Mallah.

2. Description technique

Algorithme :

- Calcul de la TFCT du signal bruité avec une fenêtre de taille N_h :

$$Z[k, l] = S[k, l] + B[k, l]$$
- Estimation de σ_B^2
- On effectue l'opération suivante :

$$\sigma_B^2[k] = \sigma_B^2[k] \times N_h$$
- Pour chaque coefficient issu de la RTF calculer :

- $RSB_{post}[k, l] = \frac{|Z[k, l]|^2}{\sigma_B^2[k]}$
- $R\hat{S}B_{prio}[k, l] = \alpha G^2(R\hat{S}B_{prio}[k, l-1])RSB_{post}[k, l] + (1 - \alpha) \max(RSB_{post}[k, l] - 1, 0)$
- $\hat{S}[k, l] = G(R\hat{S}B_{prio}[k, l], RSB_{post}[k, l])Z[k, l]$
- *Reconstruction du signal utile à partir de la RTF débruitée :*

$$\hat{s} = TFCT^{-1}(\hat{S})$$

b. Méthode basée sur l'utilisation de l'estimateur Two Step Noise Reduction (TSNR)

1. Description théorique

Dans [36], les auteurs proposent une modification de l'estimateur Decision Directed (DD). Pour ce faire il propose d'estimer le RSB *a priori* en deux étapes :

La première étape consiste à estimer le RSB *a priori* en utilisant l'estimateur DD sur la trame $l + 1$ pour calculer le gain spectral :

$$G(R\hat{S}B_{prio}[k, l], RSB_{post}[k, l])$$

avec:

$$R\hat{S}B_{prio}[k, l] = R\hat{S}B_{prio}^{DD}[k, l + 1] \quad \text{III-9}$$

La seconde étape consiste à utiliser ce gain pour calculer le RSB *a priori* sur la trame l et enfin en déduire le gain spectral final :

$$R\hat{S}B_{prio}^{TSNR}[k, l] = G^2(R\hat{S}B_{prio}[k, l], RSB_{post}[k, l])RSB_{post}[k, l] \quad \text{III-10}$$

G est une fonction de gain à valeur dans $[0; 1]$ qui supprime la part de bruit des coefficients $Z[k, l]$ en utilisant l'information des deux RSB.

2. Description technique

Algorithme :

- *Calcul de la TFCT du signal bruité avec une fenêtre de taille N_h :*

$$Z[k, l] = S[k, l] + B[k, l]$$
- *Estimation de σ_B^2*
- *On effectue l'opération suivante :*

$$\sigma_B^2[k] = \sigma_B^2[k] \times N_h$$
- *Pour chaque coefficient issu de la RTF faire :*
 - $RSB_{post}[k, l] = \frac{|Z[k, l]|^2}{\sigma_B^2[k]}$

- De nouveau pour chaque coefficient calculer :
 - $R\hat{S}B_{prio}[k, l+1] = \alpha G^2(R\hat{S}B_{prio}[k, l], RSB_{prio}[k, l])RSB_{post}[k, l+1] + (1-\alpha) \max(RSB_{post}[k, l] - 1, 0)$
 - $R\hat{S}B_{prio}^{TSNR}[k, l] = G^2(R\hat{S}B_{prio}[k, l], RSB_{post}[k, l])RSB_{post}[k, l]$
 - $\hat{S}[k, l] = G(R\hat{S}B_{prio}^{TSNR}[k, l], RSB_{post}[k, l])Z[k, l]$
- Reconstruction du signal utile à partir de la RTF débruitée :

$$\hat{s} = TFCT^{-1}(\hat{S})$$

a. Méthode basée sur l'utilisation de l'estimateur IPSE (Improved Priori SNR Estimation)

1. Description

Dans [37] les auteurs proposent d'utiliser la méthode TSNR en modifiant la dernière étape de l'algorithme. Pour cela ils font l'hypothèse selon laquelle S et B suivent des lois gaussiennes centrées (partie réelle et imaginaire) et en déduisent une estimation de $E\{|S(k, l)|^2|Z(k, l)\}$ telle que :

$$\hat{S}[k, l] = \frac{E\{|B[k]|^2\}R\hat{S}B_{prio}^{DD}[k, l]}{R\hat{S}B_{prio}^{DD}[k, l] + 1} + \frac{R\hat{S}B_{prio}^{DD}[k, l]^2}{(R\hat{S}B_{prio}^{DD}[k, l] + 1)^2} |Z[k, l]|^2 \quad \text{III-11}$$

L'algorithme devient :

- Calcul de:

$$R\hat{S}B_{prio}[k, l] = R\hat{S}B_{prio}^{DD}[k, l]$$
- Calcul de l'estimation du RSB a priori :

$$R\hat{S}B_{prio}^{IPSE}[k, l] = \frac{R\hat{S}B_{prio}^{DD}[k, l]}{R\hat{S}B_{prio}^{DD}[k, l] + 1} + \frac{R\hat{S}B_{prio}^{DD}[k, l]^2}{(R\hat{S}B_{prio}^{DD}[k, l] + 1)^2} RSB_{post}[k, l]$$

2. Description technique

L'algorithme mis en place pour filtrer un signal avec la méthode exposée ci-dessus est le suivant :

- Calcul de la TFCT du signal bruité avec une fenêtre de taille N_h :

$$Z[k, l] = S[k, l] + B[k, l]$$

- On effectue le calcul suivant :

$$\sigma_B^2[k] = \sigma_B^2[k] \times N_h$$

- Pour chaque coefficient issu de la RTF faire :

- $RSB_{post}[k, l] = \frac{|Z[k, l]|^2}{\sigma_B^2[k]}$
- $R\hat{S}B_{prio}^{DD}[k, l] = \alpha G^2(R\hat{S}B_{prio}^{IPSE}[k, l], RSB_{prio}^{IPSE}[k, l])RSB_{post}[k, l] + (1 - \alpha) \max(RSB_{post}[k, l] - 1, 0)$
- $R\hat{S}B_{prio}^{IPSE}[k, l] = \frac{R\hat{S}B_{prio}^{DD}[k, l]}{R\hat{S}B_{prio}^{DD}[k, l] + 1} + \frac{R\hat{S}B_{prio}^{DD}[k, l]^2}{(R\hat{S}B_{prio}^{DD}[k, l] + 1)^2} RSB_{post}[k, l]$
- $\hat{S}[k, l] = G(R\hat{S}B_{prio}^{IPSE}[k, l], RSB_{post}[k, l])Z[k, l]$

- Reconstruction du signal utile à partir de la RTF débruitée :

$$\hat{s} = TFCT^{-1}(\hat{S})$$

d. Méthode basée sur l'utilisation de l'estimateur non causal (NC)

1. Description théorique

En 2004 Israël Cohen décide d'améliorer l'estimation du RSB *a priori* en s'affranchissant de l'hypothèse de causalité, en effet il utilise des trames futures et trames passées pour réduire le niveau de bruit de la trame présente. Pour cela il développe une première méthode d'estimation dans [38]. Pour construire son algorithme, Cohen s'inspire du principe du filtre de Kalman [39] et utilise, dans un premier temps, une méthode similaire à IPSE. Il construit tout d'abord un estimateur causal du RSB *a priori*, composé:

- d'une étape de propagation ;
- et d'une étape de mise à jour.

Son étape de propagation est inspirée de la méthode decision-directed proposée par Ephraim et Mallah et son étape de mise à jour suit la même idée que IPSE (mais il modifie l'équation de mise à jour).

« Propagation » step :

$$R\hat{S}B_{prio}[k, l] = R\hat{S}B_{prio}^{DD}[k, l] \quad \text{III-12}$$

« Update » step :

$$R\hat{S}B_{prio}[k, l] = \frac{R\hat{S}B_{prio}^{DD}[k, l]}{R\hat{S}B_{prio}^{DD}[k, l] + 1} \left(1 + \frac{R\hat{S}B_{prio}^{DD}[k, l]}{(R\hat{S}B_{prio}^{DD}[k, l] + 1)} RSB_{post}[k, l] \right) \quad \text{III-13}$$

Un estimateur causal n'étant pas toujours nécessaire, et même sous optimal si cette contrainte n'est pas imposée, Cohen décide d'étendre son estimateur au cas non causal. Pour cela il prend en compte L trames futures qui selon la fréquence d'échantillonnage implique un délai plus ou moins long.

Pour estimer le RSB sur les trames futures Cohen utilise une simple moyenne telle que :

III-14

$$R\hat{S}B_{futur}[k, l] = \max \left(\frac{1}{N_{lissage}} \sum_{n=1}^{N_{lissage}} (RSB_{post}[k, l + n] - 1), 0 \right)$$

Il choisit ensuite d'insérer cette information future dans l'estimation du RSB *a priori* de la même manière qu'est prise en compte l'information passée (avec α et $(1 - \alpha)$).

On obtient alors :

- « Backward-forward » propagation

III-15

$$R\hat{S}B_{prio}^{BF}[k, l] = \max \left(\alpha G^2 \left(R\hat{S}B_{prio}[k, l - 1], RSB_{post}[k, l - 1] RSB_{post}[k, l] \right. \right. \\ \left. \left. + (1 - \alpha)(\beta R\hat{S}B_{prio}[k, l - 1] + (1 - \beta)R\hat{S}B_{futur}[k, l]) \right), 0 \right)$$

- « Update » step

III-16

$$R\hat{S}B_{prio}[k, l] = \frac{R\hat{S}B_{prio}^{BF}[k, l]}{R\hat{S}B_{prio}^{BF}[k, l] + 1} \left(1 + \frac{R\hat{S}B_{prio}^{BF}[k, l]}{(R\hat{S}B_{prio}^{BF}[k, l] + 1)} RSB_{post}[k, l] \right)$$

2. Description technique

Algorithme :

- Calcul de la TFCT du signal bruité avec une fenêtre de taille N_h :

$$Z[k, l] = S[k, l] + B[k, l]$$

- Interpolation du vecteur contenant les σ_B^2 si sa taille ne correspond pas à celle de $Z[:, l]$
- On effectue l'opération suivante :

$$\sigma_B^2[k] = \sigma_B^2[k] \times N_h$$

- Pour chaque coefficient issu de la RTF faire :

$$RSB_{post}[k, l] = \frac{|Z[k, l]|^2}{\sigma_B^2[k]}$$

- De nouveau pour chaque coefficient faire :

- $R\hat{S}B_{futur}[k, l] = \max\left(\frac{1}{L} \sum_{p=1}^L (RSB_{post}[k, l + p] - 1), 0\right)$
- $R\hat{S}B_{prio}^{BF}[k, l] = \max\left(\alpha G^2\left(R\hat{S}B_{prio}[k, l - 1], RSB_{post}[k, l - 1]\right) RSB_{post}[k, l] (1 - \alpha) (\beta R\hat{S}B_{prio}[k, l - 1] + (1 - \beta) R\hat{S}B_{futur}[k, l]), 0\right)$
- $R\hat{S}B_{prio}[k, l] = \frac{R\hat{S}B_{prio}^{BF}[k, l]}{R\hat{S}B_{prio}^{BF}[k, l] + 1} \left(1 + \frac{R\hat{S}B_{prio}^{BF}[k, l]}{(R\hat{S}B_{prio}^{BF}[k, l] + 1)} RSB_{post}[k, l]\right)$
- $\hat{S}[k, l] = G(R\hat{S}B_{prio}[k, l], RSB_{post}[k, l]) Z[k, l]$

- Reconstruction du signal utile à partir de la RTF débruitée :

$$\hat{s} = TFCT^{-1}(\hat{S})$$

2. Règle d'atténuation

Les règles d'atténuations constituent l'ensemble des fonctions de gain $G(R\hat{S}B_{prio}, RSB_{prio})$ à valeurs dans $[0; 1]$ tel que :

$$\hat{S}[k, l] = G(R\hat{S}B_{prio}[k, l], RSB_{post}[k, l]) Z[k, l] \quad \text{III-17}$$

Elles sont dérivées du calcul d'estimateurs minimisant sous certaines hypothèses un critère défini *a priori*. Il en existe donc potentiellement une infinité et nous présenterons ici les plus connues et les plus performantes. Ainsi pour chaque règle d'atténuation nous présentons la formule et des courbes qui représentent le gain qui va être appliqué au plan temps-fréquence en fonction du RSB *a priori* et du RSB *a posteriori*.

a. Wiener

$$G_{Wien}(RSB(k, l)) = 1 - \frac{1}{1 + RSB_{priori}[k, l]} \quad \text{III-18}$$

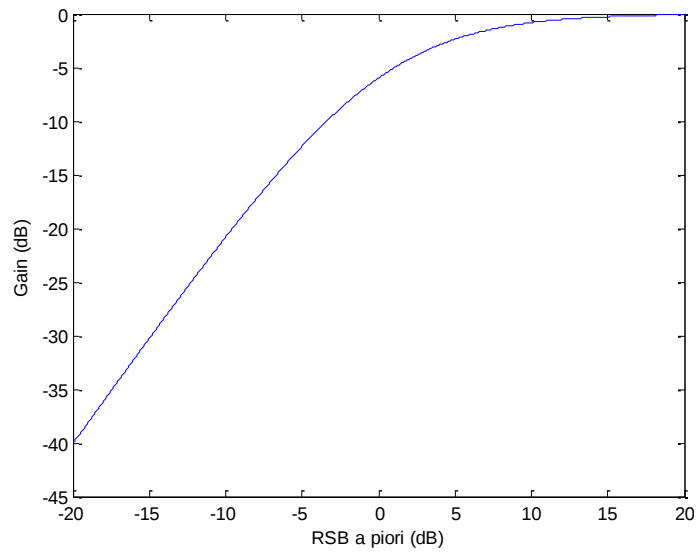


Figure 32: Courbe de gain G_{Wien}

b. L'estimateur Minimum Mean Square Error of Log-Spectral Amplitude (MMSE_LSA [40])

$$G_{LSA}(RSB_{prio}, RSB_{post}) = \frac{RSB_{prio}}{1 + RSB_{prio}} \exp\left(\frac{1}{2} \int_v^\infty \frac{e^{-t}}{t} dt\right)$$

III-19

avec:

$$v = \frac{RSB_{prio}}{1 + RSB_{prio}} RSB_{post}$$

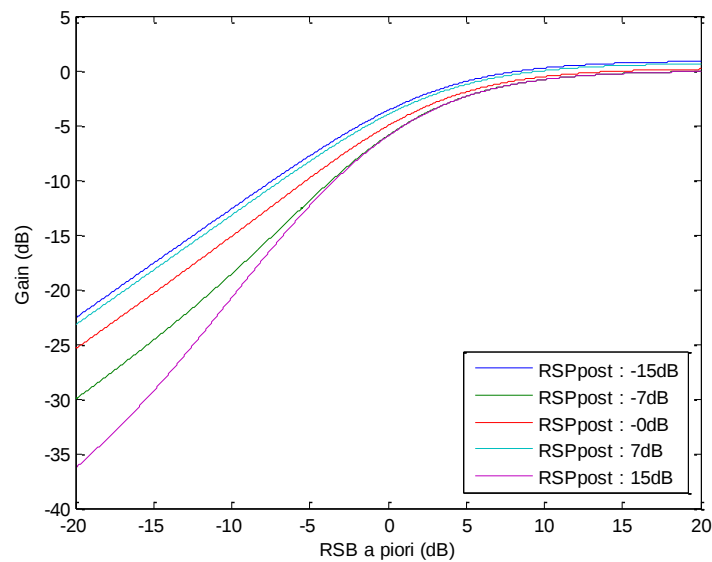


Figure 33: Courbes de gain G_{LSA}

c. Maximum a posteriori (MAP, [41])

$$G_{MAP}(RSB_{prio}, RSB_{post}) = \frac{RSB_{prio} + \sqrt{RSB_{prio}^2 + (RSB_{prio} + 1) \frac{RSB_{prio}}{RSB_{post}}}}{2(RSB_{prio} + 1)}$$

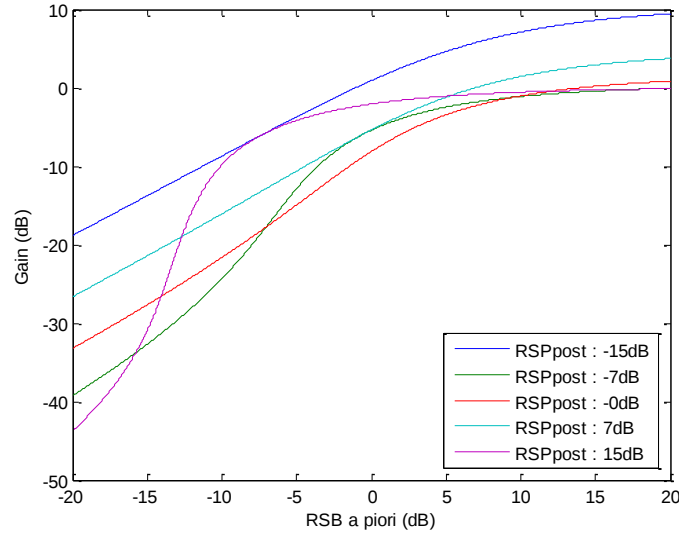


Figure 34: Courbes de gain G_{MAP}

d. JMAP [42]

$$G_{JMAP}(RSB_{prio}, RSB_{post}) = u + \sqrt{u^2 + \frac{v}{2RSB_{post}}}$$

III-20

avec : $u = \frac{1}{2} - \frac{\mu}{4\sqrt{RSB_{post}RSB_{prio}}}$

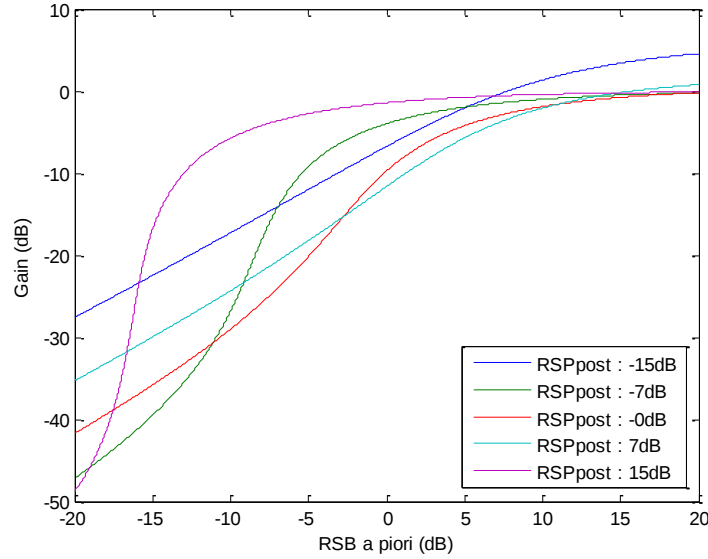


Figure 35: Courbes de gain G_{JMAP}

Ainsi ces différentes techniques de réduction du niveau de bruit sont basées sur le couplage entre une méthode d'estimation du RSB_{priori} et une règle de décision. Ces algorithmes seront comparés à un algorithme développé durant cette thèse basé sur la réduction du bruit au sein de l'*Hearingogram*.

C. Réduction du bruit au sein des signaux acoustiques sous-marins à partir de l'*Hearingogram*: *Denoised Hearingogram*

1. Principe et analyse

Nous avons exposé précédemment le principe de l'*Hearingogram*, plus particulièrement nous avons vu qu'une opération de lissage permettait une réduction du niveau de bruit. Cependant une telle opération dégrade les composantes utiles du signal. C'est pour cela que nous avons pensé à une méthode de réduction du bruit basée sur l'utilisation ondelettes. Nous allons présenter dans ce paragraphe cet algorithme : le *Denoised Hearingogram*.

L'*Hearingogram* instantané est défini comme nous l'avons vu par :

$$\theta_Z^{inst}[k, m] = \left(\sum_{n=1}^{M_{hm}} Z[k-n]h_m[n] \right)^2 \quad \begin{array}{l} \forall k = 1 \dots M_{obs} \\ \forall m = 1 \dots M \end{array} \quad \text{III-21}$$

Cependant comme pour tous les plans temps-fréquence, le bruit présent au sein de l'observation est observable sur l'*Hearingogram*. Le bruit étant supposé additif, en l'absence de signal utile au sein de l'observation, les composantes du bruit présentes sur une ligne de l'*Hearingogram* avant quadrature sont définies par :

$$B_{h_m}[k] = \sum_{n=1}^{M_{h_m}} B[k-n]h_m[n] \quad \text{III-22}$$

où M_{h_m} représente le nombre d'échantillons de la réponse impulsionnelle h_m du $m^{ième}$ filtre du banc de filtre de Mel. Voyons maintenant les différentes étapes permettant d'obtenir le *Denoised Hearingogram*, tout d'abord nous réalisons le produit de convolution discret entre le signal d'entrée Z et chaque filtre h_m constituant le banc de filtre de Mel. La sortie de chaque filtre sera filtrée dans la bande de fréquence couverte par le filtre de Mel correspondant. Ainsi la sortie de chaque filtrage correspondra à une ligne de la matrice *Hear* de dimension $M \times M_{obs}$, on a donc :

$$Hear(m, :) = Z * h_m = Z_{h_m} \quad \forall m = 1 \dots M \quad \text{III-23}$$

Cette technique propose d'opérer la réduction du bruit sur chaque ligne séparément. Sur une ligne de la matrice *Hear* avant quadrature, l'observation s'écrit de la manière suivante :

$$Z_{h_m} = S_{h_m} + B_{h_m} \quad \text{III-24}$$

où S_{h_m} et B_{h_m} représente respectivement l'*Hearingogram* avant l'étape de quadrature S et B .

Nous allons effectuer un traitement permettant de réduire le bruit avant l'opération de quadrature effectuée en fin de construction de l'*Hearingogram*.

Nous allons caractériser de manière statistique le bruit B_{h_m} au sein d'une ligne de la matrice *Hear*. On a :

$$B_{h_m} = B * h_m = \sum_{n=1}^{M_{h_m}} B[k-n]h_m[n] \quad \text{III-25}$$

La valeur moyenne de ce bruit est :

$$E\{B_{h_m}[k, m]\} = E\left\{\left(\sum_{n=1}^{M_{h_m}} B[k-n]h_m[n]\right)\right\} = \sum_{n=1}^{M_{h_m}} E\{B[k-n]\}h_m[n] \quad \text{III-26}$$

En faisant l'hypothèse d'une stationnarité du bruit, il vient :

$$E\{B_{h_m}[k, m]\} = \bar{B} \sum_{n=1}^{M_{h_m}} h_m[n] \quad \text{III-27}$$

Sachant que la réponse impulsionnelle de h_m donnée dans l'équation II-12 est centrée alors la valeur moyenne du bruit B_{h_m} sur une ligne de l'*Hearingogram* sera nulle. De plus, en faisant l'hypothèse que la réponse impulsionnelle h_M du dernier filtre du banc de filtres de Mel contient un assez grand nombre d'échantillons, et comme le produit de convolution appliqué à B est une transformation linéaire d'un vecteur de variables aléatoires, nous pouvons invoquer le théorème de la limite centrale pour dire que la séquence d'échantillons constituant B_{h_m} suit une loi normale, ainsi :

$$B_{h_m} \hookrightarrow N(0, \sigma_{B_m}^2) \quad \forall m = 1 \dots M \quad \text{III-28}$$

Il faut maintenant estimer la variance $\sigma_{B_m}^2$. Cette étape peut-être réalisée de deux façons différentes, selon les informations *a priori* disponibles sur le bruit :

- Soit nous connaissons une partie des données qui est le bruit seul, et on utilise cette connaissance *a priori* pour estimer la puissance du bruit ;
- Soit nous utilisons l'estimateur MAD (median absolute deviation) :

$$\sigma_{B_m} = C \times \text{median}(|Z_{h_m} - \text{median}(Z_{h_m})|) \quad \text{III-29}$$

Où C est une constante de normalisation qui dépend de l'observation de la distribution. Pour des données qui suivent une distribution normale, C est égale à $\frac{1}{0.6745}$ (cette constante est proposé dans [43]) et elle correspond à $\frac{1}{\phi^{-1}(0.75)}$, où $\phi^{-1}(\cdot)$ est l'inverse de la fonction de répartition pour une distribution normale. On choisit la relation III-29 pour estimer l'écart-type du bruit, pour la simple raison que le MAD est plus robuste que l'estimateur classique de l'écart-type du bruit.

Maintenant que nous avons caractérisé le bruit, il faut réaliser l'étape de réduction du niveau de ce bruit. Il existe dans la littérature plusieurs méthodes dédiées à la réduction de bruit. L'une d'entre elles est basée sur l'analyse multi-résolution, très bien adaptée à la réduction de bruit gaussien. Cette méthode développée dans [44] est basée sur 3 étapes essentielles :

- Les données bruitées sont décomposées selon l'analyse multi-résolution choisie, afin d'obtenir un jeu de coefficients d'ondelettes (étape 1) ;
- Application de la règle de seuillage sur les coefficients en ondelettes (étape 2);
- Approximation du signal d'intérêt, en appliquant la transformée en ondelettes discrètes inverse grâce aux coefficients d'ondelettes seuillés (étape 3).

Détaillons ces étapes que nous allons appliquer à Z_{h_m} , $\forall m = 1 \dots M$:

Etape 1

L'algorithme de Mallat est utilisé pour réaliser l'analyse multi-résolution. Deux types de coefficients sont alors obtenus: les approximations et les détails. Les approximations décrivent la forme globale du signal, tandis que les détails décrivent les variations plus fines. Les coefficients de détails de faible intensité véhiculent les termes perturbateurs de l'observation.

Etape 2

Un seuil agissant sur les coefficients d'ondelettes ω_m^p , obtenus pour chaque ligne de la matrice *Hear*, doit être défini. Précédemment nous avons montré que le bruit au sein de l'observation suivait une loi gaussienne de moyenne nulle et d'écart-type σ_{B_m} estimée grâce au MAD. Sous cette hypothèse, nous considérons des échantillons de bruit blanc de puissance unitaire obtenus par tirages aléatoires, nous filtrons cette séquence d'échantillons en la convoluant au banc de filtres de Mel. Après cette opération, M signaux de bruit colorés en fonction de la bande passante de chaque filtre sont obtenus. Une analyse multi-résolution est alors réalisée sur chaque bruit coloré, permettant d'accéder à la répartition, sur les composantes d'approximation et de détails, de la puissance d'un bruit blanc unitaire filtré par le banc de filtres de Mel.

Enfin, on évalue la puissance du bruit filtré dans chaque plan d'ondelettes, notée $\sigma_\eta^m[p]$, pour la suite. Cette quantité dépend du niveau p de l'analyse multi-résolution et du numéro m du filtre de Mel.

Pour seuiller nos coefficients d'ondelettes, nous utilisons le seul universel de Donoho [44], qui est une simple mesure d'entropie dépendant du nombre d'échantillons dans Z_{h_m} , permettant d'obtenir un seuil qui sera appliqué sur chaque plan d'approximation et de détails à la résolution p de l'analyse multi-résolution. Ce seuil est défini comme suit :

$$\lambda_m[p] = \sigma_{B_m} \sigma_\eta^m \sqrt{2 \ln(M_{obs})} \quad \text{III-30}$$

Où M_{obs} représente le nombre d'échantillons de Z_{h_m}

Ainsi, pour le $m^{ième}$ filtre de Mel, nous avons un seuil variable qui dépend de l'échelle p considérée. Il existe deux approches pour la règle de seuillage :

- Le seuillage dur, qui met à zéro tous les coefficients plus petits que le seuil tout en gardant les autres coefficients inchangés, il est défini tel que $\forall k = 1 \dots \frac{M_{obs}}{2^{\max(p)}}$:

$$\omega_m^p = \begin{cases} 0 & \text{si } |\omega_m^p[k]| < \alpha \lambda_m[p] \\ \omega_m^p[k] & \text{ailleurs} \end{cases} \quad \text{III-31}$$

- Le seuillage doux, qui va aussi mettre à 0 les coefficients en dessous du seuil, et va modifier les autres coefficients. En d'autres termes, ceci revient à considérer qu'au-delà du seuil une part de bruit est véhiculée par les coefficients d'ondelettes. Ce seuillage est défini ainsi :

$$\omega_m^p = \begin{cases} 0 & \text{si } |\omega_m^p[k]| < \alpha \lambda_m[p] \\ \omega_m^p[k] - \alpha \lambda_m[p] \text{sign}(\omega_m^p[k]) & \text{ailleurs} \end{cases} \quad \text{III-32}$$

avec k qui prend ses valeurs dans le même intervalle que pour le seuillage dur. Quant à, il s'agit d'une constante utilisée pour ajuster le seuil : plus le seuil est petit, plus la réduction du bruit est minime et inversement plus le seuil est fort, plus le signal d'intérêt est dégradé. Il y a donc un compromis à réaliser. Après cette étape de seuillage les coefficients différents de zéros sont supposés contenir le signal d'intérêt tandis que ceux mis à 0 ne véhiculaient *a priori* que du bruit.

Etape 3

L'analyse multi-résolution inverse est appliquée sur les coefficients restant après seuillage pour obtenir une approximation du signal d'intérêt $\tilde{S}_{h_m}$. Puis chaque ligne de la matrice temps-fréquence est élevée au carré permettant l'obtention du Denoised Hearingogram.

Les différentes étapes de la construction du Denoised Dearingogram se résument ainsi :

- Filtrage de l'observation par chaque filtre constituant le banc de filtres de Mel en réalisant cette opération nous construisons $Z_{h_m} \forall m = 1, \dots, M$;
- Evaluation d'un seuil qui sera différent pour chaque ligne;
- Réalisation d'une analyse multi-résolution de $Z_{h_m} \forall m = 1 \dots M$;

- Pour chaque ligne associée aux fréquences f_m , seuillage des coefficients d'ondelettes;
- Reconstruction du signal à partir des coefficients seuillés, pour obtenir une approximation du signal d'intérêt $\tilde{S}_{h_m} \forall m = 1 \dots M$ dans la bande de fréquence correspondant au filtre $H_{Mel}(f; m)$;
- Nous obtenons le Denoised Hearingogram en élevant au carré chaque valeur de $\tilde{S}_{h_m} \forall m = 1, \dots, M$.

L'algorithme se résume sur la Figure 36 et Figure 37 :

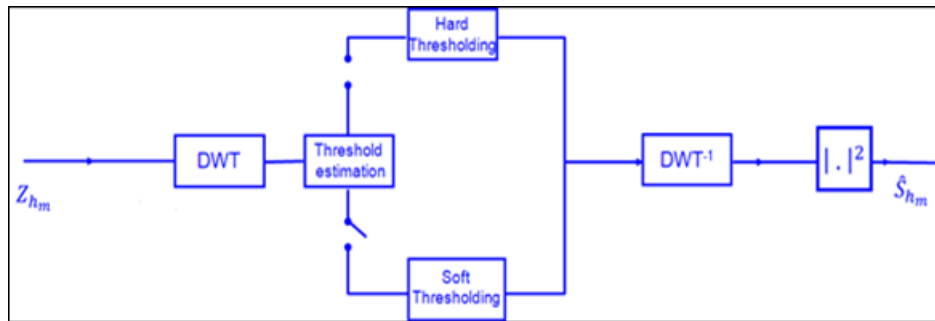


Figure 36: Schéma fonctionnel du principe de débruitage sur une ligne Z_{h_m}

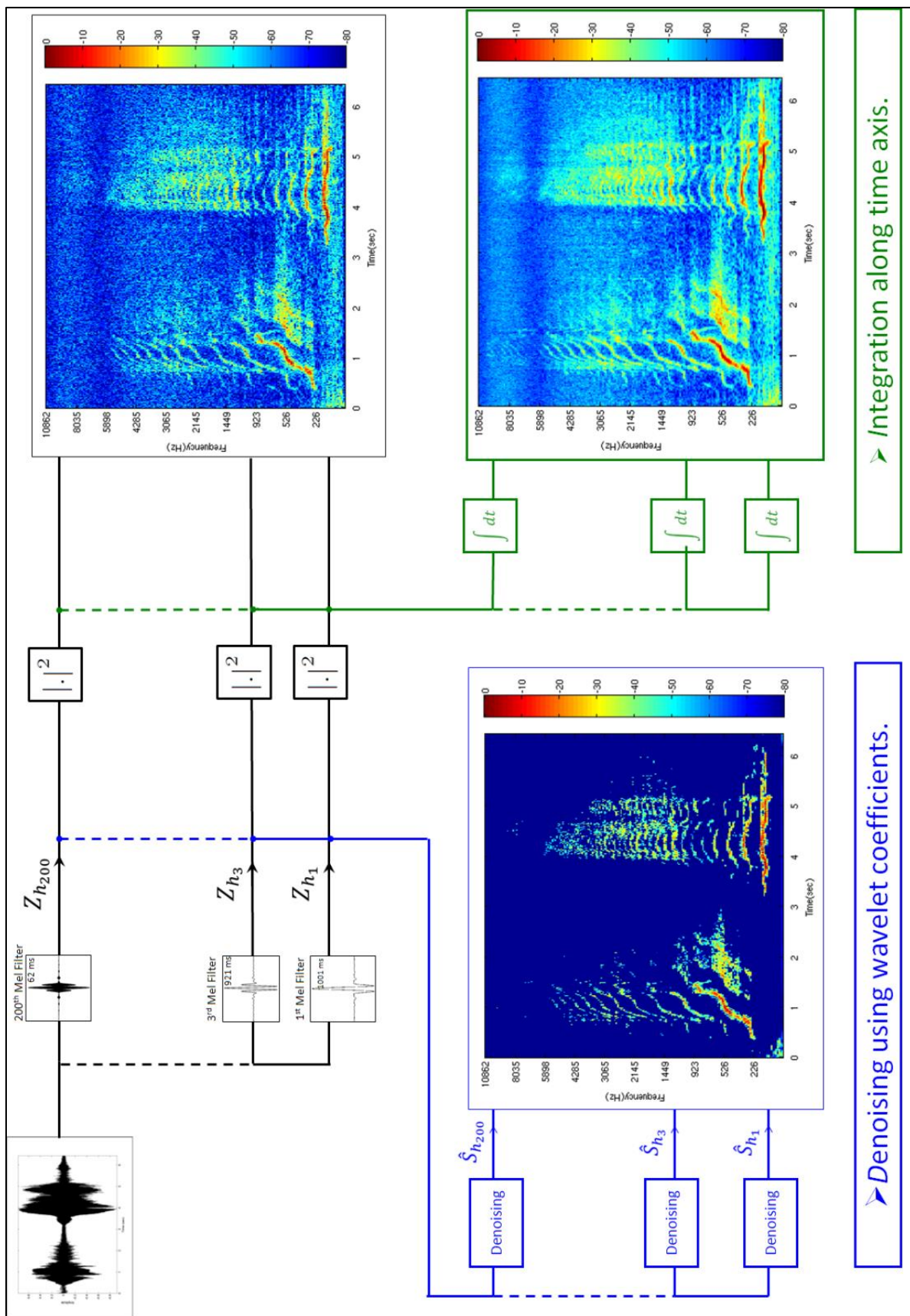


Figure 37: Schéma fonctionnel du Denoised Hearingogram

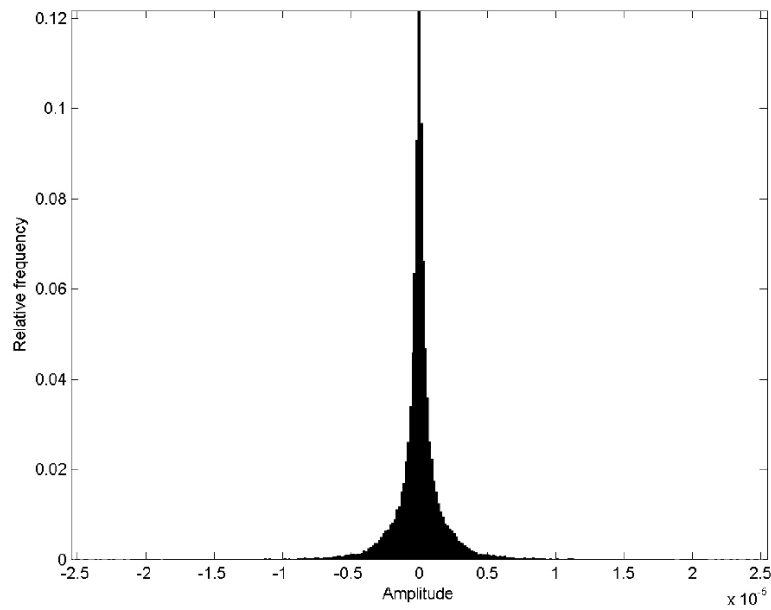


Figure 38: Estimation de la densité de probabilité de Z_{hm} correspondant au $M^{ième}$ de Mel, $M=200$

La méthode proposée s'appuyant sur le caractère gaussien des signaux après filtrage de Mel et ce sous couvert du théorème de la limite centrale, il apparaît nécessaire de valider cette hypothèse. Pour ce faire, nous avons estimé la loi de probabilité de ces signaux filtrés à l'aide de leur histogramme. Nous présentons sur la Figure 38 la loi obtenue pour la $M^{ième}$ ligne. Ce choix est conditionné par le fait que la réponse impulsionnelle associée à cette ligne présentant le moins d'échantillons, elle est la plus susceptible de ne pas respecter ce caractère gaussien. L'analyse du résultat obtenu démontre qu'il n'en est rien, justifiant de fait l'hypothèse de gaussianité.

Nous allons pour la suite présenter trois résultats obtenus sur signaux réels. Une comparaison est effectuée au sein de ces images entre l'*Hearingogram* et le *Denoised Hearingogram*.

Le premier signal est un enregistrement de vocalise de dauphin Risso. Le signal est pollué par du bruit de trafic en basses fréquences et du bruit ambiant. Nous voyons sur l'*Hearingogram* des vocalises du dauphin entre 6 kHz et 13 kHz. De plus nous voyons des clics émis par le dauphin en haute fréquence à partir de 19 kHz. Sur la version débruitée de l'*Hearingogram*, le bruit ambiant ainsi que le bruit de trafic sont très atténués, alors que les vocalises de dauphin ainsi que les clics sont préservés.

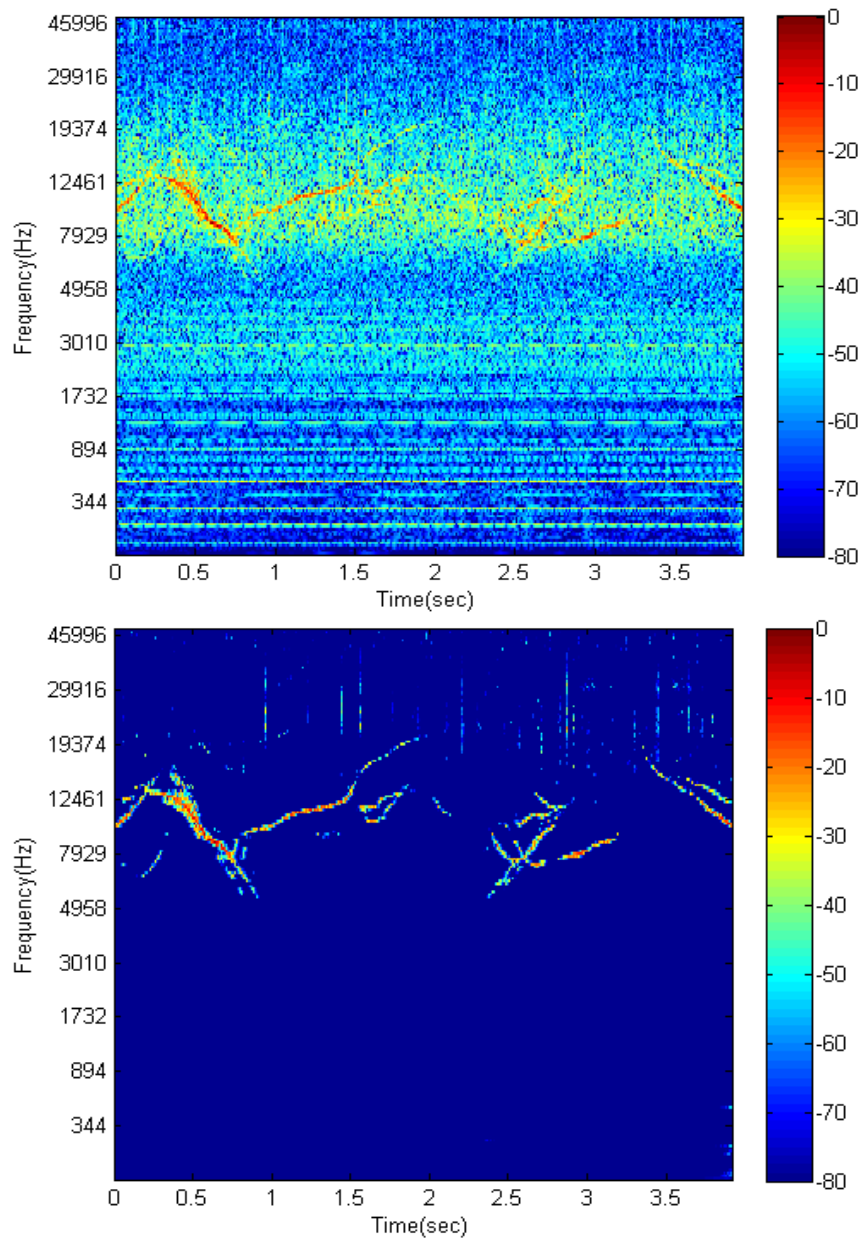


Figure 39: *Hearingogram*(en haut) et *Denoised Hearingogram* (en bas) d'un son de dauphin Risso

Le second signal représente une série de clics d'écholocation, nous voyons que même sur des signaux de durée très brève et avec une bande de fréquence assez large le signal utile est préservé alors que le bruit ambiant est fortement atténué.

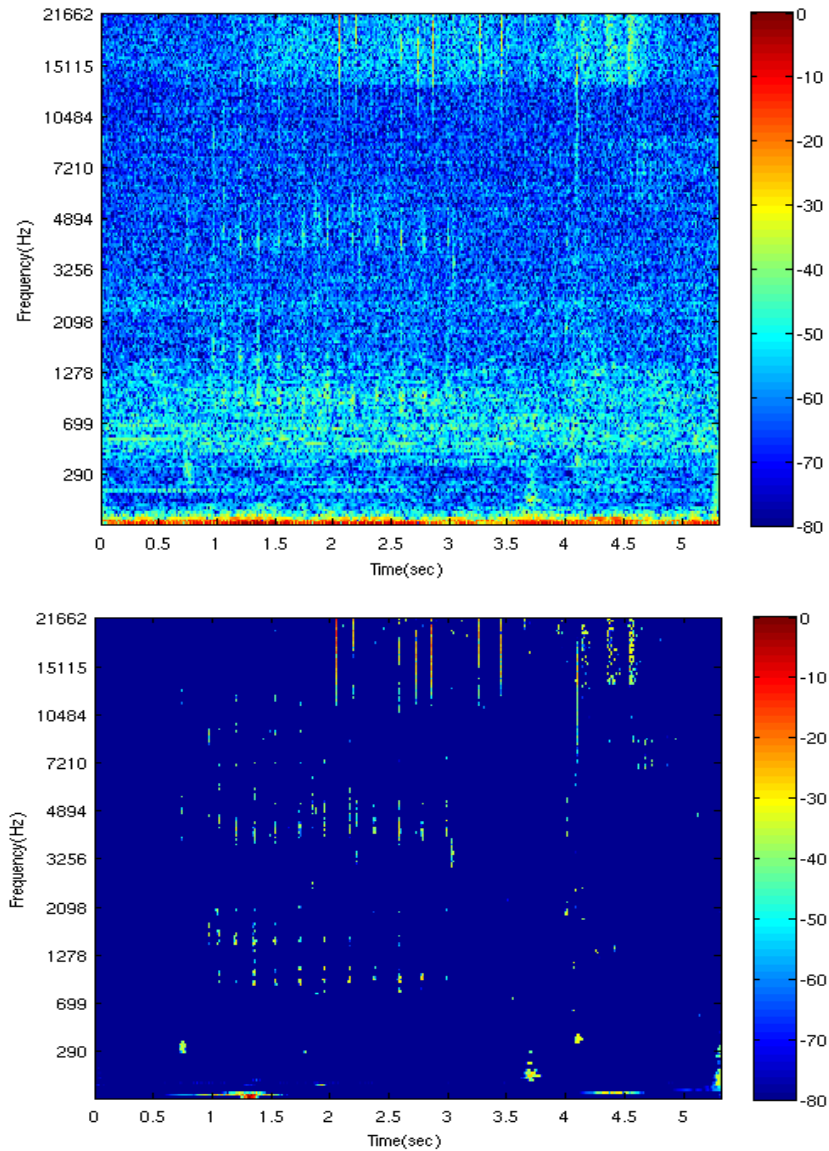


Figure 40: *Hearingogram*(en haut) et *Denoised Hearingogram* (en bas) d'écholocations d'orques

Le dernier signal à tester est celui du dauphin que nous avons présenté dans les résultats de l'*Hearingogram*. Encore une fois le bruit ambiant est fortement atténué et les différentes harmoniques du signal sont préservées. On peut d'ailleurs dire que plusieurs dauphins sont présents dans l'enregistrement, puisque nous voyons nettement sur le *Denoised Hearingogram* que plusieurs harmoniques se croisent entre elles, ce qui ne peut être émis par un seul dauphin.

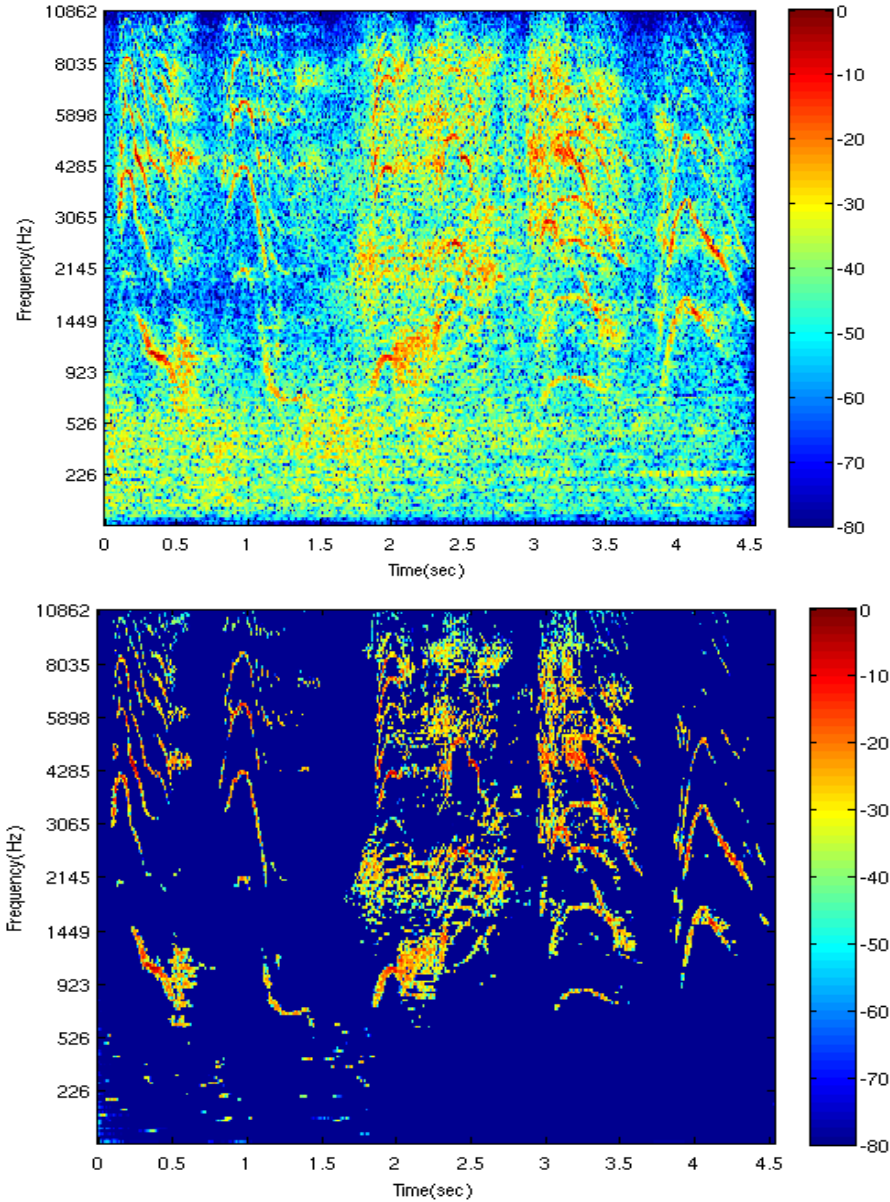


Figure 41: *Hearingogram*(en haut) et *Denoised Hearingogram* (en bas) d'un son de dauphin

Ces résultats révèlent l'efficacité et la force de l'approche. Sur les signaux observés le signal utile est préservé tandis que le bruit est fortement réduit même sans avoir de connaissance *a priori* sur le bruit. Ainsi le processus décrit nécessite peu de réglages de paramètres.

2. Reconstruction du signal utile à partir du Denoised Hearingogram

Si nous considérons $H^{Mel}(f)$ comme étant le filtre associé à tous les filtres de Mel composant le banc de filtres, ce dernier peut être assimilé à un filtre passe-bande:

$$H^{Mel}(f) = \sum_{m=1}^M H_m(f) = 1 \quad \forall f \in [f_1; f_M]$$

III-33

Avec $[mel(f_{min}); f_1[$ et $]f_M; mel(f_{max})]$ comme bandes de transition.

Pour réaliser un filtre passe-tout qui garantirait la conservation de l'énergie du signal filtré, il apparaît nécessaire d'ajouter deux filtres, $H_0(f)$ et $H_{M+1}(f)$ tel que:

$$H(f) = H_0(f) + H^{Mel}(f) + H_{M+1}(f) = 1 \quad \forall f \in \left[0; \frac{F_e}{2}\right] \quad \text{III-34}$$

Après calcul par transformée de Fourier inverse et échantillonnage, voici les réponses impulsionnelles $h_0[k]$ et $h_{M+1}[k]$ associées aux filtres que nous venons de définir:

$$h_0[k] = \frac{f_1 \text{sinc}^2(\pi f_1 k)}{F_e} \quad \text{III-35}$$

$$h_{M+1}[k] = -\frac{1}{F_e} \text{sinc}(f_0 t) (2f_M \cos(\pi f_1 t) + f_1 \text{sinc}(\pi f_1 t)) \dots \quad \text{III-36}$$

$$- \frac{1}{2f_0 F_s} (4f_M^2 \text{sinc}(2f_M t) - F_e^2 \text{sinc}(F_e t))$$

où:

$$\begin{cases} f_0 = \frac{(F_e - 2f_M)}{2} \\ f_1 = \frac{(F_e + 2f_M)}{2} \end{cases}$$

Nous présentons, sur la Figure 42, les réponses fréquentielles associées au banc de filtres, la prise en compte de la totalité de ces derniers conduisant au filtre passe-tout $H(f)$.

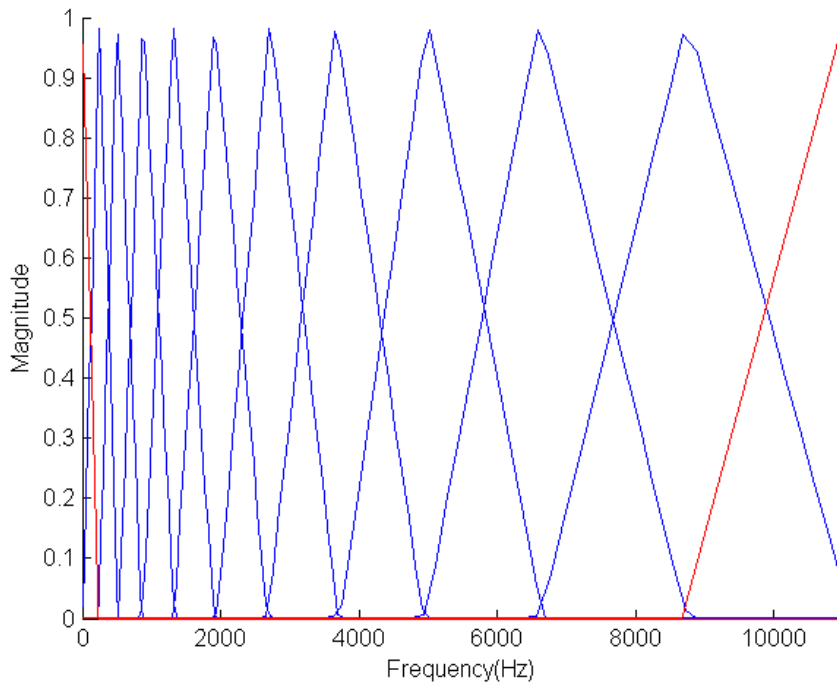


Figure 42: Banc de filtre pour la reconstruction du signal (noir: banc de filtres de Mel; rouge: filtres ajoutés afin d'assurer la conservation de l'énergie)

Dans ces conditions, la réponse impulsionnelle h associée à ce filtre passe bande peut être approximée par:

$$h(t) = TF^{-1}[H(f)] \cong \delta(t) \quad \text{III-37}$$

Où δ représente le Dirac

Ainsi par l'intermédiaire de ce banc de filtre, nous pouvons accéder à l'observation Z à partir de chaque Z_{h_m} grâce à une simple sommation. Le même raisonnement peut être appliqué pour obtenir une approximation du signal utile $\tilde{S}$ à partir des données $\tilde{S}_{h_m}$, on a :

$$\tilde{S} = \sum_{m=0}^{M+1} \tilde{S}_{h_m} \quad \text{III-38}$$

On peut résumer ainsi le processus de réduction du niveau de bruit ainsi:

- Initialisation d'un vecteur $\tilde{S}$ à zéros de la taille de l'observation Z .
- Pour $m = 0 \dots M + 1$:
 - Calcul de la réponse impulsionnelle h_m ;
 - Détermination de Z_{h_m} en réalisant le produit de convolution entre Z et h_m ;
 - Analyse multi-résolution de Z_{h_m} par l'algorithme de Mallat;
 - Seuillage des coefficients obtenus par l'analyse multi-résolution;
 - Construction de $\tilde{S}_{h_m}$ en appliquant le schéma de reconstruction de l'algorithme de Mallat;
 - Construction de façon itérative de $\tilde{S}$:

$$\tilde{S} = \tilde{S} + \tilde{S}_{h_m} \quad \text{III-39}$$

Nous pouvons résumer ce processus grâce au schéma fonctionnel suivant:

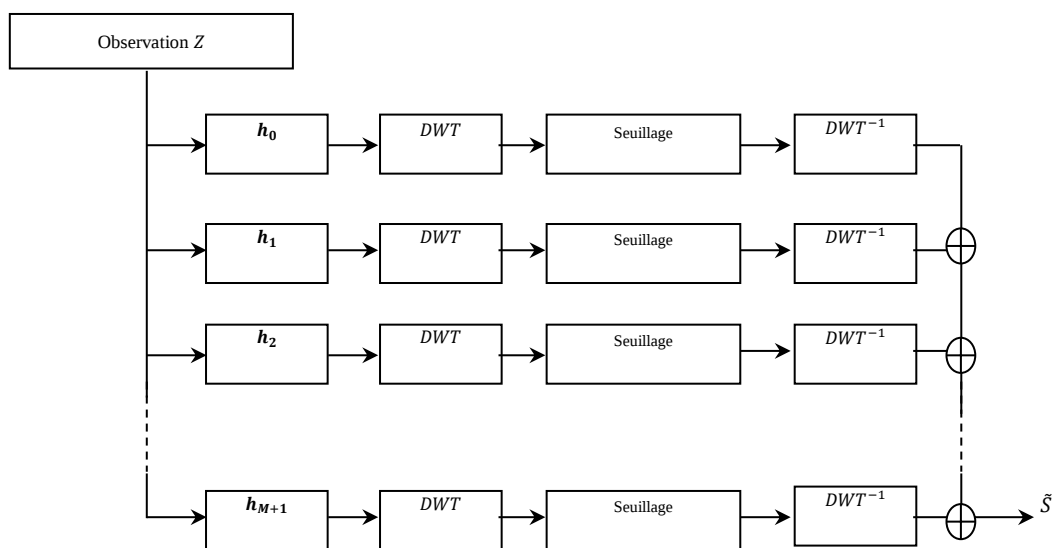


Figure 43: Schéma fonctionnel du processus de débruitage proposé

Nous allons comparer la méthode proposée avec les méthodes présentées. Afin de comparer les méthodes, nous devons donc définir des métriques de comparaison. De plus, nous devons choisir une observation de référence pour comparer les méthodes. Afin de pouvoir quantifier l'apport de l'algorithme en termes de maximisation du rapport signal à bruit, il importe que cette observation de référence soit synthétique, c'est-à-dire issue de la sommation d'un signal utile et de termes perturbateurs.

D. Comparaison des différentes techniques

a. Evaluation quantitative des performances de débruitage

Trois métriques ont été choisies pour comparer les performances des différents algorithmes de réduction du bruit de manière quantitatif.

- L'erreur quadratique moyenne, sans unité (SU) :

III-40

$$EQM = \frac{1}{N} \sum_{n=0}^{N-1} (s[n] - \hat{s}[n])^2$$

N étant la taille du signal en nombre d'échantillons.

L'EQM quantifie la différence qu'il y a entre le signal traité et le signal natif, elle doit être la plus petite possible.

- Le RSB, en dB :

III-41

$$RSB = 10 \log_{10} \left(\frac{\sum_{n=0}^{N-1} s^2[n]}{\sum_{n=0}^{N-1} (s[n] - \hat{s}[n])^2} \right)$$

Plus cette valeur est élevée, meilleure est la qualité du débruitage.

- Le RSB par segment, en dB :

III-42

$$RSB_{seg} = \frac{1}{H} \sum_{q=0}^{H-1} 10 \log_{10} \frac{\sum_{n=0}^{L-1} s^2 \left[\frac{n+qL}{2} \right]}{\sum_{n=0}^{L-1} \left(s \left[\frac{n+qL}{2} \right] - \hat{s} \left[\frac{n+qL}{2} \right] \right)^2}$$

Le RSB par segment est réputé comme un meilleur indicateur de la qualité du débruitage au sens où il est mieux corrélé à la qualité d'écoute que le RSB [45], il doit aussi être maximum.

b. Evaluation quantitative des performances des algorithmes de réduction du bruit

Nous présentons dans cette section une comparaison des différentes techniques exposées. Pour cela deux signaux test sont utilisés, l'un issu d'enregistrement à la mer, l'autre d'un enregistrement d'un morceau de piano. Les signaux seront testés avec différents RSB. Ainsi les performances de chaque technique seront évaluées pour ces différents RSB. La robustesse de ces techniques face à une mauvaise estimation du niveau de bruit sera aussi testée.

1. Signaux de référence

Les deux signaux S_1 et S_2 utilisés pour les calculs de performance, correspondant respectivement à l'enregistrement d'un morceau de piano échantillonné à 11 kHz et à l'enregistrement d'un dauphin échantillonné à 16 kHz , sont présentés ci-dessous:

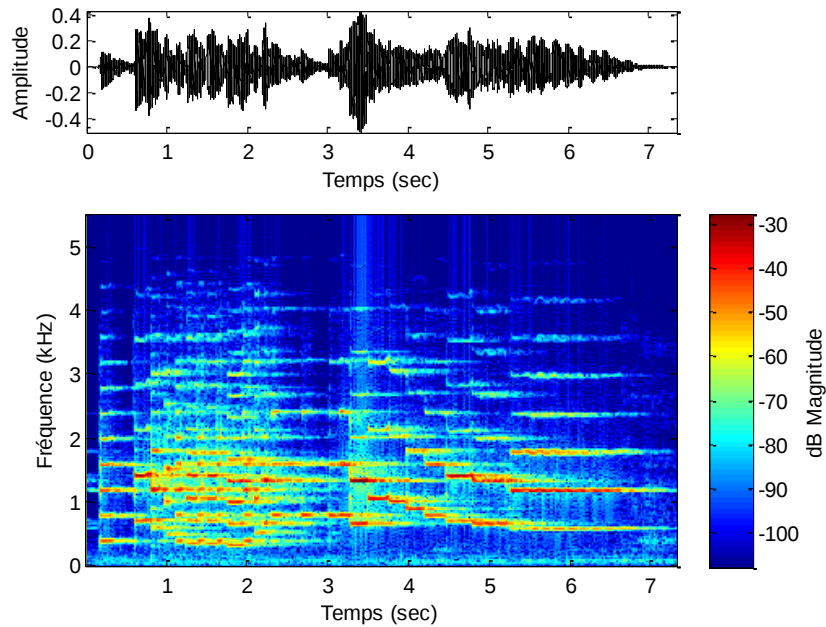


Figure 44: Spectrogramme d'un signal représentant un morceau de piano échantillonné à 11 kHz

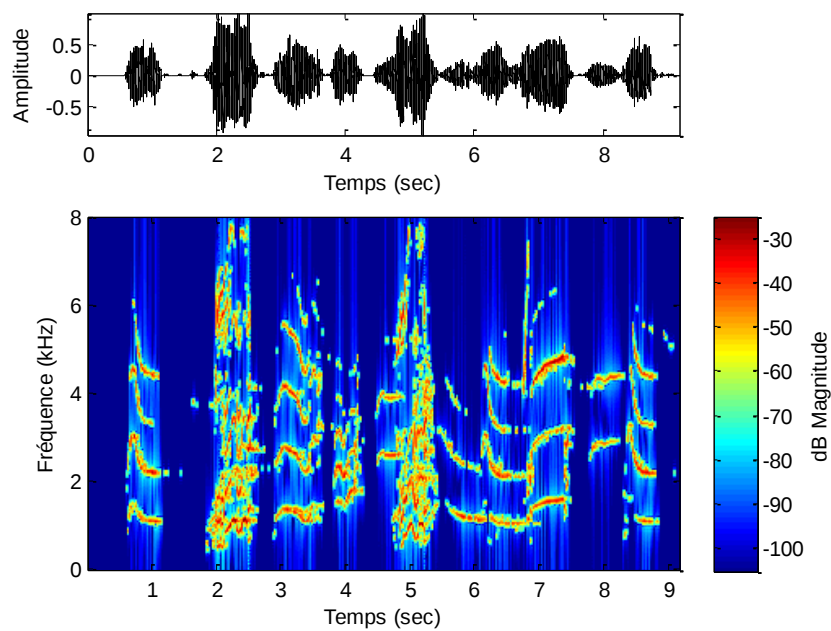


Figure 45: Spectrogramme d'un signal représentant un chant de dauphin, où la fréquence d'échantillonnage est égale à 16 KHz

Cinq observations sont construites correspondant à cinq niveaux de bruit différents. Chaque observation sera notée sous la forme S_{ij} , i désignant le numéro du signal utilisé (1 ou 2) et j le niveau de bruit (de 1 à 5). Le bruit qui est ajouté est un bruit de mer issu de signaux réels, en agissant ainsi nous pouvons contrôler le RSB.

2. Comparaisons des différents estimateurs du RSB a priori

Nous avons vu qu'il existait une multitude d'estimateurs du RSB *a priori* ainsi que de nombreuses règles d'atténuation, il y a donc un très grand nombre de combinaisons possibles. Etant donné que nous ne pourrions pas tester toutes les combinaisons nous sélectionnerons les meilleurs estimateurs et les règles d'atténuation les plus performantes.

Dans cette partie nous allons évaluer les performances des différents estimateurs du RSB *a priori* (DD, TSNR, IPSE et NC). Pour cela nous choisissons une unique règle d'atténuation, celle de Wiener, qui sera associée aux différents estimateurs. Nous calculerons ensuite certains scores produits par chaque combinaison sur chaque signal pour en déduire les meilleurs estimateurs. Les scores calculés seront le RSB par segment et l'EQM. L'estimation des paramètres statistiques du bruit est faite sur le bruit seul avant qu'il soit ajouté au signal.

Ci-dessous les différentes mesures de performance obtenues :

	Scores	DD	TSNR	IPSE	NC
S11	RSBseg (dB)	18.3	18.6	18.9	18.8
	EQM	5.1e-005	4.4e-005	4.1e-005	4.2e-005
S12	RSBseg (dB)	14.4	14.6	14.8	14.8
	EQM	1.3e-004	1.2e-004	1.1e-004	1.1e-004
S13	RSBseg (dB)	10.6	10.9	11.1	11.1
	EQM	3.3e-004	2.9e-004	2.7e-004	2.9e-004
S14	RSBseg (dB)	6.3	6.8	6.9	6.8
	EQM	9.3e-004	8.0e-004	7.4e-004	8.8e-004
S15	RSBseg (dB)	2.8	3.3	3.1	3.7
	EQM	0.0021	0.0021	0.0018	0.0020
S21	RSBseg (dB)	11.7	14.2	13.7	13.4

	EQM	4.9e-004	2.2e-004	2.5e-004	2.8e-004
S22	RSBseg	7.6	9.5	9.1	8.7
	EQM	0.0017	8.7e-004	9.4e-004	0.0012
S23	RSBseg	4.3	5.5	5.3	4.7
	EQM	0.0047	0.0029	0.0031	0.0043
S24	RSBseg	1.6	2.4	2.1	1.8
	EQM	0.01	0.0085	0.0081	0.01
S25	RSBseg	-0.3	0.32	-0.45	0.29
	EQM	0.016	0.017	0.016	0.017

Tableau 2 : Performances des différents algorithmes de réduction du bruit

L'estimateur TSNR produit les meilleures performances globales, tandis que l'estimateur NC donne les meilleurs résultats visuel et auditif, ces deux estimateurs seront donc conservés pour la suite de la comparaison.

3. Comparaison des différentes règles d'atténuation

Cette comparaison est difficile à mettre en œuvre et à interpréter. En effet, les écarts de performance entre les différentes règles sont très dépendants de la méthode d'estimation du RSB *a priori* utilisée. Ainsi, pour une méthode d'estimation donnée une règle d'atténuation va s'avérer bien meilleure qu'une autre tandis qu'en changeant la méthode d'estimation l'écart deviendra négligeable. Il faut aussi garder à l'esprit que certaines règles sont mieux adaptées à certains types de signaux.

Nous allons néanmoins les comparer en utilisant l'estimateur NC pour en déduire celles qui seront capables de produire les meilleures performances. Le choix de l'estimateur NC s'explique par le fait qu'il laisse peu de bruit résiduel, et que c'est souvent sur ce point que les règles d'atténuations agissent, en comblant de manière détournée les faiblesses d'un premier estimateur. Or nous voulons que notre comparaison soit la plus détachée possible de la qualité de l'estimation du RSB *a priori*.

	Scores	Wiener	MMSE_LSA	MAP	JMAP
S11	RSBseg	18.7	18.9	18.7	18.7
	EQM	4.4e-005	4.0e-005	4.3e-005	4.3e-005
S12	RSBseg	14.8	14.9	14.7	14.8
	EQM	1.1e-004	1.1e-004	1.1e-004	1.1e-004
S13	RSBseg	10.9	10.9	10.9	10.9
	EQM	2.9e-004	2.8e-004	2.9e-004	2.9e-004
S14	RSBseg	6.7	6.6	6.6	6.7
	EQM	8.9e-004	7.5e-004	8.8e-004	8.6e-004
S15	RSBseg	3.5	2.5	3.3	3.5
	EQM	0.0022	0.0020	0.0022	0.0022
S21	RSBseg	13.4	13.2	13.8	13.9
	EQM	2.7e-004	2.6e-004	2.4e-004	2.3e-004
S22	RSBseg	8.7	8.7	9.0	9.2
	EQM	0.0012	0.0010	0.0010	9.5e-004
S23	RSBseg	4.7	4.8	4.9	5.2
	EQM	0.0043	0.0034	0.0038	0.0034
S24	RSBseg	1.8	1.5	1.8	1.9
	EQM	0.010	0.0089	0.010	0.0097
S25	RSBseg	0.29	-0.99	0.11	0.35
	EQM	0.017	0.016	0.017	0.017

Tableau 3: Mesures de performances des méthodes de débruitage utilisées, en vert les meilleures performances et en rouge les performances les moins bonnes.

Compte tenu des scores obtenus nous conserverons la règle d'atténuation JMAP. De plus, nous choisissons de conserver la règle d'atténuation de Wiener car c'est la seule qui prend en compte seulement l'estimation du RSB a priori, et qui n'a donc aucun pouvoir correcteur sur cette estimation

(l'estimation du RSB étant la seule information dont elle dispose). Cet effet de correction est visible sur les courbes de gains présentes en amont. On remarque que deux échantillons différents, dont les RSB *a priori* estimés sont égaux mais dont les RSB *a posteriori* sont différents, ne seront pas atténués de la même manière. L'échantillon qui possède le plus grand RSB *a posteriori* sera atténué plus fortement que l'autre. Ceci est l'un des phénomènes qui contribue à réduire le bruit musical⁴ en réduisant plus fortement les pics de bruit car ils auront un RSB *a posteriori* plus grand. En contrepartie si l'on considère que l'estimation du RSB *a priori* est juste alors pour un même RSB l'atténuation devrait être la même, c'est pour cette raison que l'effet correcteur peut être contre-productif.

4. Résultats

Ci-dessous les scores moyens obtenus par chacune des méthodes sélectionnées sur les 10 signaux de test.

<i>Scores moyens</i>	RSB	RSB _{seg}	EQM
NC+ Wien.	10,9780	8,3811	0,0037
NC+JMAP	11,3218	8,5686	0,0035
TSNR+ Wien.	11,4462	8,5922	0,0033
TSNR+JMAP	11,1213	8,4919	0,0035

Tableau 4: Résultats moyens des méthodes de débruitage testées

5. Points forts et points faibles des méthodes comparées

NC+ Wien.	
+	-
Bonnes performances globales	Moyennement adapté au débruitage de signaux à émergences verticales
Faible niveau de bruit résiduel	
Bonne qualité d'écoute	
Assez rapide	

⁴ Résidu de bruit après traitement qui produit effet musical à l'écoute du signal débruité.

Robuste

NC+JMAP	
+	-
Bonnes performances globales	Moyennement adapté au débruitage de signaux à émergences verticales
Faible niveau de bruit résiduel	
Bonne qualité d'écoute	
Assez rapide	
Robuste	

TSNR+ Wien.	
+	-
Bonnes performances globales	Moyennement adapté au débruitage de signaux à émergences verticales
Robuste	Bruit musical
Rapide	Mauvaise qualité d'écoute

TSNR+ JMAP	
+	-
Bonnes performances globales	Moyennement adapté au débruitage de signaux à émergences verticales
Faible niveau de bruit résiduel	Faible bruit musical
Rapide	
Robuste	

6. Comparaison entre le *Denoised Hearingogram* et l'état de l'art

Pour illustrer le traitement proposé, nous avons construit un signal test. Ce dernier est proposé sur la Figure 46 et a été échantillonné à une fréquence égale à 44100 Hz . Ce signal test a été perturbé par un bruit de type gaussien-gaussien, qui est un modèle utilisé pour modéliser le bruit de mer, la densité de probabilité est présenté sur la Figure 47 et l'observation résultante sur la Figure 48. Ce signal a été créé pour représenter différents types de gabarits de signaux que nous pouvons trouver dans le milieu sous-marin. De droite à gauche nous voyons une vocalisation, trois signaux impulsifs, un choc et sa trainée, un signal large bande, trois signaux impulsifs différents des premiers et une vocalisation discontinue.

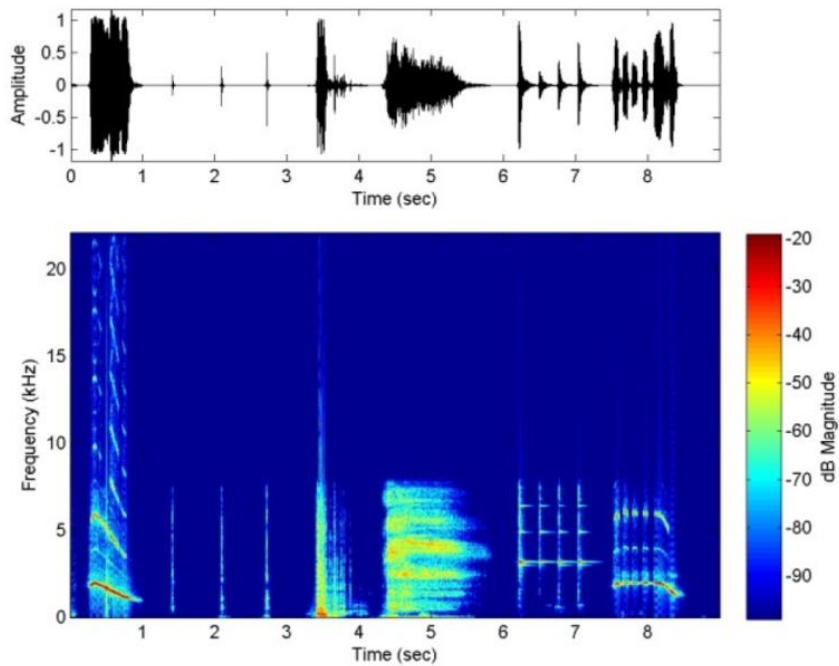


Figure 46: Signal test (représentation temporelle et le spectrogramme associé)

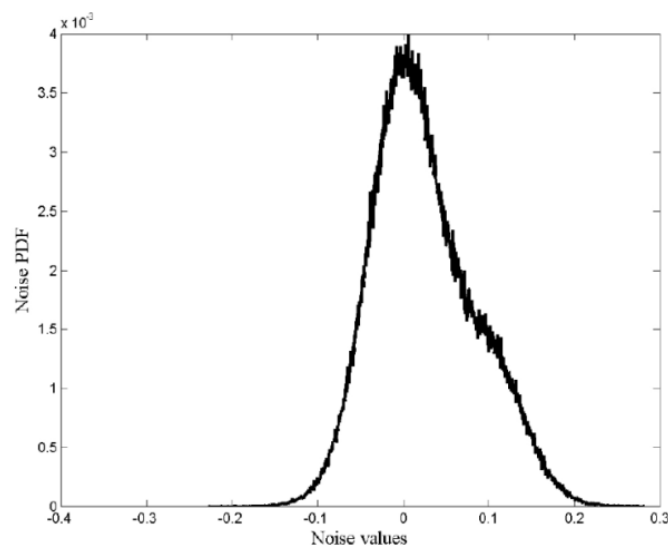


Figure 47: Densité de probabilité du bruit

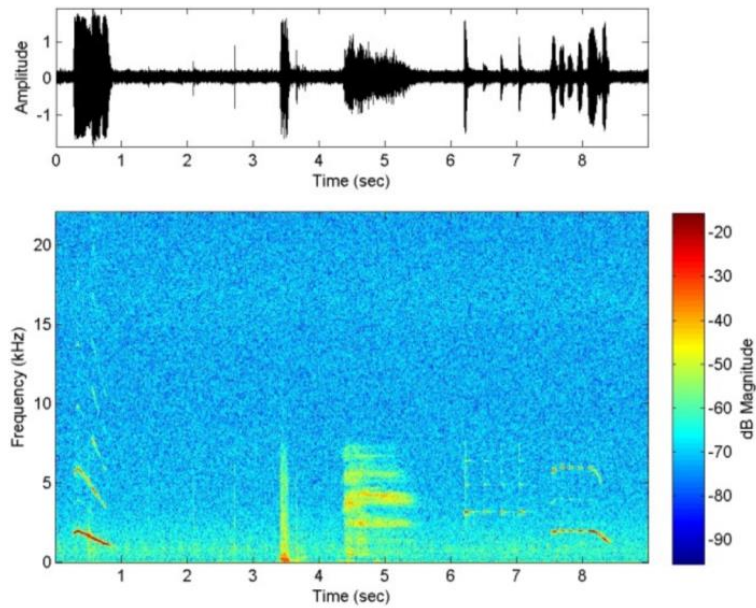


Figure 48: Observation bruitée et son spectrogramme associé

En appliquant l'algorithme Denoised Hearingogram, avec 200 filtres de Mel, sur le signal construit, on obtient le signal présenté sur la Figure 49.

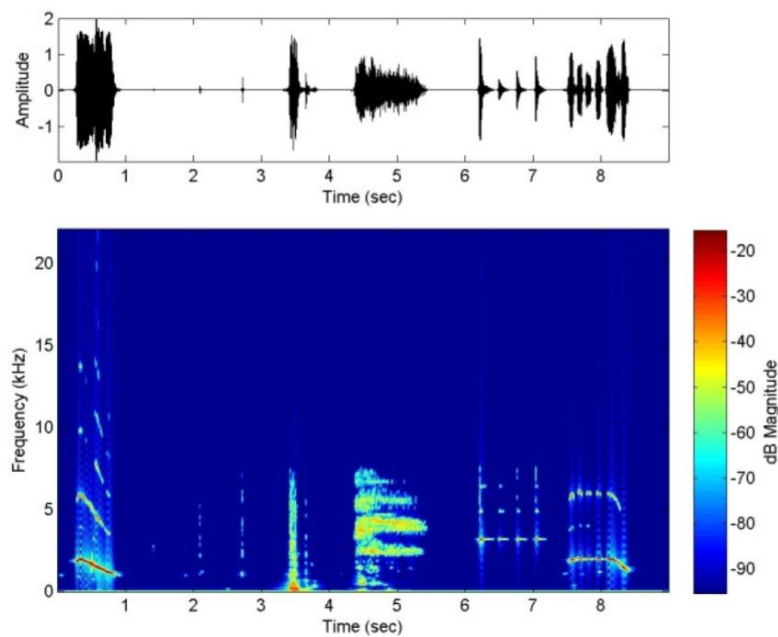


Figure 49: Signal test débruité par la méthode faisant intervenir le *Denoised Hearingogram* (signal temporel et spectrogramme associé)

Ce résultat montre l'efficacité de l'algorithme Denoised Hearingogram, tous les types de signaux sont bien conservés, à l'exception du premier signal impulsif qui avait un RSB trop défavorable pour être restauré, dans le même temps le bruit de fond est fortement atténué.

Dans le but d'évaluer notre méthode, nous la comparons avec deux algorithmes présentés précédemment. Le premier est l'estimateur DD associé à la règle d'atténuation MMSE-LA et le

second algorithme est basé sur la règle d'atténuation de Wiener associé avec l'estimateur non causal de I.Cohen [38]. Les algorithmes sont paramétrés de la manière suivante :

- $\alpha = 0.98$ pour l'estimateur DD
- $\alpha = 0.9, \alpha_{bis} = 0.98, \beta = 2$ et $\xi_{min} = -25dB$ pour le second algorithme

Nous avons calculé le spectrogramme avec une fenêtre de Hanning de 11 ms avec 50% de recouvrement. La puissance du bruit est estimée avec le MAD comme dans le *Denoised Hearingogram*.

Nous présentons, dans le tableau, les mesures de performance que nous avons définie dans la partie précédente pour les différents types de signaux :

	Observation	LSA-DD	Wien-NC	DH
RSB	16.02	28.20	29.10	20.30
RSBseg	-3.90	6.94	13.05	8.50
EQM	44*10 ⁻⁴	2*10 ⁻⁴	2*10 ⁻⁴	2*10 ⁻⁴

Tableau 5: Résultat sur les vocalisations

	Observation	LSA-DD	Wien-NC	DH
RSB	-10.47	2.09	3.20	3.30
RSBseg	-43.41	-26.70	-10.14	-8.40
EQM	43*10 ⁻⁴	2*10 ⁻⁴	2*10 ⁻⁴	2*10 ⁻⁴

Tableau 6: Résultat sur les signaux impulsifs

	Observation	LSA-DD	Wien-NC	DH
RSB	9.40	15.70	16.40	14.50
RSBseg	-12.90	-1.59	3.15	2.30
EQM	41*10 ⁻⁴	10*10 ⁻⁴	8*10 ⁻⁴	15*10 ⁻⁴

Tableau 7: Résultat sur le choc et sa trainée

	Observation	LSA-DD	Wien-NC	DH
RSB	7.94	12.17	12.24	9.16
RSBseg	-1.43	5.78	7.45	5.69
EQM	41*10 ⁻⁴	16*10 ⁻⁴	15*10 ⁻⁴	28*10 ⁻⁴

Tableau 8: Résultat sur le signal à bande large

	Observation	LSA-DD	Wien-NC	DH
RSB	7.69	20.35	22.04	18.13
RSBseg	-11.68	3.32	8.62	6.02
EQM	43*10 ⁻⁴	2*10 ⁻⁴	2*10 ⁻⁴	6*10 ⁻⁴

Tableau 9: Résultat sur signaux impulsif, autre type

	Observation	LSA-DD	Wien-NC	DH
RSB	13.07	24.77	25.85	19.28
RSBseg	-9.50	4.07	11.28	10.14
EQM	42*10 ⁻⁴	3*10 ⁻⁴	2*10 ⁻⁴	19*10 ⁻⁴

Tableau 10: Résultats sur vocalisations discontinues

	Observation	LSA-DD	Wien-NC	DH
RSB	-42.07	-24.61	-11.12	-24.28
RSBseg	-43.58	-25.48	-5.96	6.9
EQM	42*10 ⁻⁴	7.5*10 ⁻⁵	3.4*10 ⁻⁶	1.5*10 ⁻⁷

Tableau 11: Résultat sur le bruit

Une analyse de ces résultats révèle que la performance moyenne de notre processus est meilleure que le LSA-DD, mais moins bonne que la méthode Wien-NC, sauf pour les cas de signaux impulsifs et les cas bruités. Il est à noter que le *Denoised Hearingogram* (DH) nécessite peu de réglages de paramètres comparés à la méthode Wien-NC où quatre paramètres sont ajustés selon l'observation considérée. C'est un inconvénient notable pour une étape de réduction du bruit dans un système opérationnel qui est complètement automatique.

Pour illustrer le processus sur des données réelles, nous avons choisi de l'appliquer sur l'enregistrement contenant des écholocations d'orque présenté précédemment. Ce signal est présenté sur la Figure 50, avec le spectrogramme associé.

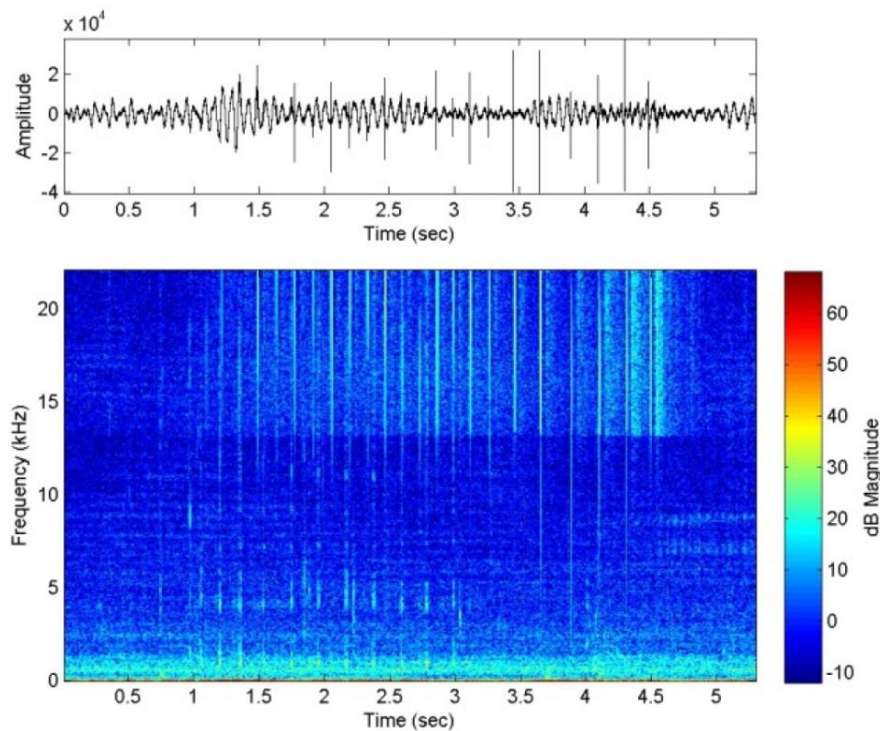


Figure 50: Signal temporel d'écholocations d'orque et son spectrogramme

Bien que ces écholocations correspondent à des signaux très impulsifs, elles sont bien préservées après réduction du bruit, et permettent donc une détection et une interprétation automatique plus facile. Nous avons utilisé les mêmes paramètres que dans l'expérience précédente.

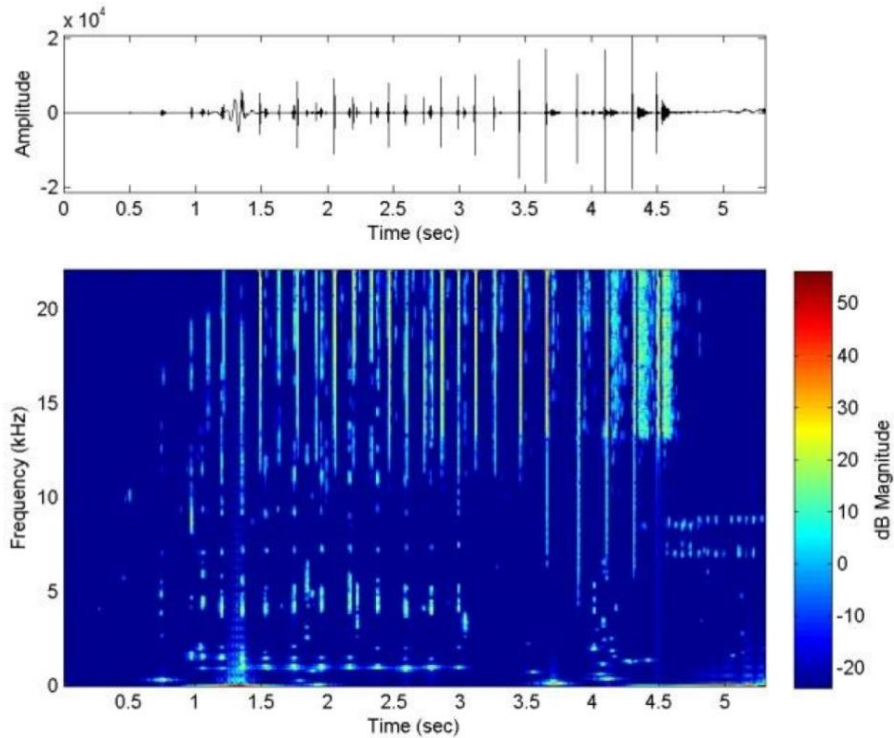


Figure 51 : Signal d'echolocations d'orque débruité par la méthode faisant intervenir le *Denoised Hearingogram* (signal temporel et spectrogramme associé)

Cette partie a permis de faire un état de l'art des techniques de représentation temps- fréquence. De plus, une nouvelle représentation basée sur la physiologie humaine a été exposé, l'*Hearingogram* et sa version où le bruit a été réduit le *Denoised Hearingogram*.

Après avoir abordé la représentation des signaux, nous allons voir comment l'identification de signaux acoustiques sous-marins peut être réalisée.

Chapitre 3: Applications à la détection sous-marine dans le contexte du sonar passif

I. Reconnaissance des signaux acoustiques sous-marins

La faculté d'apprendre est essentielle à l'être humain pour reconnaître une voix, une personne, un objet... On peut naturellement distinguer deux types d'apprentissage :

- L'apprentissage par cœur qui consiste à mémoriser des informations telles qu'elles sont,
- L'apprentissage par généralisation où l'on se construit un modèle à partir d'exemples qui nous permettra de reconnaître de nouveaux exemples.

L'apprentissage automatique est une tentative de comprendre et de reproduire le deuxième type d'apprentissage dans des systèmes automatiques.

Dans le contexte de la reconnaissance des signaux sous-marins, nous disposons d'une base de données. Dans cette dernière est représenté chaque événement détecté à l'aide de descripteurs (voir partie 4). La base de données se présente donc comme un tableau de données de taille $N \times D$, donc une matrice où chaque ligne représentera la description d'un exemple (soit N exemples) et chaque colonne un descripteur (donc D descripteurs au total) donné pour tous les exemples d'apprentissage, on parle aussi d'individu. Cette procédure est aussi utilisée dans plusieurs domaines, c'est notamment le cas pour la reconnaissance des codes-barres, ou de l'ADN.

Le but d'un système automatique de reconnaissance acoustique sous-marine est d'attribuer une classe à chaque exemple que nous traitons, en s'appuyant sur les descripteurs disponibles. Deux situations sont possibles :

- La classification supervisée, c'est le cas où les classes sont connues *a priori*. Ainsi dans une base de données, pour chaque exemple il faut renseigner la classe à laquelle il appartient parmi les classes utilisées dans l'apprentissage. L'opération d'étiquetage de la base d'apprentissage nécessite souvent l'aide d'un expert. La définition des classes n'est pas un problème simple dans le cas des signaux acoustiques sous-marins. La difficulté étant de trouver le bon niveau de granularité dans le choix des classes.
- La classification non-supervisée, dans ce cas-là nous ne connaissons pas *a priori* les classes à affecter aux individus. Le but ici est de trouver une organisation du nuage de points correspondant aux individus, en K régions, appelées clusters, on parle donc de clustering. Nous ferons appel à ces méthodes dans certains contextes précis, que nous présenterons plus tard.

Nous utiliserons pour l'identification de la plupart des signaux acoustiques sous-marins une classification supervisée. L'une des grandes difficultés lors de la classification des signaux sous-marins est la diversité des signaux. En effet, deux signaux provenant d'une même source acoustique, peuvent avoir des signatures différentes. Plusieurs raisons expliquent ce phénomène:

- Le milieu marin est délimité par le fond et la surface de l'eau, qui constituent des interfaces permanentes sur lesquelles vont se réfléchir les ondes sonores. Ceci va entraîner une multiplication des échos, comme on peut le voir sur la Figure 52, qui peuvent perturber la réception notamment s'il s'agit de transmissions de données. Plus le nombre de réflexions sera important, plus l'intensité acoustique diminuera et le temps de trajet augmentera. La trainée d'échos est considérée comme parasite pour le récepteur. Le signal le plus rapide sera évidemment celui empruntant le trajet direct (dans la mesure où ce dernier existe).

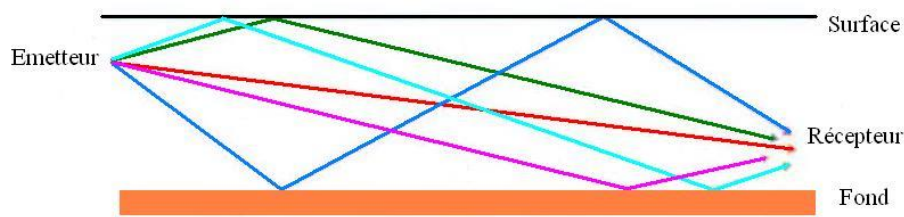


Figure 52: Représentation des multi-trajets pour une onde sonore

- Les pertes sont aussi dues à ce que l'on appelle la divergence géométrique. Les pertes, qui se font notamment au niveau de l'intensité de l'onde sonore, sont dues à un effet géométrique de divergence et à l'absorption de l'énergie acoustique par le milieu de propagation lui-même. Au cours de la propagation, l'énergie acoustique émise se conserve. Cependant, elle se répartit sur une surface qui augmente au cours de la progression de l'onde. L'intensité acoustique diminue proportionnellement à l'inverse de cette surface, c'est le phénomène de perte par divergence géométrique.

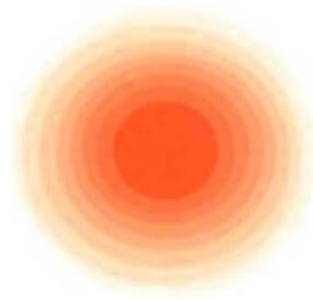


Figure 53: Divergence sphérique

Sur la Figure 53, les surfaces, pour un milieu infini homogène et une source omnidirectionnelle de faibles dimensions, sont des sphères de rayons de plus en plus importants. L'intensité diminue donc au cours de la propagation.

- Un autre facteur impactant est l'absorption acoustique. Le milieu et la fréquence en sont les principaux paramètres. Le milieu est dissipatif et absorbe une partie de l'énergie de l'onde (à cause de la viscosité du milieu en l'occurrence). Le coefficient d'amortissement évolue fortement avec la fréquence et ses ordres de grandeurs sont très variables.

Voici une représentation du coefficient d'amortissement de l'eau de mer en fonction de la fréquence, à plusieurs températures et pour une salinité de 0.35 %. [46]

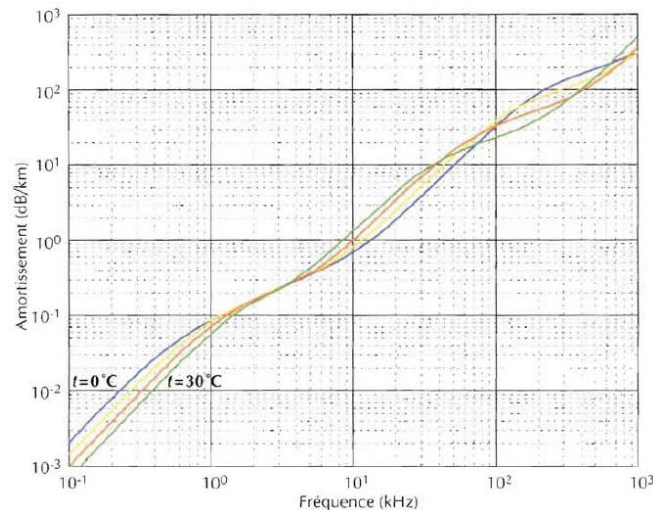


Figure 54: Amortissement du son dans l'eau de mer en fonction de la fréquence [46]

La Figure 54 permet de comprendre que plus la fréquence est élevée, plus l'amortissement sera important. A 10 kHz, le coefficient d'amortissement est de l'ordre de 1 dB/km, ce qui permet d'atteindre des distances de quelques dizaines de kilomètres, alors que lorsque la fréquence est aux alentours du Mégahertz, le coefficient est proche de 500 dB/km, ce qui est considérable. À ces fréquences, les systèmes de détection sont limités à moins de 100 mètres de portée.

Afin de prévoir les pertes de propagation et les performances des systèmes acoustiques sous-marins, un calcul peut être réalisé en première approche pour une dispersion sphérique. Il permet d'estimer le niveau de pertes en décibels en fonction de la distance et du coefficient d'amortissement :

$$PT = 20\log(R) + \alpha R$$

avec R la distance, en kilomètres, parcourue et α le coefficient d'amortissement.

- La réverbération acoustique est un autre facteur potentiellement déformant du signal émis. Une partie de l'énergie acoustique se propage vers le récepteur en suivant d'autres chemins que les rayons propres, sur les interfaces où dans le volume d'eau et avec des temps de propagation généralement supérieurs. A ce titre, on parle de réverbération « de surface », « de fond » et « de volume ». L'effet est une « traine temporelle » prenant la forme d'un signal large bande, observable sur la représentation temps-fréquence des signaux les plus énergétiques. Si l'œil averti d'un analyste sait associer le signal source et sa réverbération, un système automatique devra de même résoudre ce phénomène naturel.

Il existe encore d'autres facteurs qui rendent donc le problème de reconnaissance des signaux acoustiques sous-marins non trivial. Pour cette présentation, nous nous sommes limités aux principales difficultés. Pour ces raisons la définition des descripteurs est un problème crucial, ainsi que la définition des classes et de la base d'apprentissage. Nous reviendrons sur ces différents points dans la dernière partie.

II. Principe d'un système de reconnaissance automatique

La classification automatique vise à assigner une classe à un objet sans l'intervention d'un opérateur. Dans notre problème nous cherchons donc à étiqueter un signal qui a été extrait automatiquement. Le principe général d'un système de reconnaissance inclut deux étapes :

- Une étape d'apprentissage, où l'on définit les frontières permettant de séparer les différentes classes.
- Une étape de test où l'on évalue la performance de notre classifieur.

La phase d'apprentissage est composée de 3 étapes :

- La définition de la classe attribuée à chaque exemple (ou individu) ;
- Un calcul de descripteurs, permettant dans notre application d'identifier un signal détecté par un vecteur de descripteurs ;
- Une sélection des descripteurs les plus pertinents au niveau de la classification ;
- L'apprentissage des classifieurs, à partir des attributs sélectionnés on obtiendra des frontières de décision, ainsi on pourra classer les nouveaux exemples.

Lors de l'étape de test nous avons besoin de calculer les descripteurs sélectionnés lors de la phase d'apprentissage et décider de l'appartenance du signal à une classe en utilisant les frontières calculées. Ces différentes étapes sont résumées sur la Figure 55.

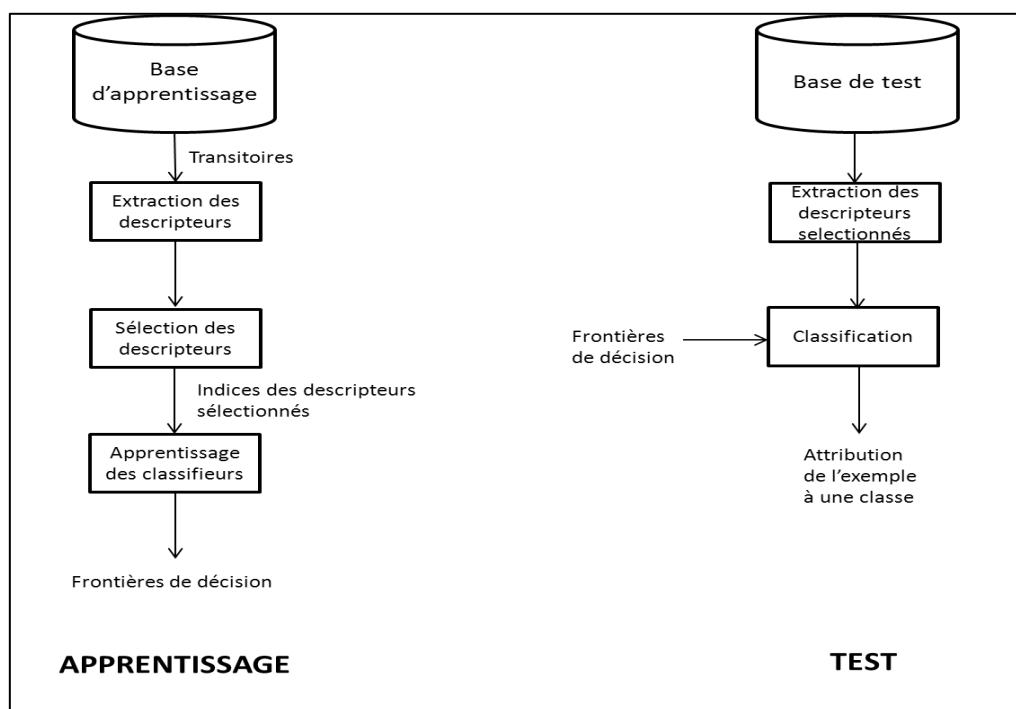


Figure 55: Description des étapes de la classification supervisée

III. Classification non supervisée

A. Principe

La classification non supervisée consiste à regrouper de manière automatique les données en cluster sans réaliser d'apprentissage ou quelque traitement *a priori* sur les données. L'hypothèse centrale qui régit ces types d'algorithmes est qu'il existe bien des clusters au sein de nos données, autrement dit qu'elles sont séparables. Ainsi, des échantillons d'exemples très proches doivent appartenir au même groupe, et donc avoir la même classe. Inversement, une frontière de séparation de deux groupes doit se trouver dans une zone dans laquelle peu d'individus sont présents. Dans le cas extrême où les données sont réparties de manière uniforme, on voit que l'hypothèse de séparabilité est mise en défaut et, de fait, un algorithme de clustering ne donnera aucun résultat satisfaisant sur un tel type de données. Schématiquement, nous pouvons ainsi dire que la vocation du clustering est de regrouper ce qui se ressemble, il s'agit d'un mécanisme de coalescence.

Plusieurs familles d'algorithmes de clustering existent :

- hiérarchiques: produisant un ensemble de partitions imbriquées appelés dendrogrammes;
- partitives: le résultat est une partition en un nombre fixé de groupes, donné ou calculé automatiquement par l'algorithme;
- floues: permet d'attribuer des valeurs de probabilités d'appartenance des individus à chaque groupe.

Nous utilisons des algorithmes de classification non supervisée en intégrant l'information sur la nature des signaux à classifier. Plus particulièrement, sous la mer il y a une grande variété de signaux. Notamment les espèces biologiques qui émettent une grande quantité de signaux, par exemple pour se localiser ou pour chasser ils utilisent des clics d'écholocation. Lors de la partie 1 il a été vu que cette étude se réalisait dans un contexte d'identification d'un pavé temps-fréquence à la fois, c'est-à-dire que nous tentons de classer chaque pavé individuellement après segmentation sans tenir compte des pavés voisins. Le problème est qu'avec les signaux impulsifs, c'est souvent une information de contexte qui permet de les reconnaître, par exemple lorsqu'un mammifère marin chasse il émet un train de signaux impulsifs avec une certaine période, et c'est l'ensemble du train de clics qui permet l'identification de chaque signal impulsif composant le train. Ainsi ces algorithmes permettent d'associer la même étiquette à tous les signaux identifiés comme appartenant au même groupe par les différents algorithmes. De plus, nous verrons que ces trains de clics sont souvent caractérisés par la durée inter-clics, c'est-à-dire que la distance entre deux clics est assez régulière. Nous allons donc aussi extraire cette information afin de réunir ces clics en famille.

Nous allons dresser un état de l'art des différentes techniques de clustering, puis nous présenterons des scores permettant de juger de manière quantitative la qualité de la classification non-supervisée réalisée. Enfin, nous présenterons une technique permettant d'extraire la durée inter-clics d'une famille de clics, qui se nomme le temps-rythme.

B. Etat de l'art

1. Classification ascendante hiérarchique

Comme les autres méthodes de l'analyse des données, dont elle fait partie, la classification a pour but d'obtenir une représentation schématique simple d'un tableau rectangulaire de données dont les colonnes, suivant l'usage, sont des descripteurs de l'ensemble des observations, placées en lignes.

L'objectif le plus simple d'une classification est de répartir la population des individus en groupes d'observations homogènes, chaque groupe étant bien différencié des autres. Le plus souvent, cependant, cet objectif est plus complexe ; on veut, en général, obtenir des sections à l'intérieur des groupes principaux, puis des subdivisions plus petites de ces sections, et ainsi de suite. En bref, on désire avoir une hiérarchie, c'est-à-dire une suite de partitions emboîtées, de plus en plus fines, sur l'ensemble des observations initiales.

Une telle hiérarchie peut avantageusement être résumée par un arbre hiérarchique qu'on appelle dendrogramme. Sur la Figure 56, les nœuds (m, n, p, q) symbolisent les diverses subdivisions de la population ; les éléments de ces subdivisions étant les objets (a, b, c, d, e) placés à l'extrémité inférieure des branches qui leur sont reliées.

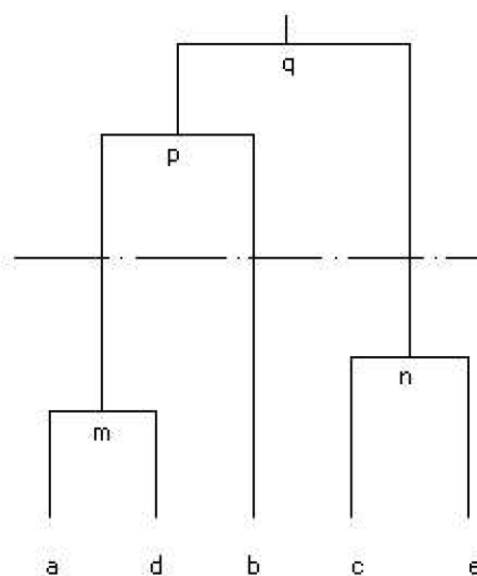


Figure 56: Exemple de dendrogramme portant sur cinq objets a, b, c, d, e . Les points m, n, p, q sont les nœuds de l'arbre. Le trait horizontal mixte indique un niveau de troncature définissant une partition en trois classes.

Le niveau des nœuds est sensé indiquer un degré de similitude entre les objets correspondants. Ainsi, les objets a et d se ressemblent plus que les objets c et e . Enfin nous pouvons remarquer que si nous coupons l'arbre à un niveau intermédiaire entre n et p , on obtient une partition en trois classes de l'ensemble étudié, à savoir $\{a, d\}, \{b\}, \{c, e\}$. En faisant varier ce niveau de troncature on obtient les diverses partitions constituant la hiérarchie.

Les différentes mesures mènent naturellement à différentes déclinaisons de la méthode de construction des partitions imbriquées. Notons $C_i \forall i = 1 \dots K$ une classe parmi les classes potentielles et $x_i \forall i = 1 \dots N$ un individu parmi les individus de la base. Parmi les différentes mesures, nous trouvons :

- Lien simple (saut minimum ou *single linkage*) :

$$D(C_i, C_j) = \min_{x_i \in C_i, x_j \in C_j} d(x_i, x_j),$$

III-1

la distance entre les deux amas est alors la distance la plus courte entre deux individus de ces amas ;

- Lien complet (saut maximum ou *complete linkage*) :

$$D(C_i, C_j) = \max_{x_i \in C_i, x_j \in C_j} d(x_i, x_j), \quad \text{III-2}$$

la distance entre les deux amas est alors la distance la plus grande entre deux individus de ces amas ;

- Lien moyen (*average linkage*) :

$$D(C_i, C_j) = d(c_i, c_j), \quad \text{III-3}$$

où $c_i = \frac{1}{\text{card}(C_i)} \sum_{x_i \in C_i} x_i$ et $c_j = \frac{1}{\text{card}(C_j)} \sum_{x_j \in C_j} x_j$ sont les moyennes respectives des amas C_i et C_j .

La distance entre les deux amas correspond dans ce cas à la distance entre les barycentres respectifs de ceux-ci.

Il existe d'autres types de lien tel que le type de Ward. Dans [47] on peut trouver un tableau nous résumant les différents types de lien.

L'inconvénient de cette méthode est qu'elle n'utilise que des critères d'optimisation locaux qui n'induisent pas forcément une optimisation globale des résultats. De plus, les regroupements sont définitifs, donc ne pouvons pas appliquer des post-traitements à cet algorithme. Par contre, on n'a pas besoin de connaître forcément le nombre de clusters désirés et il n'y a pas de fonction d'initialisation contrairement à l'algorithme K-means (décrit ci-dessous).

2. K-means

Cette méthode, présentée dans [48] est de type discriminative. Elle est l'approche la plus connue et utilisée dans les différentes communautés scientifiques utilisant le clustering. Le principe est intuitif : étant donnés la distribution des individus dans l'espace de description et un nombre fixé K de groupes, l'objectif est de minimiser la dispersion des individus relativement à un ensemble d'individus représentatifs de ces groupes.

Les individus x_i sont représentés par un vecteur de $\mathbb{R}^D$, et l'ensemble des individus est alors décrit par une matrice $X \in \mathbb{R}^{N \times D}$. Du point de vue du modèle, l'algorithme des K-means est basé sur la minimisation d'une erreur quadratique relativement à ces prototypes qui se formalise par :

$$\min_{c, C} Q_{KM}(c, C) = \min_{c, C} \sum_{k=1}^K \sum_{x_i \in C_k} \|x_i - c_k\|_2^2 \quad \text{III-4}$$

où c_k est le prototype du groupe C_k .

Le résultat est une partition de l'espace des données en clusters séparés. La qualité de la solution dépend fortement de l'initialisation. De plus, la sensibilité à l'initialisation est d'autant plus grande que la dimensionnalité des données est grande.

3. Spectral Clustering

Cette méthode [49] est une autre approche de type partitionnement, elle permet de prendre en compte la structure naturelle des données. En réalité, il s'agit d'un algorithme de type K-means appliqué à l'ensemble des individus projetés dans un sous-espace particulier. Cet espace de projection de dimension n_k est construit de telle sorte que des paquets d'individus proches s'agrègent de façon séparable dans chacune des dimensions.

4. DBScan

L'objectif [50] est explicitement de capturer les zones de fortes densités, définissant ainsi un groupe. Il s'agit d'une approche exclusivement algorithmique qui se fonde sur une modélisation particulière du concept de zone dense, et qui parcourt l'ensemble des individus afin de déterminer si ceux-ci appartiennent ou non à une telle zone.

5. Fuzzy K-Means

Les Fuzzy K-Means [51] sont une généralisation des K-means se basant sur des éléments de la théorie des ensembles flous. Le principe est toujours de minimiser la dispersion des individus relativement aux prototypes, mais en pondérant cette dernière par le degré d'appartenance de l'individu au groupe. Du point de vue du critère objectif, on présente les K-means flous comme la minimisation du critère de l'erreur quadratique semblable à l'algorithme K-means, mais évaluée pour chaque individu relativement à l'ensemble des prototypes :

$$\begin{aligned} \min_{c,u} Q_{FKM}(c,u) &= \min_{c,u} \sum_{k=1}^K \sum_{x_i \in C_k} u_{ik}^\beta \|x_i - c_k\|_2^2 \\ \text{s. c } \sum_{k=1}^K u_{ik} &= 1 \quad \forall x_i \in X \\ u_{ik} &\geq 0 \quad \forall x_i \in X, \forall k \in [1 \dots n_k] \end{aligned} \quad \text{III-5}$$

où $\beta \geq 1$ est un paramètre fixé dans l'objectif et c_k est le prototype du groupe C_k . $u = \{u_{ik}\}$ est l'ensemble des degrés d'appartenance des individus aux groupes. En particulier, u_{ik} indique le degré d'appartenance de l'individu x_i au groupe C_k .

C. Mesure de performance pour la classification non supervisée

La validation des groupes formés suite à l'application des algorithmes de classification non supervisée est un problème non-trivial. Les algorithmes de clustering tentent de trouver le meilleur modèle

séparant les données selon un nombre fixé de cluster. Cela ne veut pas dire que l'on a trouvé le meilleur modèle selon les données en notre possession car il se peut que le nombre de clusters ne soit pas le bon dans la réalité.

La méthode retenue pour s'approcher au mieux du nombre optimal de clusters est de tester le clustering pour plusieurs valeurs et de calculer des indicateurs de performance, puis finalement garder le nombre de clusters qui maximise les valeurs des différents indicateurs.

Plusieurs indicateurs de performance ont été proposés dans la littérature pour la classification non supervisée :

- Le coefficient de partitionnement (PC) qui mesure le « recouvrement » entre deux groupes. Il est défini par [51] comme :

$$PC(c) = \frac{1}{N} \sum_{i=1}^c \sum_{j=1}^N (u_{ij})^2 \quad \text{III-6}$$

Le nombre optimal de cluster est considéré être celui qui entraîne le maximum de cette valeur.

- L'entropie de classification (CE), définie par [52] qui est une mesure du caractère flou de la partition, il est similaire au précédent :

$$CE(c) = -\frac{1}{N} \sum_{i=1}^c \sum_{j=1}^N u_{ij} \log(u_{ij}) \quad \text{III-7}$$

Le nombre optimal de cluster est celui entraînant la valeur maximum.

- Index de partition (SC) : [52] est le ratio de la somme de compacité et de la séparation des clusters. C'est une somme de la mesure de validité pour chaque cluster normalisé par la cardinalité floue de chaque cluster :

$$SC(c) = \sum_{i=1}^c \frac{\sum_{j=1}^N (u_{ij})^2 \|x_j - v_i\|^2}{N_i \sum_{k=1}^c \|v_k - v_i\|^2} \quad \text{III-8}$$

Plus la valeur de SC est basse, meilleur est la partition.

- Index de séparation (S) : [52] Au contraire de l'index de partition (SC), l'index de séparation utilise la distance minimum de séparation pour la validité de la partition :

$$S(c) = \sum_{i=1}^c \frac{\sum_{j=1}^N (u_{ij})^2 \|x_j - v_i\|^2}{N \min_{i,k} \sum_{k=1}^c \|v_k - v_i\|^2} \quad \text{III-9}$$

Comme précédemment plus la valeur de S est basse, meilleur sera la partition.

- Xie and Beni's Index (XB) : [53] l'objectif ici est de quantifier le ratio de la variation total au sein d'un cluster et la séparation des différents clusters :

$$XB(c) = \sum_{i=1}^c \frac{\sum_{j=1}^N (u_{ij})^2 \|x_j - v_i\|^2}{N \min_{i,k} \sum_{k=1}^c \|x_k - v_i\|^2} \quad \text{III-10}$$

Le nombre optimal de clusters doit minimiser la valeur de ce score.

La combinaison de toutes ces mesures permet d'obtenir une bonne approximation du nombre de groupes à utiliser pour discriminer au mieux les données.

Nous allons maintenant présenter une nouvelle méthode nous permettant de regrouper tous les signaux d'un train de clics au sein d'un même groupe. Cette technique est nommée le temps-rythme.

D. Temps-rythme

On a vu précédemment que l'activité biologique représente une grande partie des signaux sous-marins. De plus, certains types de bateaux produisent avec leurs hélices un phénomène que l'on appelle la cavitation. Ce phénomène est un signal composé de plusieurs signaux impulsifs comme précédemment. Ces impulsions sont le plus souvent émises en trains rythmés, et ce rythme peut nous aider grandement dans l'identification de signaux sous-marins. La difficulté de la classification tient au fait que plusieurs trains de clics d'origines différentes peuvent se mélanger au sein du même signal, ce qui complique bien évidemment l'interprétation finale.

C'est pour cela que nous utiliserons un algorithme qui nous permet d'obtenir un plan temps-rythme [54], [55].

- Tout d'abord, on modélise le train de clics par le temps d'arrivée de chaque clic, appelé TOA⁵. Un train de N clics est décrit alors par une somme d'impulsions de type Dirac:

$$g(t) = \sum_{n=0}^{N-1} \delta(t - t_n) \quad \text{III-11}$$

avec δ la distribution de Dirac et t_n le temps d'arrivée du $n^{ième}$ clic.

- Le rythme, appelé ICI⁶, est analysé par une fonction d'autocorrélation à suppressions d'harmoniques (ASH), définie ainsi:

$$D(\tau) = \int_{-\infty}^{+\infty} g(t)g(t - \tau)e^{\frac{2i\pi t}{\tau}} dt \quad \text{III-12}$$

On substitue $g(t)$ par son expression et on obtient après calcul:

⁵ Time of arrival
⁶ Inter-click interval

$$D(\tau) = \sum_{n=1}^{N-1} \sum_{m=0}^{n-1} \delta(\tau - (t_n - t_m)) e^{\frac{2i\pi\tau}{\tau}} dt \quad \text{III-13}$$

- Le résultat se présente sous la forme d'une carte montrant l'évolution du rythme des trains de clics en fonction du temps, pour réaliser cette carte nous calculons l'ASH dans des fenêtres glissantes le long du temps. Ainsi cette transformée est définie par:

$$D(t, \tau) = \int_{s \in W(t, \tau)} g(s) g(s + \tau) e^{\frac{2i\pi s}{\tau}} ds \quad \text{III-14}$$

Où $W(t, \tau) = \left[t - \frac{\mu\tau}{2}, t + \frac{\mu\tau}{2} \right]$ représente la fenêtre glissante et μ représente un nombre réel positif. Ainsi le résultat de l'analyse temps-rythme s'exprime sous la forme d'une image représentant le spectre des ICI en fonction du temps.

Sur la Figure 57 nous pouvons voir la représentation temps-rythme d'un signal synthétique contenant 3 familles :

- Le premier train de clics est constitué de 17 chocs avec un ICI égal à 0.5s. Le TOA du premier membre du train de clics est 1s et le dernier est 9s.
- Le deuxième train de clics est constitué de 7 chocs avec un ICI égal à 1.3s. Le TOA du premier membre du train de clics est 6.1s et le dernier est 13.9s.
- Le troisième train de clics est constitué de 7 chocs avec un ICI égal à 2.4s. Le TOA du premier membre du train de clics est 14s et le dernier est 28.4s.

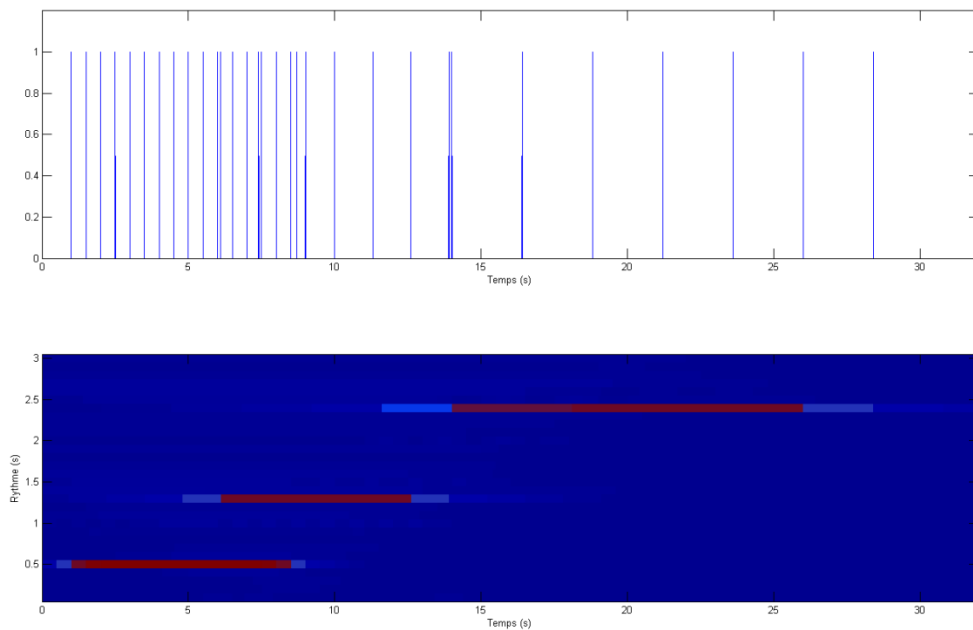


Figure 57 : Représentation temps-rythme (en bas) avec le modèle associé (en haut) sur un signal simulé contenant 3 trains de clics différents

Les résultats de cette technique sont très bons sur cet exemple particulier, en effet nous distinguons clairement les trois trains de clics différents sur notre plan temps-rythme. Ainsi cette information d'ICI est très utile pour réaliser de la classification non supervisée afin de regrouper les clics au sein d'une même famille. Cependant un point faible de cette méthode est que le paramétrage n'est pas universel, ainsi dans un système automatique d'identification et donc sur signaux réels les résultats sont parfois moins bons.

Après avoir présenté différents algorithmes de classification non-supervisée, les principes de la classification sont exposés dans la partie suivante et plus particulièrement ceux des machines à vecteur support.

IV. Classification supervisée

A. Introduction

L'apprentissage supervisé concerne le cas où les données d'entrée sont organisées en classes connues à l'avance. C'est le cas de notre problème où nous disposons d'observations, appelées exemples d'apprentissage, qui sont associées à des classes de signaux. L'objectif de la classification supervisée est principalement de définir des règles, qui peuvent être de différentes natures, permettant d'associer des observations à des classes prédéfinies à l'aide d'un expert. Cette classification est faite à partir de variables qualitatives ou quantitatives caractérisant ces observations.

Ces techniques ont été utilisées dans beaucoup de domaines :

- Reconnaissance de formes : chiffres manuscrits, visages ...
- Catégorisation de textes : classification d'e-mails, de pages web...
- Diagnostic médical : Evaluation des risques de cancer, détection d'arythmie cardiaque.

Plusieurs méthodes de classification supervisée existent dans la littérature, tel que les réseaux de neurones ou le modèle de mélange gaussien. Ces dernières années les machines à vecteur support ont connu un succès évident. Une présentation de cette méthode est faite dans la section suivante.

B. Machines à vecteurs supports (SVM)

1. Le choix des SVM

Nous avons choisi d'utiliser un classifieur de type SVM car après étude bibliographique et pratique de différents classifieurs supervisés, il s'est avéré que ce classifieur est le mieux adapté à notre problématique, ne faisant aucune hypothèse approximative sur la forme des densités de probabilité des données, contrairement aux autres approches.

2. Principe et calcul des SVM

Les SVM sont par définition des classificateurs binaires qui visent à séparer les exemples de chaque classe C_1 ou C_2 au moyen d'un hyperplan choisi de manière à maximiser la marge de séparation entre les deux classes, où seuls certains exemples d'apprentissage participent au calcul de la frontière de décision. La Figure 58 illustre le principe des SVM.

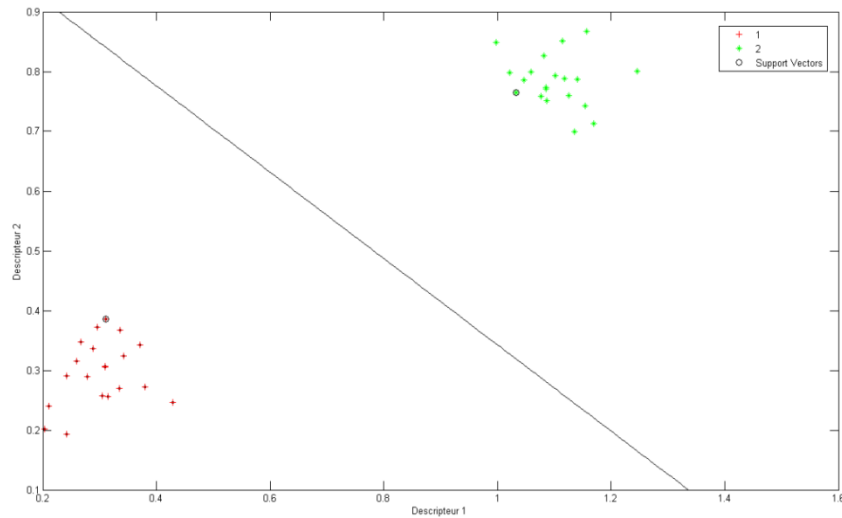


Figure 58: Illustration du principe des SVM

Formulation primale du problème SVM

Soit $\vec{x} \in \mathbb{R}^D$, on suppose l'existence de la loi inconnue $P(\vec{x}, y)$ à valeurs dans $(\mathbb{R}^D, \{-1, 1\})$. L'objectif est de construire un estimateur de la fonction de décision idéale :

$$D: \mathbb{R}^D \rightarrow \{-1, 1\}$$

qui minimise pour toutes les observations $\vec{x}$ la probabilité d'erreur $P(D(\vec{x}) \neq y|\vec{x})$.

Nous nous plaçons tout d'abord, dans le cas où les données sont séparables et linéaires. Il existe alors une fonction de décision linéaire, de la forme :

$$D(\vec{x}) = \text{signe}(f(\vec{x})) \quad \text{IV-1}$$

avec $f(\vec{x}) = \vec{v}^T \vec{x} + a$

avec $\vec{v} \in \mathbb{R}^D$ et $a \in \mathbb{R}$ classant correctement toutes les observations de l'ensemble d'apprentissage $\{D(\vec{x}_i) = y_i, i \in [1, N]\}$.

Le but est de trouver un hyperplan qui va maximiser la marge afin d'augmenter les probabilités d'une bonnes classifications des nouveaux exemples. L'étape concernant la maximisation de la marge peut-être vue ainsi :

$$\max_{\vec{v}, a} \left(\min_{i \in [1, n]} d(\vec{x}_i, \Delta(\vec{v}, a)) \right) \quad \text{IV-2}$$

Où la marge m est égale à $\min_{i \in [1, n]} d(x_i, \Delta(\vec{v}, a))$. On peut alors réécrire ce problème comme un problème d'optimisation sous contraintes :

$$\begin{cases} \max_{\vec{v}, a} m \\ \text{avec } \min_{i \in [1, n]} \frac{|\vec{v}^T \vec{x}_i + a|}{\|\vec{v}\|} \geq m \end{cases} \quad \text{IV-3}$$

Ce problème est mal posé car si $(\vec{v}, a)$ est solution alors $(\overline{kv}, ka)$, avec $\mathbb{R} > 0$, l'est aussi. Voilà pourquoi nous effectuons le changement de variables suivant :

$\vec{w} = \frac{\vec{v}}{m\|\vec{v}\|}$ et $b = \frac{a}{m\|\vec{v}\|}$, ainsi le problème se réécrit de la manière suivante :

$$\begin{cases} \max_{\vec{w}, b} m = \frac{1}{\|\vec{w}\|} \\ \text{avec } y_i(\vec{w}^T \vec{x}_i + b) \geq 1 \quad ; \forall i = 1 \dots n \end{cases} \quad \text{IV-4}$$

Ainsi on formule le problème des SVM de la façon suivante :

Un séparateur à vaste marge linéaire est un discriminateur de la forme :

$$D(\vec{x}) = \text{signe}(\vec{w}^T \vec{x} + b) \quad \text{IV-5}$$

où $\vec{w} \in \mathbb{R}^D$ et $b \in \mathbb{R}$ sont donnés par la résolution du problème suivant :

$$\begin{cases} \min_{\vec{w}, b} \left(\frac{1}{2} \|\vec{w}\|^2 \right) \\ \text{avec } y_i(\vec{w}^T \vec{x}_i + b) \geq 1 \quad ; i = 1 \dots n \end{cases} \quad \text{IV-6}$$

Il est à noter qu'on utilise le carré de la norme pour faciliter la résolution du problème, le coefficient $\frac{1}{2}$ est présent pour la même raison.

Dans le cas des données linéairement séparables on peut solutionner directement ce problème avec des algorithmes de résolution tel que Gauss-Seidel [56]. Cependant il est intéressant de passer par la formulation duale de ce problème car cette dernière fait apparaître une matrice de Gram [57], qui est une matrice représentant la distance entre chaque exemple d'apprentissage, ce qui nous permettra d'introduire plus facilement l'utilisation des noyaux.

Formulation duale du problème SVM

Pour résoudre un problème d'optimisation convexe sous contraintes affines, on utilise le Lagrangien. Dans le cas des SVM, le Lagrangien s'écrit :

$$L(\vec{w}, b, \vec{\alpha}) = \frac{1}{2} \|\vec{w}\|^2 - \sum_{i=1}^n \alpha_i (y_i (\vec{w} \vec{x}_i + b) - 1) \quad \text{IV-7}$$

où les α_i sont les multiplicateurs de Lagrange associés aux contraintes.

On peut exprimer à partir de là les conditions d'optimalité de Karush, Kuhn Tucker (KKT) [58] [59] qui permettront de caractériser la solution du problème primal $(\vec{w}^*, b^*)$ et les multiplicateurs de Lagrange $\vec{\alpha}^*$:

- Stationnarité : $\frac{\partial L}{\partial \vec{w}} = 0 \rightarrow \vec{w}^* = \sum_{i=1}^n \alpha_i^* y_i \vec{x}_i$ IV-8

- Complémentarité : $\alpha_i^* (y_i (\vec{w}^* \vec{x}_i + b^*) - 1) = 0 \quad i = 1, \dots, n$ IV-9

- Admissibilité primale : $(\vec{w}^* \vec{x}_i + b^*) \geq 1 \quad i = 1, \dots, n$ IV-10

- Admissibilité duale $\alpha_i^* \geq 0 \quad i = 1, \dots, n$ IV-11

Les conditions de complémentarité permettent de définir l'ensemble v_s des indices des contraintes qui à l'optimum sont les multiplicateurs de Lagrange α_i^* qui sont strictement positifs :

$$v_s = \{i \text{ tel que } y_i (\vec{w}^* \vec{x}_i + b^*) = 1 \mid i = 1, \dots, n\} \quad \text{IV-12}$$

On parlera pour ces indices de contraintes saturées ou actives, alors que pour les indices ne vérifiant pas cette contrainte, leur multiplicateur de Lagrange α_i^* sera égal à 0.

Ce qui signifie donc que seuls les indices correspondant aux contraintes saturées participent au calcul de la solution, on parle alors de vecteurs supports, car seuls ces vecteurs interviennent dans la construction de l'hyperplan optimal. Les autres données n'interviennent pas dans le calcul de l'hyperplan optimal. En d'autres termes si on enlève les individus n'étant pas des vecteurs supports de nos données d'apprentissage, l'hyperplan optimal reste inchangé.

De ce qui précède, le problème dual des SVM dans le cas de données linéairement séparables s'écrit :

$$\begin{cases} \max_{\vec{w}, b, \vec{\alpha}} \left(\frac{1}{2} \|\vec{w}\|^2 - \sum_{i=1}^n \alpha_i (y_i (\vec{w}^T \vec{x}_i + b) - 1) \right) \\ \vec{w} - \sum_{i=1}^n \alpha_i y_i \vec{x}_i = 0 ; \sum_{i=1}^n \alpha_i y_i = 0 ; \alpha_i \geq 0 \end{cases} \quad i = 1, \dots, n \quad \text{IV-13}$$

Après élimination de la variable primale $\vec{w}$, on a la formulation duale du problème SVM :

$$\begin{cases} \min_{\vec{\alpha}} \left(\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_j \alpha_i y_i y_j \vec{x}_j^T \vec{x}_i - \sum_{i=1}^n \alpha_i \right) \\ \text{avec } \sum_{i=1}^n \alpha_i y_i = 0 \text{ et } \alpha_i \geq 0 \quad i = 1, \dots, n \end{cases} \quad \text{IV-14}$$

Au final l'hyperplan solution des SVM s'écrit :

$$f(\vec{x}) = \sum_{i \in v_s} \alpha_i y_i \langle \vec{x}_i, \vec{x} \rangle + b \quad \text{IV-15}$$

On remarque comme précédemment que la frontière de décision est calculée uniquement sur quelques vecteurs qui sont les vecteurs supports.

Ces deux formulations du problème SVM sont équivalentes, diverses méthodes existent pour résoudre les deux formulations du problème. La question reste encore ouverte sur le choix de la formulation à résoudre. Cependant, le problème majeur de cette formulation est que les données réelles ne sont que très rarement linéairement séparables, en tout cas elles ne le sont pas dans notre domaine d'étude.

Pour modéliser le fait que les données ne sont pas linéairement séparables, on insère dans le problème primal des variables d'écart positives ξ_i pour que les contraintes deviennent moins rigides. Ainsi le problème primal se réécrit :

$$\begin{cases} \min_{\vec{w}, b} \left(\frac{1}{2} \|\vec{w}\|^2 + C \sum_{i=1}^n \xi_i \right) \\ \text{avec } y_i (\vec{w}^T \vec{x}_i + b) \geq 1 - \xi_i \quad ; i = 1 \dots n \\ \xi_i \geq 0 \end{cases} \quad \text{IV-16}$$

où $\vec{\xi} = [\xi_1, \dots, \xi_n]^T$ et $C > 0$

C est un coefficient de pénalisation des contraintes permettant de contrôler le compromis entre le fait de maximiser la marge, au prix d'accepter certaines erreurs lors de l'apprentissage et éviter le sur-apprentissage, et minimiser les erreurs de classification commises sur l'ensemble de

l'apprentissage. On parle alors de classification à marge souple. Notons qu'il est souvent préférable de tolérer certaines erreurs, au bénéfice d'une marge plus grande car certains individus de nos données d'apprentissage peuvent être des données aberrantes.

En tenant compte des variables d'écart et de la constante C que nous avons introduites, la formulation du problème dual s'obtient de la même manière que précédemment en écrivant le Lagrangien et en exprimant les conditions de Karush, Kuhn et Tucker (KKT). Après calcul on obtient la formulation suivante :

$$\left\{ \begin{array}{l} \max_{\vec{\alpha}} \left(\sum_{i=1}^n \alpha_i - \sum_{k=1}^n \sum_{l=1}^n \alpha_k \alpha_l y_k y_l \vec{x}_k^T \vec{x}_l \right) \\ \alpha_i \geq 0 \\ 0 \leq \alpha_i \leq C \text{ et } \sum_{i=1}^n \alpha_i y_i = 0 \end{array} \right. \quad \forall i = 1, \dots, n \quad \text{IV-17}$$

Finalement on obtient que l'hyperplan solution des SVM s'écrit :

$$f(\vec{x}) = \sum_{i \in v_S, \vec{i} \in BSV} \alpha_i y_i \langle \vec{x}_i, \vec{x} \rangle + b \quad \text{IV-18}$$

Où $BSV = \{i \text{ tel que } \alpha_i = C \mid i = 1, \dots, n\}$ et v_S est décrit par IV-12.

Par rapport au cas où les données sont linéairement séparables, une contrainte a été rajoutée sur les α_i , en effet ils sont maintenant bornés supérieurement par C qui représente l'influence maximale que peut avoir un exemple d'apprentissage sur le calcul de la frontière optimale. En réécrivant les conditions KKT, on retrouve la même solution pour $\vec{w}$, à la différence près qu'il n'y a pas que les vecteurs supports qui participent à la solution il y a aussi les vecteurs supports se trouvant à l'intérieur de la marge, appelés erreurs de marge, qui sont associés aux multiplicateurs de Lagrange qui sont tels que $\alpha_i = C$. Ces vecteurs sont appelés BSV (Bounded Support Vectors). On peut déduire aussi des conditions KKT que les variables d'écart ξ_i sont nulles pour tous les vecteurs supports associés à des multiplicateurs α_i tels que $0 < \alpha_i < C$.

Pour plus de détails concernant les calculs, nous invitons le lecteur à consulter [60], [61].

3. SVM non-linéaires

a. Principe

Malgré une base théorique solide, les SVM restent toutefois fortement limitées par la restriction aux séparateurs linéaires. Il est en effet rare que des données réelles soient providentiellement réparties de chaque côté d'un hyperplan. L'idée pour réaliser cette opération est de projeter les données dans un espace de plus grande dimension. Ainsi dans cet espace, les données auront une plus grande

probabilité d'être linéairement séparable. Pour illustrer ce principe, prenons l'exemple du ou exclusif, appelé XOR, qui est une fonction logique décrite de la manière suivante :

Descripteur 1	Descripteur 2	Label
0	0	0
0	1	1
1	0	1
1	1	0

Tableau 12 : Table de vérité du Ou exclusif

Le constat fait sur la Figure 59 est que les données ne sont pas linéairement séparables, aucune droite ne pourra séparer les données.

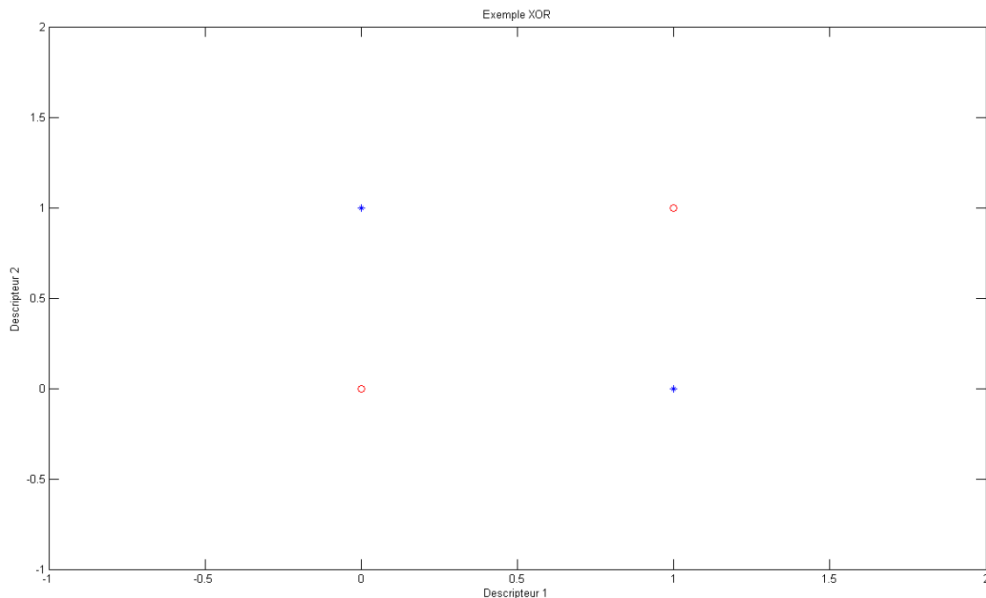


Figure 59: Représentation graphique de l'exemple du XOR

Effectuons la transformation suivante sur les données :

$$(x_1, x_2) = (x_1, x_2, x_1 \cdot x_2)$$

IV-19

Le tableau précédent se transforme ainsi :

Descripteur 1	Descripteur 2	Descripteur 3	Label
0	0	0	0
0	1	0	1
1	0	0	1
1	1	1	0

Tableau 13 : Table de vérité du Ou exclusif après application d'une transformation

Ainsi le problème devient linéairement séparable car il existe un plan qui sépare les données de façon linéaire. Ainsi dans la transformation effectuée dans l'exemple ci-dessus nous avons utilisé un noyau, la définition de cette notion est abordée dans la prochaine partie.

b. Noyaux

Dans le cas des SVM à marge souple, le fait d'admettre des éléments mal classés, ne peut pas donner toujours une bonne généralisation pour un hyperplan même si ce dernier est optimisé. Nous pouvons observer ceci sur la Figure 60 où la frontière de décision idéale serait plutôt de forme circulaire.

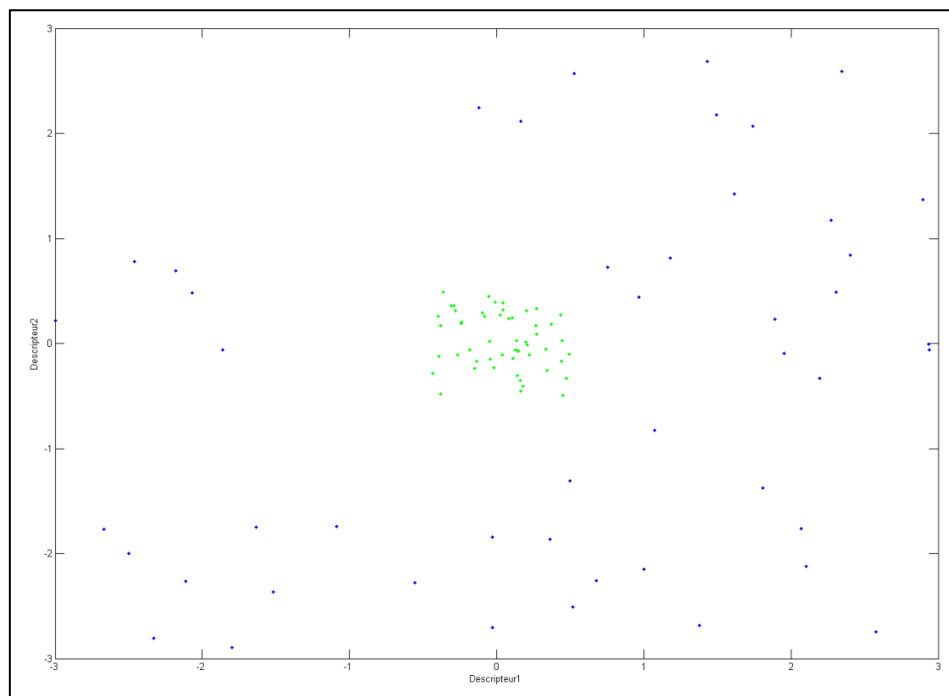


Figure 60: Exemple de données non linéaires. Problème de discrimination binaire avec en vert les individus appartenant à la classe 1 et en bleu les individus appartenant à la classe 2.

La détermination d'une telle fonction non linéaire est très difficile voire impossible. En projetant les données dans un espace où la fonction séparerait linéairement les exemples, on peut utiliser le formalisme des SVM vu précédemment, basé sur la détermination d'une fonction linéaire. Ainsi en introduisant une application :

$$\Phi: \mathbb{R}^D \rightarrow E$$

L'algorithme SVM, que nous avons décrit précédemment, appliqué aux données $\Phi(\vec{x}_i)$ dans l'espace E produira des surfaces de décision non-planes dans l'espace de départ $\mathbb{R}^D$. Cette surface dépendra donc du choix de l'application Φ .

Cette procédure est rendue très efficace grâce à l'astuce du noyau⁷. Cette astuce nous permet d'effectuer les calculs nécessaires dans l'espace de départ $\mathbb{R}^D$ sans passer explicitement dans l'espace des caractéristiques E . Ceci est dû au fait que dans les calculs des SVM les données apparaissent sous forme de produit scalaire $\langle x_i, x_j \rangle$, il suffit de trouver une façon efficace de calculer $\langle \Phi(\vec{x}_i), \Phi(\vec{x}_j) \rangle$. Nous définissons pour cela une fonction appelée noyau définie ainsi :

$$k(\vec{x}_i, \vec{x}_j) = \langle \Phi(\vec{x}_i), \Phi(\vec{x}_j) \rangle \quad \text{IV-20}$$

Ainsi toute la présentation des SVM faite durant les parties précédentes reste valable en remplaçant simplement $\langle x_i, x_j \rangle$ par $k(\vec{x}_i, \vec{x}_j)$. La nouvelle fonction de décision est donc définie par le signe de :

$$f(\vec{x}) = \sum_{i=1}^{n_s} \alpha_i y_i k(\vec{x}_i, \vec{x}) + b \quad \text{IV-21}$$

Ainsi l'avantage d'une telle approche est qu'il n'est pas nécessaire de connaître Φ explicitement. Il suffit d'utiliser des noyaux qui respectent certaines conditions.

La fonction $k(\vec{x}_i, \vec{x}_j)$ peut être vue comme une matrice symétrique G dite de Gram [61] qui représente les distances entre tous les exemples :

$$G = \begin{bmatrix} k(\vec{x}_1, \vec{x}_1) & \cdots & k(\vec{x}_1, \vec{x}_n) \\ \vdots & \ddots & \vdots \\ k(\vec{x}_n, \vec{x}_1) & \cdots & k(\vec{x}_n, \vec{x}_n) \end{bmatrix} \quad \text{IV-22}$$

Pour qu'une fonction k soit un noyau, il faut qu'il respecte les conditions de Mercer [62] c'est-à-dire que la matrice G doit être semi-définie positive⁸. La construction de tels noyaux peut-être réalisée par nos soins, mais il existe dans la littérature scientifique des noyaux qui sont largement étudiés.

Ci-dessous une liste non-exhaustive des noyaux les plus utilisés :

- Noyau linéaire : si les données sont linéairement séparables on n'utilise pas de noyaux car on n'a pas de besoin de changer d'espace, et le produit scalaire suffit donc pour définir la fonction de décision :

⁷

On peut trouver dans la littérature le nom anglais de kernel trick.

⁸ Une matrice $M \in \mathcal{M}(N, N)$ est symétrique semi-définie positive si l'ensemble de ses valeurs propres sont positives ou nulles, donc si son spectre $S_p(M) \in \mathbb{R}^+$.

$$k(\vec{x}_i, \vec{x}_j) = \langle \vec{x}_i, \vec{x}_j \rangle \quad \text{IV-23}$$

- Noyau polynomial homogène : le noyau polynomial de degrés p correspond à une transformation ϕ par laquelle les composantes des vecteurs transformés $\Phi(\vec{x})$ sont tous les monômes d'ordre δ formés à partir des composantes de $\vec{x}$. Ce noyau est défini ainsi :

$$k(\vec{x}_i, \vec{x}_j) = \langle \vec{x}_i, \vec{x}_j \rangle^p \quad \text{IV-24}$$

Nous pouvons calculer la dimension de l'espace des caractéristiques en fonction de la dimension D de l'espace de départ et du degré p du noyau polynomial :

$$\dim(E) = \binom{p + d - 1}{\delta} \quad \text{IV-25}$$

- Noyau polynomial inhomogène : l'idée est la même que pour le noyau polynomial homogène sauf que nous ajoutons une constante afin de prendre en compte tous les monômes de degrés inférieurs à δ , ainsi la dimension de l'espace des caractéristiques sera plus élevée que dans le cas homogène. Ce noyau est défini ainsi :

$$k(\vec{x}_i, \vec{x}_j) = (1 + \langle \vec{x}_i, \vec{x}_j \rangle)^p \quad \text{IV-26}$$

- Noyau RBF (Radial Basis Functions) : ces fonctions sont radiales, elles ne dépendent que de la distance entre leurs arguments,

$$\Phi(\vec{x}, \vec{y}) = \Phi(\|\vec{x} - \vec{y}\|)$$

Le noyau Gaussien RBF applique ainsi une gaussienne sur la distance entre les exemples. On montre dans ce cas que l'espace des caractéristiques est de dimension infinie. Ce noyau est défini comme suit :

$$k(\vec{x}_i, \vec{x}_j) = \exp\left(-\frac{\|\vec{x}_i - \vec{x}_j\|^2}{D\sigma^2}\right) \quad \text{IV-27}$$

4. Choix des paramètres

a. Influence du paramètre C

Le paramètre C est un paramètre particulier en ce sens que lui seul intervient directement dans la fonction à minimiser, lors de la résolution du problème SVM à marge souple. Il est appelé paramètre de pénalisation. Si on analyse de plus près le comportement, on peut remarquer que :

- Lorsque $C \rightarrow \infty$, la tolérance aux erreurs de classification devient de plus en plus rigide et on retombe sur le problème des SVM à marge dure, on risque le sur-apprentissage.

- Lorsque $C \rightarrow 0$, le système tolère les erreurs jusqu'à ne plus pouvoir distinguer les deux classes. On remarque d'ailleurs que si on prend $C = 0$, on aura $\alpha_i = 0 \forall i = 1 \dots N$, il n'y aura donc plus de vecteurs supports et la fonction de décision ne dépendra plus des données d'apprentissage.

Ainsi la valeur de C optimale est un compromis entre la maximisation de la marge et la tolérance aux erreurs de classification⁹.

b. Choix du noyau

Dans la mise en œuvre des SVM, le choix du noyau et de ses paramètres, reste un problème ouvert. Il faut adapter le noyau aux données, comme nous l'avons vu plus haut dans l'exemple du XOR. Nous avons choisi de nous concentrer sur le noyau RBF gaussien, car parmi les noyaux connus il est celui qui permet d'obtenir les résultats les plus corrects sur nos données. Dans la littérature, le noyau RBF a été le plus souvent utilisé, pour différentes raisons:

- Le nombre d'hyperparamètres induit par ce noyau est faible comparé par exemple aux noyaux polynomiaux.
- Les difficultés numériques sont réduites. En effet $0 < k(\vec{x}_i, \vec{x}_j) \leq 1$ pour un noyau RBF gaussien alors que pour un noyau polynomial les valeurs peuvent être infinies ou nulles.

En revanche, il est à noter que lorsque le nombre de descripteurs est très grand, en général supérieur à 500, il peut être préférable d'utiliser simplement un noyau linéaire, car la dimension de l'espace de départ est assez grande et ainsi la probabilité de trouver un hyperplan linéaire est plus grande.

c. Influence du paramètre σ

Lorsque on utilise un noyau de type RBF gaussien, celui-ci traduit une mesure de similarité basée sur la distance entre les exemples de la base d'apprentissage. Si nous analysons le comportement de cette valeur on se rend compte que lorsque :

- $\sigma \rightarrow \infty$: la mesure de similitude tend vers 1 ainsi ceci fait croître la similarité entre les exemples. Ainsi lorsque σ devient trop grand on peut voir que la mesure de similarité ne permet plus de distinguer des exemples de classes différentes. L'algorithme crée donc des frontières incohérentes avec les données comme nous pouvons le voir sur la Figure 64.
- $\sigma \rightarrow 0$: la mesure de similitude tend vers 0, ainsi la mesure de similarité devient pratiquement nulle entre chaque exemple. Lorsque σ est trop faible, la mesure de similarité devient trop sélective et la fonction de décision doit être construite à partir de beaucoup de vecteurs supports pour couvrir tout l'espace. Ainsi on se retrouve dans une situation où'on risque d'effectuer un sur-apprentissage, au détriment donc de la capacité de généralisation de notre classifieur, ce qui est le cas sur la Figure 62.

Nous avons vérifié ce comportement sur un jeu de données modélisant le problème de l'échiquier qui est un exemple de données artificielles non-linéairement séparables. Sur la Figure 61 nous montrons la position du problème de l'échiquier avec la frontière idéale.

⁹ Plus connu sous le nom d'outliers.

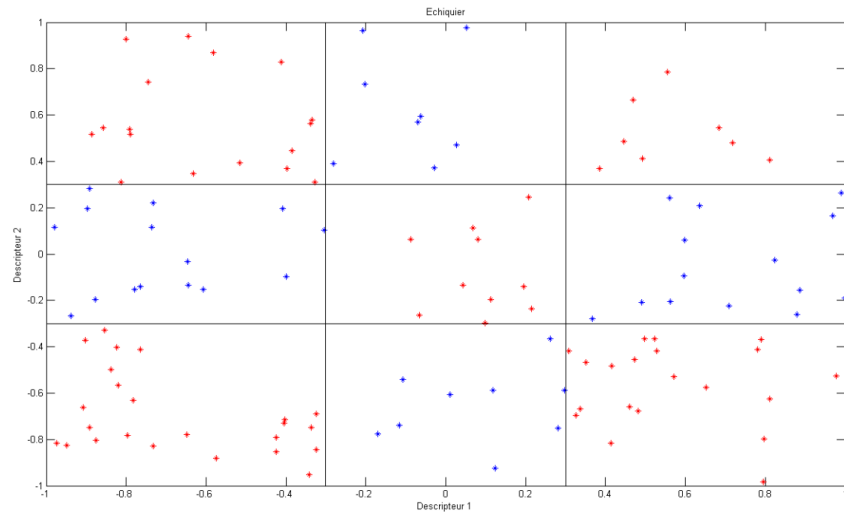


Figure 61: Problème de l'échiquier avec représentation de la fonction de décision idéale. Les individus appartenant à la classe 1 sont en bleu et les individus appartenant à la classe 2 sont en rouge.

Nous avons réalisé un apprentissage sur ces données en utilisant un noyau RBF gaussien pour différentes valeurs de σ . Nous avons fixé la valeur de C à 100 afin de voir l'influence du paramètre σ .

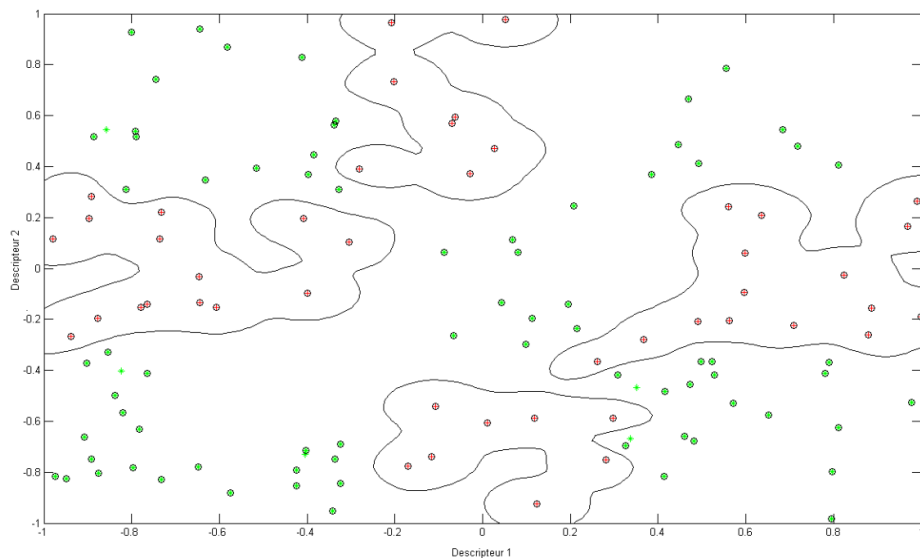


Figure 62: Surfaces de décision obtenues par apprentissage SVM à noyau gaussien RBF avec $\sigma=0.1$

Nous observons sur ces figures que le comportement prédit est vérifié, en effet sur la Figure 62, σ est trop petit ainsi de nombreux exemples deviennent vecteurs supports et nous sommes donc dans une situation de sur-apprentissage. Sur la Figure 64, σ est trop grand ainsi la mesure de similarité ne permet plus de distinguer les exemples de classes opposées, ainsi la frontière tracée n'a aucun sens. Un bon compromis est trouvé avec $\sigma=0.3$ sur la Figure 63, où la frontière de décision est cohérente et se rapproche de la frontière idéale.

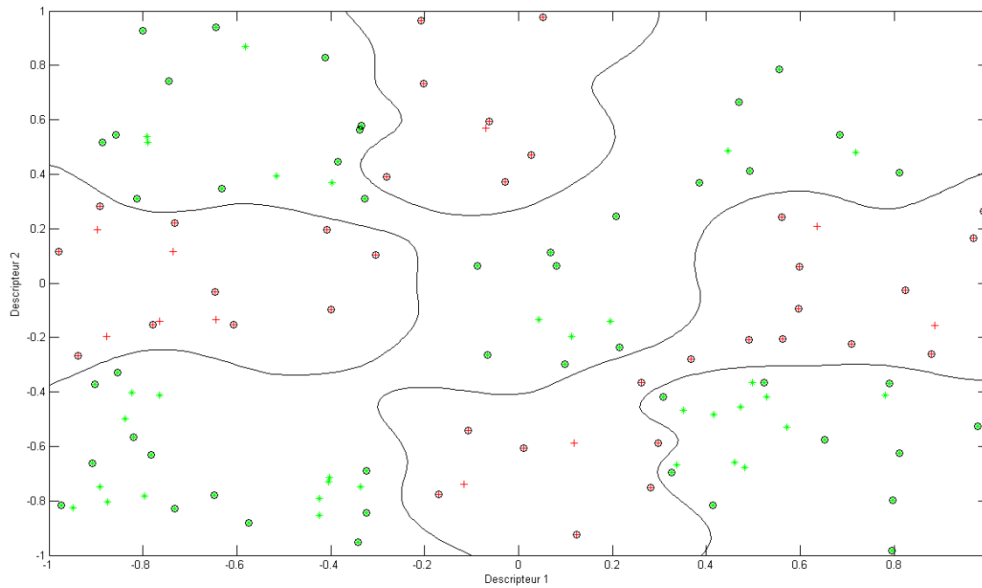


Figure 63: Surfaces de décision obtenues par apprentissage SVM à noyau gaussien RBF avec $\sigma=0.3$

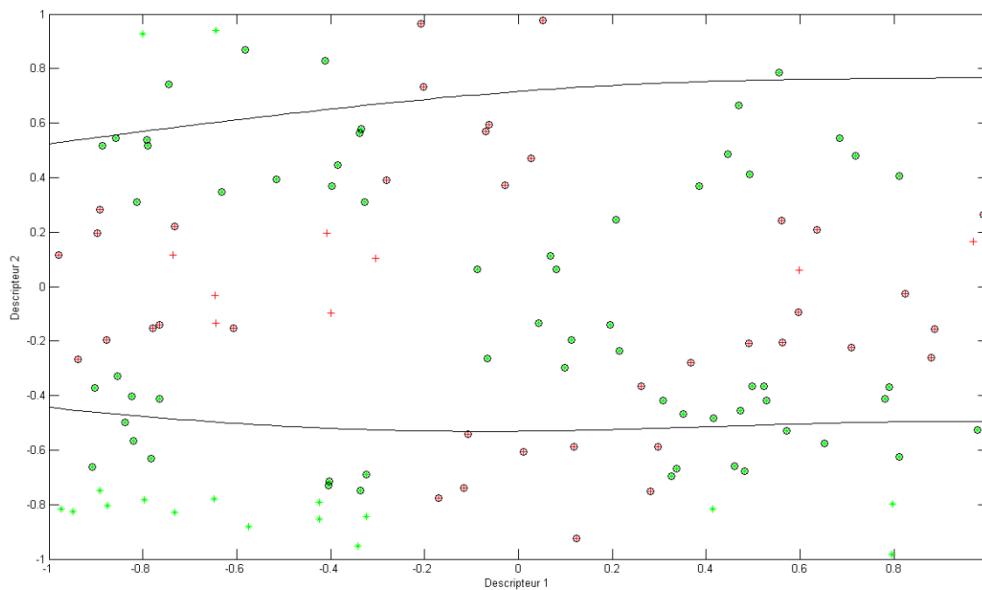


Figure 64: Surfaces de décision obtenues par apprentissage SVM à noyau gaussien RBF avec $\sigma=3$

d. Stratégie de recherche des paramètres optimaux

Il est à noter que les différents paramètres que nous avons vus ne sont pas indépendants, c'est-à-dire que nous ne pouvons pas les optimiser séparément, nous devons le faire de manière conjointe. Pour cela, il existe plusieurs moyens de recherche pour trouver les différents paramètres optimaux. La stratégie la plus courante est la recherche par maillage. Elle consiste à évaluer les performances des SVM pour différentes valeurs discrètes des paramètres, pour un critère donné.

La recherche par maillage¹⁰ est la méthode la plus généralement employée. Elle consiste à évaluer les performances du classifieur SVM appris sur un ensemble fini de V valeurs appartenant à l'ensemble $\Lambda = \{\theta_i, i \in [1, \dots, V]\}$. Soit $P(k_\theta)$ la mesure de performance du noyau k_θ , l'algorithme consiste donc à retenir la valeur $\hat{\theta}$ telle que :

$$\hat{\theta} = \underset{\theta \in \Lambda}{\operatorname{argmax}} P(k_\theta) \quad \text{IV-28}$$

Pour les choix des valeurs de Λ , on utilise généralement pour chaque paramètre un ensemble de valeurs également réparties dans un intervalle donné. Λ est alors le produit cartésien de ces ensembles et constitue un maillage de l'espace des paramètres sur un intervalle donné. Il est courant d'utiliser des valeurs réparties de manière logarithmique. Cependant, la recherche par maillage souffre de deux défauts majeurs :

- On fait face à une grande combinatoire de paramètres à tester dès que le nombre de paramètres à régler dépasse un ou deux. Ce qui peut être couteux en termes de temps de calcul.
- Si nous supposons que notre critère de performance est fiable, cette stratégie ne garantit pas de trouver le maximum global car il se peut que ce dernier ne se trouve pas sur la grille du maillage que nous avons défini pour une échelle donnée. Prendre un maillage fin garantit presque sûrement de tomber sur le maximum global, mais au prix d'une grande augmentation du temps de calcul, c'est pour cela qu'il faut réaliser un compromis entre nombre de valeurs à tester dans notre grille et temps de calcul.

5. SVM multi-classes

Nous avons vu précédemment que par construction les SVM étaient des classifieurs binaires, permettant de séparer deux classes. Or dans notre problématique de classification nous disposons de plusieurs classes. Nous devons donc adopter une stratégie afin d'adapter ce classificateur à la discrimination multi-classe. Plusieurs stratégies ont été élaborées sur la base d'une décomposition d'un problème en une collection de sous-problèmes binaires, dont il convient de combiner ensuite les résultats pour déterminer la solution multi-classe finale [57], [63], [64]. Les deux méthodes de ce type les plus utilisées dans la littérature sont la stratégie « un contre tous » et la stratégie « un contre un ». De plus, un autre type de méthode est la création de graphes de décision, ce dernier type est basé sur l'utilisation astucieuse de la méthode un contre un. Enfin, des formulations des SVMs pour un problème multi-classes ont été proposé, nous en présentons succinctement quelques-unes.

a. Un contre tous

Cette approche est la plus simple des méthodes de décomposition [65], [66]. Elle consiste à utiliser un classifieur binaire pour chaque classe. Nous choisissons une classe k et nous créons une seconde classe qui est la réunion de toutes les autres. Ainsi, le $k^{\text{ième}}$ classifieur binaire aura pour fonction de distinguer les éléments de la classe k de tous les autres éléments des autres classes. Ainsi pour

¹⁰ Grid-search

affecter un exemple, on le présente à K classifieurs binaires, et on fusionne les différentes sorties de chaque classifieur. La classe affectée à l'exemple sera celle pour laquelle le classifieur a renvoyé la distance à la marge la plus élevée. Il convient de signaler que cette méthode implique d'effectuer des apprentissages aux répartitions entre catégories très déséquilibrées, ce qui soulève des difficultés pratiques. Sur la Figure 65, nous pouvons voir une illustration de la méthode un contre tous, sur un problème à trois classes. La zone en grise est une zone d'incertitude, c'est-à-dire que le classifieur ne peut pas prendre de décision.

b. Un contre un

L'approche « un contre un » [67] vise à élaborer un classifieur pour chaque paire de classes possibles. Le classifieur est indicé par le couple (k, l) avec $1 \leq k \leq K$ est destiné à distinguer la classe k de la classe l . Ce qui nous donne C_K^2 classifieurs différents, la décision pour un nouvel exemple s'obtient traditionnellement en utilisant la technique du vote majoritaire [68]. Cependant la technique du vote majoritaire pose un problème majeur, des indéterminations peuvent intervenir, c'est-à-dire un cas où plusieurs classifieurs ont le même nombre de vote. Sur la Figure 66, nous pouvons voir une illustration de la méthode un contre un, sur le même problème que précédemment. La zone en grise est une zone d'incertitude, nous voyons qu'elle est beaucoup moins importante que dans le cas un contre tous.

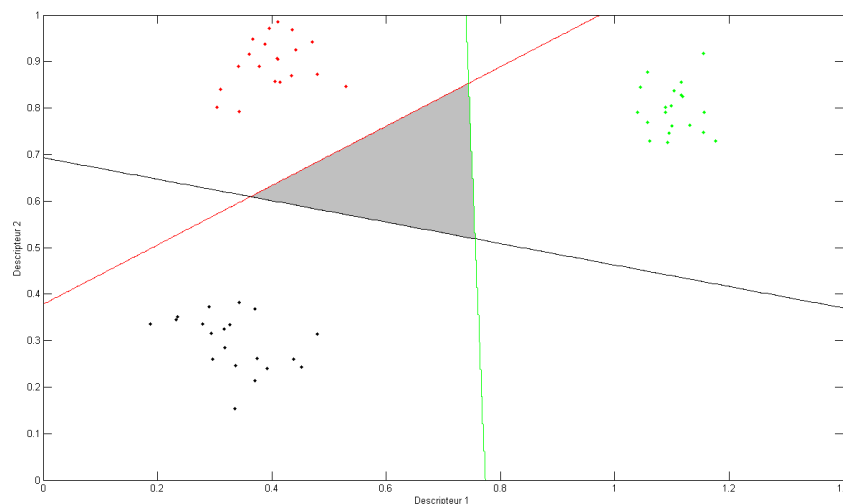


Figure 65: Illustration du principe un contre tous, en gris se trouve la zone d'indétermination

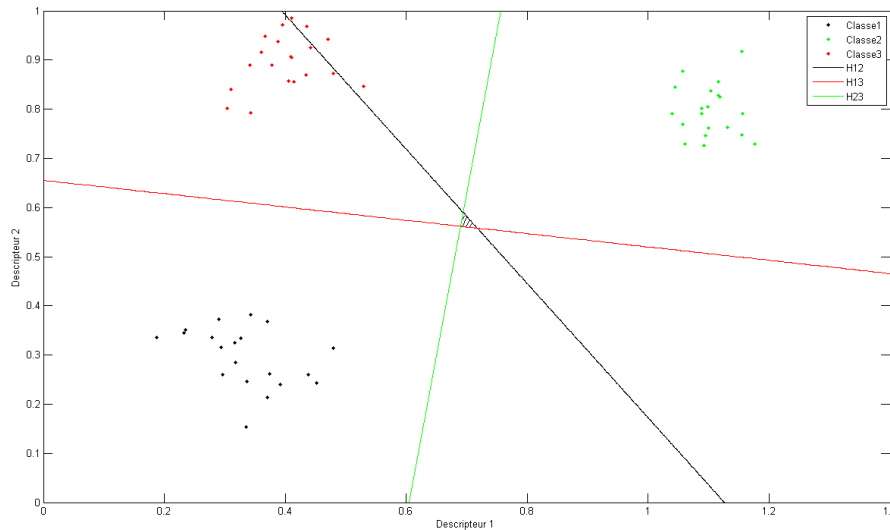


Figure 66: Illustration du principe un contre un, la zone d'indétermination est hachurée au centre

c. Graphe de décisions

La première méthode fondée sur un graphe de décision est la DAGSVM [69], cette méthode s'appuie sur un graphe de décision orienté¹¹. On obtient, comme dans la méthode 1 contre 1, $\frac{K(K-1)}{2}$ hyperplans. Puis au lieu d'utiliser la technique du vote majoritaire on construit un graphe de décision de la manière suivante :

On définit une mesure E_{kl} qui représente la capacité de généralisation des différents hyperplans pour chaque classifieur SVM binaire :

$$E_{kl} = \frac{N_{vs}}{N_{exemples}} \quad \text{IV-29}$$

Nous construisons le graphe de décision selon les étapes suivantes :

- Créer une liste L contenant toutes les classes,
- Si L contient une seule classe, créer un nœud étiqueté de cette classe et l'algorithme s'arrête.
- Calculer pour chaque paire de classes (k, l) la capacité de généralisation E_{kl} .
- Rechercher les deux classes k et l dont E_{kl} est maximum, on crée alors un nœud N , avec l'étiquette (k, l) .
- Créer un graphe de décision à partir de la liste $L - \{k\}$ ¹², de la même manière, et on l'attache au fils gauche de N .
- On effectue la même opération à partir de la liste $L - \{l\}$, de la même manière, et on l'attache au fils droit de N .

¹¹

¹² C'est-à-dire un graphe où chaque nœud est une décision binaire
Cela signifie que l'on a retiré la classe k à la liste L

Sur la Figure 67, une illustration de la méthode est proposé.

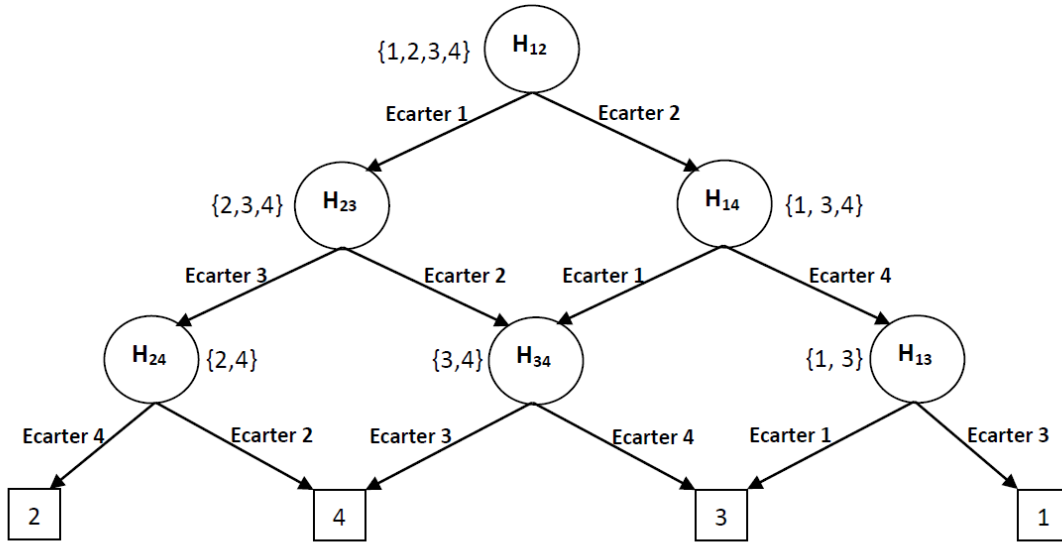


Figure 67: Exemple de graphe de décision, pour un problème à 4 classes

A l'issue de l'algorithme, on obtient un graphe de décision, un nouvel exemple à classier est confronté tout d'abord à l'hyperplan racine. Si la décision est positive on continue avec le fils droit¹³, sinon on continue avec le fils gauche et ainsi de suite jusqu'à atteindre une feuille¹⁴. Cette dernière représentera la classe finale de l'exemple. Cette méthode est confrontée uniquement à $K - 1$ classifieurs ce qui la rend très rapide en classification par rapport aux méthodes 1 contre 1 et 1 contre tous.

Dans [70] les auteurs proposent une autre méthode basée cette fois-ci sur les dendrogrammes. L'approche est la suivante, on calcule le centre de gravité des éléments de la base d'apprentissage appartenant à la même classe, ainsi on calcule les distances entre chaque classe et on crée un dendrogramme. A chaque nœud du dendrogramme il y a un classifieur SVM de type un contre un. En utilisant cette approche c'est comme si nous réalisons une fusion de classes et qu'à chaque nœud du dendrogramme nous séparions les classes de nouveau. Cette approche est intéressante dans le sens où nous prenons en compte toutes les données.

d. M-SVM

Soit $((\vec{x}_i, y_i))_{1 \leq i \leq N} \in (\mathbb{R}^D \times \{1, \dots, C\})^N$. Une M-SVM à C classes est un modèle discriminant à grande marge obtenu en minimisant l'hyperplan $\sum_{k=1}^C h_k = 0$ de $\mathcal{H}$ une fonction objectif J_{M-SVM} de la forme:

¹³

¹⁴ A chaque sommet de l'arbre on a une décision binaire. Soit on choisit le fils droit, soit le fils gauche.

Une feuille représente la fin de l'arbre.

$$J_{M-SVM}(h) = \sum_{i=1}^m l_{M-SVM}(y_i, h(\vec{x}_i)) + \lambda \|\bar{h}\|_{\bar{\mathcal{H}}}^2$$

Les deux éléments distinguant les différents M-SVM sont donc la fonction l_{M-SVM} et le choix de la norme sur $\bar{\mathcal{H}}$. Plusieurs types de M-SVM ont été développés dans la littérature, nous nous contenterons de citer les 3 principaux modèles:

- Modèle de Weston et Watkins [71] [72]
- Modèle de Crammer et Singer [73] [63]
- Modèle de Lee [74]

Cette partie n'est pas développée car nous avons préféré utiliser l'approche de décomposition du problème multiclassé en plusieurs problèmes bi-classes. Nous expliquons pourquoi au sein de la conclusion.

e. Conclusion

La présentation des différentes techniques montre qu'il n'existe pas de nos jours de formulation multiclassé faisant office de référence pour les SVM. On préfère utiliser les schémas de décomposition que nous avons présentés précédemment. Ceci est dû au fait que ces méthodes sont plus rapides en terme de temps de calcul, et que les résultats expérimentaux [71] [75] montrent une stagnation des performances. Dans [75] les auteurs prennent parti pour l'approche 1 contre tous car c'est la plus simple des méthodes multiclassées, et les différentes approches se valent si l'on affine correctement les paramètres des SVM. Friedman dans [68] a écrit: « La leçon la plus importante à tirer de l'exercice ci-dessus est que les performances relatives des différentes approches peuvent fortement dépendre du problème particulier auquel elles sont appliquées. Comme tous les autres aspects de la méthodologie de l'apprentissage, aucune approche ne domine toutes les autres dans toutes les situations ». C'est pour ces raisons que nous nous sommes concentrés sur le réglage des paramètres, qui est une étape essentielle comme nous l'avons vu précédemment. Dans la partie suivante nous allons nous concentrer sur la description des données.

Chapitre 4 : Caractérisation des signaux acoustiques sous-marins

I. Descripteurs

La question de représentation des données d'entrée du classifieur SVM est abordée dans ce chapitre. Soit un problème de discrimination ayant un ensemble d'apprentissage $S = \{(\vec{x}_i, y_i)\}_{i=1, \dots, N}$, où les exemples $\vec{x}_i \in R^D$ sont décrits par D composantes correspondant chacune à un descripteur, on a :

$$\vec{x}_i = [x_{i,1}, \dots, x_{i,D}]^T$$

et sont associés à un label correspondant à une classe $y_i \in \{1 \dots K\}$. L'ensemble S représente donc un tableau où chaque ligne est un vecteur de descripteurs pour un exemple donné et chaque colonne, appelée f_i , $\forall i = 1 \dots D$ est la valeur d'un descripteur pour tous les exemples de la base d'apprentissage, la dernière colonne du tableau correspond au label de classification y_i .

A. Représentation de l'information par des vecteurs de descripteurs

Nous avons présenté dans la partie précédente l'architecture du classificateur SVM, ainsi que les points relatifs à leur mise en œuvre sur un problème impliquant plusieurs classes. Dans cette partie, nous n'avons pas abordé le sujet concernant la nature des données d'apprentissage, qui conditionne le choix de notre espace de départ. Cependant, ce choix est essentiel car nous pouvons essayer de construire le meilleur classifieur possible, si nos données ne sont pas séparables dans l'espace de départ nous obtiendrons de mauvais résultats de classification. Autrement dit plus les régions associées aux différentes classes se chevauchent dans l'espace d'entrée, plus le problème sera difficile à traiter et la probabilité d'erreur plus grande.

C'est pour cela, que dans cette partie, nous allons présenter une collection de descripteurs audio, choisis pour leurs capacités à séparer au mieux les différentes classes étudiées.

B. Normalisation

Les descripteurs calculés sont de différentes natures, ainsi leur dynamique peut être très différente. Pourtant, lorsque nous utilisons les différents noyaux lors du calcul de la frontière de décision, nous remarquons que les descripteurs sont mis en concurrence au travers de sommes mettant en jeu une pondération uniforme. Ainsi, si un descripteur a une moyenne très largement supérieure à un autre descripteur, l'influence du descripteur le plus faible sera pratiquement nulle dans l'expression du noyau. C'est pour cette raison qu'une étape de normalisation rendant les données sans dimension physique est essentielle.

Dans la littérature scientifique il existe plusieurs méthodes de normalisation, en voici quelques-unes :

- Homogénéiser les statistiques du premier et du second ordre pour chaque descripteur [76].
On note $x_{n,d}$ la composante d'indice d du vecteur exemple $\vec{x}_n$, on estime alors la moyenne

que l'on notera μ_d et l'écart-type σ_d du descripteur d grâce aux estimateurs classiques qui sont :

$$\mu_d = \frac{1}{n} \sum_{i=1}^n x_{i,d} \quad \text{I-1}$$

$$\sigma_d^2 = \frac{1}{n-1} \sum_{i=1}^n (x_{i,d} - \mu_d)^2 \quad \text{I-2}$$

Les descripteurs normalisés ont donc pour expression :

$$\hat{x}_{i,d} = \frac{x_{i,d} - \mu_d}{\sigma_d} \quad \text{I-3}$$

Cette méthode est couramment utilisée dans la littérature scientifique cependant, elle fait l'hypothèse de la gaussianité des données, alors que cette hypothèse n'est pas toujours vérifiée.

- Normaliser les données, de telle sorte à les réduire dans l'intervalle $[0,1]$. Ainsi la dynamique de chaque descripteur sera la même. Si $f_i \in [a, b]$ alors :

$$\hat{f}_i = \frac{f_i - a}{b - a} \in [0,1] \quad \text{I-4}$$

- Une autre méthode [77] consiste à remplacer la valeur du descripteur $x_{i,d}$ par la valeur de la fonction de répartition $(x_{i,d})$, estimée sur l'ensemble des exemples. Ainsi, on garantit que toutes les composantes auront une distribution quasi-uniforme dans l'intervalle $[0,1]$.
- Enfin, une méthode basée sur l'interquartile range (IQR) permet d'éviter d'inclure les valeurs aberrantes dans la normalisation de nos données à un intervalle restreint $([-1;1]$ ou $[0;1])$. Au lieu de normaliser les données en utilisant le minimum et le maximum, on utilise l'IQR $[10; 90]$ c'est à dire que le minimum est remplacé par la valeur qui laisse 10% des valeurs en dessous d'elle et le maximum est remplacé par la valeur qui laisse 90 % des données au-dessous d'elle.

Une étude comparative a été réalisée dans [77] concernant les différentes méthodes de normalisation de données.

C. Détail des descripteurs utilisés

Dans cette partie, nous exposons les descripteurs utilisés dans le système de reconnaissance automatique. C'est grâce à ces derniers que chaque forme sera représentée.

1. Descripteurs temporels

Les descripteurs suivants sont basés sur la forme d'onde du signal audio :

- Le taux de passage par zéro (ZCR) [78],
- Les moments statistiques temporels d'ordre 1 à 4:
 - le centroïde
 - la largeur spectrale
 - l'asymétrie (skewness)
 - la platitude (kurtosis)

2. Descripteurs spectraux

Les descripteurs spectraux sont calculés à partir du spectre obtenu par la Transformée de Fourier Discrète (TFD), qui est définie, sur une trame de N échantillons, de la façon suivante:

$$X[k] = \sum_{n=0}^{N-1} x[n] e^{-2j\pi k \frac{n}{N}} \quad \forall k \in [0, \dots, N-1] \quad \text{I-5}$$

Le calcul de la TFD est précédé de la pondération du signal de trame par une fenêtre de Hanning, qui limite l'étalement des pics spectraux et nous permet d'éviter le phénomène de Gibbs. En pratique, nous utilisons $|X(k)|$ dans les descripteurs ci-dessous:

- Les moments statistiques spectraux: on considère dans ce type de descripteur notre spectre d'amplitude comme une densité de probabilité sur lequel nous allons calculer les moments d'ordre 1 à 4:
 - le centroïde spectral, décrivant le centre de gravité du spectre,
 - la largeur spectrale, décrivant l'étendue du spectre autour de sa moyenne,
 - l'asymétrie spectrale (Skewness), représentant la symétrie du spectre autour de sa moyenne,
 - la platitude spectrale (Kurtosis), elle est d'autant plus grande que le spectre est "peaky" autour de sa moyenne, pour un spectre de forme gaussienne sa valeur est nulle.
- Descripteurs MPEG-7: nous exploitons deux descripteurs de la norme standard MPEG-7 [79]:
 - le rapport spectral,
 - la platitude spectrale.
- Le flux spectral [80], représentant une variation spectrale entre trames consécutives.

3. Descripteurs cepstraux

Le cepstre du signal $x[n]$ s'obtient par la transformée de Fourier inverse du logarithme du spectre d'amplitude $|X[k]|$:

$$c[q] = \sum_k \log |X[k]| e^{2j\pi q \frac{n}{N}} \quad \forall q \in [0, \dots, N-1] \quad \text{I-6}$$

Dans une modélisation source-filtre du signal:

$$s(n) = g * h(n) \quad \text{I-7}$$

où $g(n)$ est l'excitation et $h(n)$ le filtre. Il est montré dans [81] que les coefficients cepstraux correspondant aux basses quéfrenes¹⁵ q représentent la contribution du filtre $h(n)$. Il s'agit aussi d'une version lissée de l'enveloppe spectrale et c'est pour cette raison que nous l'utiliserons.

Nous allons utiliser comme descripteurs une variante des coefficients cepstraux qui se nomment les Mel-Frequency Cepstral Coefficients (MFCC). Ils s'obtiennent en considérant, pour le calcul du cepstre, une représentation fréquentielle selon une échelle perceptive appelée l'échelle des fréquences Mel, que nous avons définie dans la partie 2. Pour ce faire, nous utilisons un banc de filtres triangulaires Mel. Nous intégrons le spectre d'amplitude $|X(k)|$ par bandes de Mel, pour obtenir un spectre d'amplitude modifié $\tilde{a}_m$, $m = 1 \dots M_l$, où $\tilde{a}_m$ représente l'amplitude dans la bande m . Les MFCC s'obtiennent alors par une transformée en cosinus discrète inverse (de type 2) du logarithme de $\tilde{a}_m$:

$$\tilde{c}(q) = \sum_m^{M_l} \log(\tilde{a}_m) \cos\left(q\left(m - \frac{1}{2}\right)\frac{\pi}{M_l}\right) \quad \text{I-8}$$

Nous utilisons un banc de filtres composés de 16 bandes de MEL ($M_l = 16$), ainsi nous obtenons 16 coefficients MFCC.

4. Descripteurs perceptuels

a. Loudness spécifique relative (L_d)

La loudness spécifique [82] est définie dans la bande critique b_c par :

$$L(b_c) = E(b_c)^{0.23} \quad \text{I-9}$$

où $E(b_c)$ est l'énergie du signal dans la bande b_c . Nous mesurons en fait la loudness spécifique relative:

$$L_d(b_c) = \frac{L(b_c)}{L_T} \quad \text{I-10}$$

avec $L_T = \sum_{sb} L(sb)$ étant la loudness totale. En faisant cela nous rendons la loudness indépendante des conditions d'enregistrement du signal. En effet, il se peut que pour un signal de même nature l'énergie totale soit différente en fonction de la distance et du milieu de propagation, nous voulons que les descripteurs soient insensibles à ces conditions.

¹⁵

Pour rappeler le fait que l'on effectue une transformation inverse à partir du domaine fréquentiel, les dénominations des notions sont des anagrammes de celles utilisées en fréquentiel. Ainsi le spectre devient le cepstre, la fréquence une quéfrence, un filtrage un lifrage.

b. Sharpness

La sharpness [82] représente une version "perceptuelle" du centroïde spectral calculée à partir de la loudness spécifique selon:

$$S_h = 0.11 \frac{\sum_{b_c} g(b_c) L_d(b_c)}{L_T} \quad \text{I-11}$$

avec $g(b_c)$ définie par

$$g(b_c) = \begin{cases} 1 & \text{si } b_c < 15 \\ 0.066 e^{0.171 b_c} & \text{si } b_c \geq 15 \end{cases} \quad \text{I-12}$$

c. Largeur perceptuelle

Il s'agit d'une mesure de l'écart entre la loudness spécifique maximale et la loudness totale, elle est définie dans [82] par:

$$S_p = \left(\frac{L_T - \max_{b_c} L_d(b_c)}{L_T} \right)^2 \quad \text{I-13}$$

D. Discussions

Le choix des descripteurs est essentiel dans la mise en place d'un système de classification. En effet, ces derniers vont permettre la séparation entre les différentes classes. En entrée du système de classification les signaux ne seront représentés que par leurs descripteurs, d'où leur importance capitale. Cependant un jeu de descripteurs peut convenir pour bien séparer deux classes mais peut ne pas convenir pour deux autres classes différentes. De plus, nous pouvons penser de prime abord qu'augmenter le nombre de descripteurs est forcément bénéfique pour la classification, mais cette idée est fautive. En effet, plusieurs descripteurs peuvent avoir un effet contre-productif est donc introduire un bruit de classification. Pour illustrer cette idée, prenons l'exemple d'un classifieur qui a pour but de séparer les individus en hommes d'un côté et femmes de l'autre. Alors, si la description de chaque individu se fait par le génome alors l'information principale sera noyée dans une grande quantité d'information qui ne sera pas pertinente, alors que seulement un descripteur suffit à classer correctement cette population. C'est pour cette raison qu'il est nécessaire d'effectuer une étape de sélection des descripteurs, afin de ne conserver pour chaque problème binaire de classification, que les descripteurs nous permettant de discriminer au mieux les deux classes mises en jeu au sein de ce problème. Nous allons présenter dans la prochaine section différentes techniques permettant de réaliser cette étape nécessaire.

II. Sélection des descripteurs

A. Nécessité d'une étape de sélection des descripteurs

Nous devons effectuer une sélection des descripteurs car il est fort probable que les descripteurs que nous avons à disposition ne soient pas la combinaison de descripteurs qui engendreraient le meilleur taux de classification. En effet, nous pouvons penser que plus il y a de descripteurs meilleure sera la classification, or cette idée reçue est fausse, ceci est dû principalement à ce que l'on nomme la malédiction de la dimension [83]. Souvent nous introduisons du bruit dans l'information considérée par le classifieur, c'est-à-dire des descripteurs qui seront contre-productifs pour la tâche d'identification. En effet si nous introduisons un descripteur qui est redondant avec un autre ce dernier nous pénalisera. Faire une sélection humaine serait une tâche difficile car il a été montré dans [84] que deux descripteurs non-pertinents quand on les prend individuellement peuvent être très pertinents lorsqu'ils sont exploités ensemble. On peut généraliser ce constat, comme cela a été démontré dans [85], où on voit que les k pires descripteurs, ceux qui ont obtenu le moins bon score selon un critère de performance donné, peuvent se révéler meilleurs ensemble que les k meilleurs descripteurs, on parle de phénomène d'interpertinence. Ainsi l'humain ne peut pas à son échelle prendre en compte toutes les dépendances possibles entre les descripteurs.

De plus, une sélection des descripteurs nous permet une réduction de la complexité calculatoire et de gestion de la mémoire, en enlevant les descripteurs non pertinents et/ou redondants. Enfin elle peut apporter pour l'opérateur une meilleure compréhension d'un problème par l'interprétation des descripteurs les plus pertinents.

Il est intéressant d'approfondir la notion de pertinence qui est abstraite. Dans la littérature nous pouvons trouver différentes définitions, la plus connue est celle que nous trouvons dans [86]. Selon cette dernière un descripteur peut être très pertinent, peu pertinent et non pertinent :

- Très pertinent : Un descripteur est très pertinent si son absence entraîne une détérioration significative de la performance du système d'identification utilisé.
- Peu pertinent : Un descripteur est dit peu pertinent s'il n'est pas très pertinent (voir ci-dessus) et si il existe un sous-ensemble tel que si on ajoute à ce sous ensemble le descripteur on remarque une augmentation significative de la performance du système.
- Non pertinent : un descripteur est non pertinent s'il n'est ni très pertinent, ni peu pertinent. Il faut alors retirer ces descripteurs de l'ensemble d'apprentissage.

Ainsi nous pouvons définir la sélection de descripteurs comme un processus de recherche permettant de trouver un sous-ensemble de descripteurs pertinents parmi l'ensemble de départ. Cette notion de pertinence dépend des objectifs et des critères du système. Dans [87], une illustration du processus de la sélection des descripteurs est proposée :

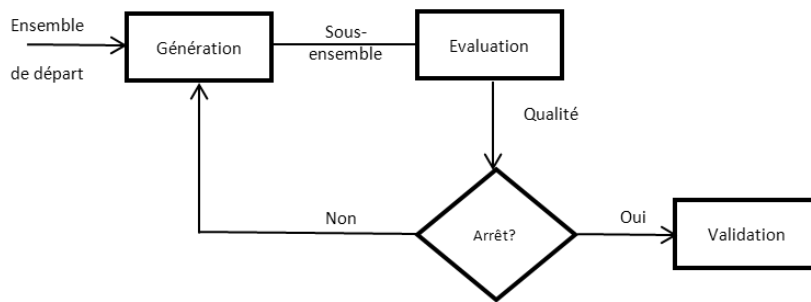


Figure 68 : Procédure générale d'un algorithme de sélection des descripteurs

B. Différentes stratégies de recherches

1. Best Individual N (BIN)

Cette stratégie de recherche consiste à évaluer chaque descripteur indépendamment des autres, selon un critère d'évaluation. Elle ne prend pas en compte les différentes interactions entre les descripteurs. L'avantage de cette stratégie est le temps de calcul, elle nous permet aussi d'obtenir un classement des différents descripteurs.

2. Sequential (SEQ)

Le but de cette stratégie est d'évaluer un sous-ensemble de descripteurs selon un critère donné. Différentes approches sont possibles :

- Approche par force brute : cette méthode est la plus intuitive, elle assure de nous retourner le meilleur sous-ensemble de descripteurs selon un critère donné. Cependant la complexité de calcul est trop grande. Plus précisément, la recherche exhaustive d'un sous-ensemble au sein d'un ensemble de D descripteurs se fait avec 2^D opérations. Lorsque D devient trop grand la recherche exhaustive devient impossible.
- Une autre approche a été développée, elle est basée sur des algorithmes de recherches intelligents, ces différents algorithmes sont itératifs et chaque itération permet de sélectionner ou de rejeter une ou plusieurs caractéristiques. Il en existe trois principaux types, la différence entre ces méthodes repose sur l'initialisation et la procédure de recherche. Voici la description de ces algorithmes :
 - Sequential Forward Selection (SFS), c'est la première méthode proposée comme algorithme de recherche [88]. Pour constituer le meilleur sous-ensemble de descripteurs, cet algorithme part d'un ensemble vide de descripteurs. A chaque itération, le meilleur descripteur parmi ceux qui restent sera sélectionné, supprimé de l'ensemble de départ et ajouté au sous-ensemble des descripteurs sélectionnés. Le processus continue jusqu'à un critère d'arrêt.
 - Sequential Backward Selection (SBS) : cette méthode a été proposée dans [89]. Elle est similaire à la précédente, à la différence que cette méthode commence avec

l'ensemble des descripteurs et à chaque itération, le descripteur le plus mauvais sera retiré.

- Sequential Forward Floating Selection (SFFS) [88] : est une extension naturelle qui utilise le SFS et le SBS comme algorithmes de recherche, en incluant et excluant certains descripteurs selon la direction dominante de recherche. Il existe aussi le Sequential Backward Floating Selection (SBFS). La méthode SFFS est considérée comme la meilleur méthode de recherche sous optimale [90] (sachant que seule la recherche exhaustive est optimal)

3. Optimisation des paramètres (PO)

La troisième stratégie repose sur une procédure d'optimisation. En pondérant chaque descripteur d d'un exemple $\vec{x}$ par un vecteur de poids $\vec{\omega}_d$, on minimise un critère donné par mises à jour successives de $\vec{\omega}_d$, jusqu'à convergence de l'algorithme. Ainsi le $\vec{\omega}_d$ peut-être assimilé à une notion de pertinence pour chaque descripteur et nous pouvons ainsi effectuer une sélection des descripteurs en ne conservant que les poids les plus importants. En prenant en compte tous les descripteurs à la fois au sein d'une itération, cette stratégie nous permet de prendre en compte les dépendances et les redondances entre les différents descripteurs. Cette approche est plus avantageuse en terme de temps de calculs comparée à la recherche séquentielle. La difficulté de cette méthode est de trouver un critère qui est dérivable par rapport à $\vec{\omega}_d$. De plus, cette méthode n'a pas encore été justifiée d'un point de vue théorique.

C. Différentes taxonomies d'algorithmes

1. Classement ou sélection

Lors de la sélection des descripteurs, nous cherchons à réduire la dimension de notre ensemble de descripteurs, en gardant le plus d'information pertinente possible de l'ensemble départ. Cependant cette définition reste vaste et nous pouvons nous poser au moins deux questions:

- Connaissons-nous la dimension de l'ensemble sélectionné?
- Cherchons-nous à déterminer automatiquement ce nombre conjointement au sous-ensemble sélectionné?

Ces deux approches de sélection sont dans la pratique totalement différente. En effet, il y a:

- Une approche de type classement qui vise à ranger les descripteurs par ordre croissant de pertinence.
- Une approche de type sélection qui vise à extraire de l'ensemble original un sous-ensemble de descripteurs pertinents, dont la taille est déterminée manuellement ou automatiquement.

2. Différentes familles d'algorithmes de sélection des descripteurs : les filtres, les enrouleurs et les embarqués

a. Filtres

Le modèle de type filtre a été le premier utilisé pour la sélection de descripteurs. Dans celui-ci, le critère d'évaluation est utilisé pour estimer la pertinence, au moyen d'un score, d'un descripteur en se basant sur les propriétés des exemples de la base d'apprentissage. Cette étape peut être utilisée comme une étape de prétraitement avant la phase d'apprentissage, car généralement l'évaluation se fait indépendamment du classifieur [86].

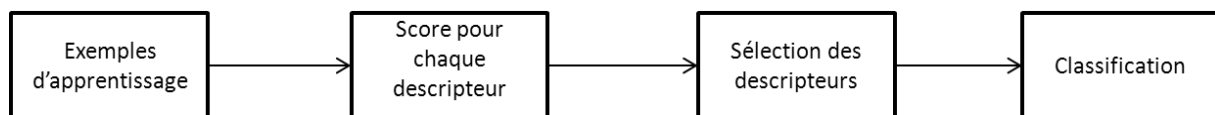


Figure 69: Schéma de la sélection de descripteurs de type filtre

b. Enveloppeurs

La sélection des descripteurs par enveloppeur donne généralement de meilleurs résultats que l'approche par filtre, comme on peut le voir dans [91] et [87]. Généralement nous utilisons le taux de bonnes classifications comme critère pour évaluer les différents sous-ensembles, nous sélectionnons donc à la fin de l'algorithme le sous-ensemble ayant engendré le taux de classification le plus élevé. Ainsi la base d'apprentissage est séparée en deux parties, une partie pour apprendre et une partie pour tester le sous-ensemble sélectionné. Avec cette méthode le sous-ensemble sélectionné est dépendant du classifieur.

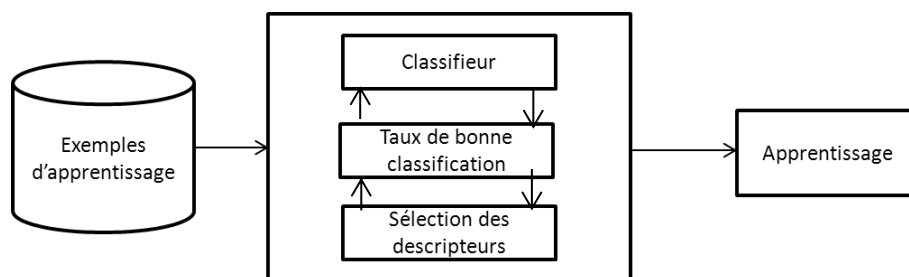


Figure 70: Schéma principe de la sélection de descripteurs de type enrouleur

c. Embarqués

Les méthodes de type embarqué sont différentes dans leur philosophie des approches de type filtre et enveloppante car elles incorporent le mécanisme de la sélection de variables lors du processus d'apprentissage. A la différence de la méthode de type enveloppante elles peuvent se servir de tous

les exemples d'apprentissage pour établir le système. Elles bénéficient d'une rapidité plus élevée que les méthodes de type enveloppantes.

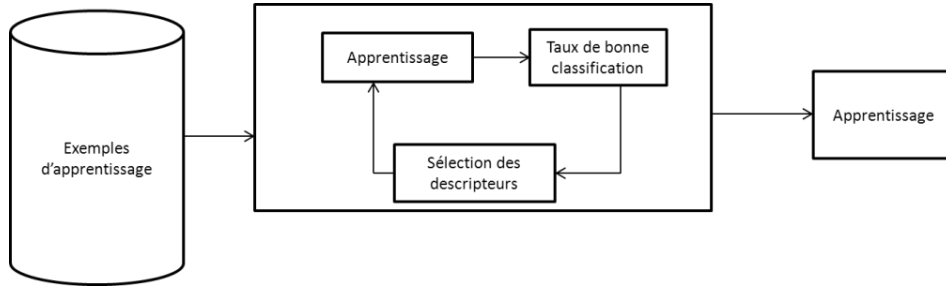


Figure 71: Schéma de principe de la sélection de descripteurs de type embarqué

D. Etat de l'art de différents algorithmes

1. Critère de Fisher

Le critère de Fisher permet de mesurer la séparabilité entre les différents groupes ([92], [93]). Ce dernier est défini par :

$$F(i) = \frac{\sum_{c=1}^K n_c (\mu_c^i - \mu^i)^2}{\sum_{c=1}^K n_c (\sigma_c^i)^2} \quad \text{II-1}$$

où n_c représente le nombre d'éléments composant la classe c , μ_c^i et σ_c^i représente respectivement la moyenne et l'écart-type du $i^{\text{ème}}$ descripteur pour les exemples de la classe c , enfin μ^i est la moyenne des valeurs que prend l' $i^{\text{ème}}$ descripteur sur l'ensemble des classes.

2. Minimum Redundancy Maximum Relevance (MRMR)

L'algorithme de sélection des descripteurs MRMR [94] considère l'information mutuelle pour évaluer le score d'un descripteur utilisé pour réaliser la sélection. Il considère donc dans un même temps les descripteurs pertinents et les descripteurs redondants. Pour définir plus précisément ces deux principes nous définissons l'information mutuelle $I(X, Y)$ entre deux variables aléatoires X et Y :

$$I(X, Y) = \sum_{x, y} p_{X, Y}(x, y) \log \frac{p_{X, Y}(x, y)}{p_X(x) p_Y(y)} \quad \text{II-2}$$

La pertinence du sous-ensemble S de descripteurs pour la classe c est définie par la moyenne des valeurs des informations mutuelles entre chacun des descripteurs f_i et la classe c :

$$D(S, c) = \frac{1}{\text{card}(S)} \sum_{f_i \in S} I(f_i, c) \quad \text{II-3}$$

La redondance des descripteurs dans le sous-ensemble S est définie comme la valeur moyenne des de l'information échangée mutuellement entre le descripteur f_i et le descripteur f_j :

$$R(S) = \frac{1}{(\text{card}(S))^2} \sum_{f_i, f_j \in S, i \neq j} I(f_i, f_j) \quad \text{II-4}$$

Finalement, on obtient le critère MRMR en combinant II-3 et II-4 :

$$MRMR = \max_S (D - R) \quad \text{II-5}$$

3. Diversité marginale maximale (MMD)

a. Distance de Kullback-Liebler

La MMD est basée sur la divergence de Kullback-Liebler [95] qui est définie par:

$$\text{divKL}(P_X || P_Y) = \sum_{x \in S} p_X(x) \log \frac{p_X(x)}{p_Y(y)} \quad \text{II-6}$$

avec X et Y deux variables aléatoires et P_X et P_Y leurs densités de probabilités. Cette divergence est une mesure d'entropie relative entre deux densités de probabilités, le problème est qu'elle n'est pas symétrique car nous avons $\text{div}(P_X || P_Y) \neq \text{divKL}(P_Y || P_X)$. C'est pour cela que nous utiliserons la distance de Kullback-Liebler [95] qui est définie par:

$$dKL(P_X || P_Y) = \text{divKL}(P_Y || P_X) + \text{divKL}(P_X || P_Y) \quad \text{II-7}$$

b. Définition du critère

Pour un problème de classification multi-classes à K classes dans un espace à D dimensions, on cherche à estimer au mieux les probabilités *a priori* $P(C_k)$ des différentes classes. Le moyen le plus simple, à condition que nous ayons une base d'apprentissage représentatif de la réalité, est d'utiliser la formule suivante:

$$\hat{P}(C_k) = \frac{\text{card}(\vec{x}_i \in C_k)}{N} \quad \text{II-8}$$

Ensuite pour chaque descripteur f_i et pour chaque classe C_k , nous cherchons à estimer la probabilité $P(f_i | C_k)$, pour cela nous utilisons l'histogramme que nous nommerons $h_{f_i, k}$. Il est important de noter que chaque histogramme doit contenir le même nombre d'éléments. En pratique, il est très difficile de déterminer ce nombre appelé le pas de quantification de l'algorithme [96]. Il existe deux approches classiques pour déterminer ce pas de quantification n_{hist} :

$$\bullet \quad n_{hist} = \lceil \sqrt{n} \rceil, \quad \text{II-9}$$

avec $\lceil . \rceil$ représentant l'arrondi supérieur.

$$\bullet \quad n_{hist} = 1 + \frac{10 \log(n)}{3}, \quad \text{II-10}$$

Enfin si nous avons un *a priori* sur les données nous pouvons nous même déterminer le nombre et le centre de nos classes, on parle alors de l'alphabet de l'histogramme.

Nous calculons ensuite l'histogramme moyen sur toutes les classes pour un descripteur donné:

$$h(f_i) = \sum_{k=1}^C h_{f_i,k} \quad \text{II-11}$$

Enfin nous pouvons calculer le score MMD défini par:

$$J_{MMD}(f_i) = \sum_{k=1}^C P(C_k) h_{f_i,k}^T \log(h_{f_i,k} / h_{f_i}) \quad \text{II-12}$$

où ./ représente la division élément par élément.

Il a été montré par Vasconcelos [97] que la meilleure solution pour un problème de sélections des axes les plus discriminants est de choisir les axes qui maximisent la diversité marginale maximale.

Cette méthode utilise une stratégie de recherche de type BIN, ainsi nous ne tenons pas compte des interactions entre les différents descripteurs et nous pouvons rencontrer ainsi le phénomène d'interpetinence. C'est pour cela que dans la partie suivante nous avons proposé un nouvel algorithme basé sur le calcul MMD.

E. Extension du critère MMD sur plusieurs dimensions

Pour remédier au problème de non-prise en compte des dépendances entre les différents descripteurs, nous proposons une extension du critère MMD sur plusieurs dimensions. Nous voulons être capables de calculer le critère MMD pour un sous-ensemble de descripteurs. Avant de décrire l'algorithme, introduisons quelques notations. Soient $\Delta = \{1; \dots; D\}$ et $\Omega = \{1; \dots; N\}$ respectivement l'ensemble des indices des descripteurs et l'ensemble des indices des exemples. On appelle $\Delta^{(l)}$ le sous-ensemble sélectionné à la $l^{ème}$ itération de l'algorithme. Nous utilisons pour cet algorithme une recherche séquentielle en utilisant l'algorithme SFFS.

Ensuite pour chaque descripteur f_i et pour chaque classe C_k , nous cherchons à estimer la probabilité $P(f_i|C_k)$, pour cela nous allons donc utiliser l'histogramme que nous nommerons $h_{f_i,k}$.

Pour étendre le principe, nous devons donc estimer la probabilité jointe $P(\Delta^{(l)}|C_k)$, pour cela nous estimons la probabilité $P(f_i|C_k)$ pour chaque descripteur f_i , notée comme précédemment $h_{f_i,k}$, composant le sous-ensemble $\Delta^{(l)}$ sélectionné. Ainsi la probabilité $P(\Delta^{(l)}|C_k)$ sera estimée à l'aide d'un histogramme multidimensionnel avec une dimension égale à la taille du sous-ensemble $\Delta^{(l)}$, on notera ce dernier $h_{\Delta^{(l)},k}$. En pratique il faut noter que les histogrammes doivent être construits avec le même alphabet, afin de pouvoir réaliser une table de contingence pour créer $h_{\Delta^{(l)},k}$. Plus les

données sont maîtrisées, plus l'alphabet à choisir sera facile et meilleure sera l'estimation de $P(\Delta^{(l)}|C_k)$. Il est à noter que dès que $\Delta^{(l)}$ atteint de grandes dimensions il y aura beaucoup de zéros dans l'objet nous permettant d'estimer $h_{\Delta^{(l)},k}$, alors afin d'éviter les problèmes de stockage et de mémoire quand la dimension croît nous utilisons une approche de type *creuse* au sein de notre algorithme.

De plus, $P(\bar{C}, \Delta^{(l)})$ sera estimée par $h_{\Delta^{(l)}}$, défini comme un histogramme multidimensionnel moyen:

$$h_{\Delta^{(l)}} = \sum_{k=1}^C h_{\Delta^{(l)},k} \quad \text{II-13}$$

Ainsi on obtient l'extension du score MMD sur plusieurs dimensions, l'EMMD défini comme suit :

$$J_{MMD}(\Delta_l) = \sum_{k=1}^C P(C_k) h_{\Delta^{(l)},k}^T \log(h_{\Delta^{(l)},k}/h_{\Delta^{(l)}}) \quad \text{II-14}$$

Finalement on calcule le score EMMD pour le sous-ensemble à la $l^{ième}$ étape et on répète cette opération jusqu'à convergence de l'algorithme SFFS. Après la convergence on sélectionne le sous-ensemble qui nous permet d'obtenir le score maximum.

F. Test et résultats

Nous avons effectué des tests sur trois bases provenant toutes les trois du dépôt public UCI [98], qui est un site permettant de trouver des bases de données libres de droit afin d'éprouver les algorithmes d'apprentissage statistique. La popularité de ce dépôt permet de comparer objectivement les résultats des divers auteurs sur une tâche commune. Voici une présentation succincte des trois bases de données utilisées :

- La base lymphoma est une base liée au domaine la bioinformatique et plus précisément l'analyse de données de puces à ADN. Cette base est caractérisée par un très grand nombre de descripteurs basés sur le code génétique. La nécessité d'identifier parmi cette collection de gènes, ceux qui ont une influence sur le phénomène observé est d'ailleurs en grande partie à l'origine de l'essor des techniques de sélection automatique de descripteurs. Le problème décrit par la base contient 96 exemples caractérisés par 4026 descripteurs exprimant le code génétique, et répartis entre les cas sains et les cas malins manifestant la présence d'un lymphome des cellules B. Cette base permettra d'évaluer dans le contexte de la sélection des descripteurs, la fiabilité des algorithmes en présence de nombreux descripteurs fortement redondants.
- La base modélisant le problème de Monk fut la première base de données internationale permettant la comparaison des algorithmes d'apprentissage. Pour plus de détails le lecteur pourra se référer à [99].
 - La base Pima Indians est une étude réalisée sur 768 femmes indiennes. C'est un problème de classification binaire où la classe est égale à 1 quand la patiente montre

les signes du diabète qui sont définis selon les critères de l'organisation mondiale de la santé, -1 dans le cas contraire.

Le protocole d'évaluation est le suivant pour la base Lymphoma :

- Sélection des descripteurs effectuée sur un ensemble d'apprentissage contenant n_{appr} exemples tirés aléatoirement parmi les N exemples de la base. Cette dernière contient n_1 et n_2 exemples pour chacune des deux classes.
- Apprentissage d'un classificateur SVM sur le même ensemble d'apprentissage dont on a sélectionné les R descripteurs les plus pertinents.
- Evaluation de la performance du classifieur sur un ensemble contenant n_{test} individus, avec les R descripteurs sélectionnés à l'étape précédente.

Pour l'évaluation des taux de bonnes classifications dans le cas de la base Pima Indians et de la base représentant le problème de Monk nous avons utilisé une procédure de validation croisée avec 10 parties (voir explication partie 5).

Base	$N(n_1, n_2)$	n_{appr}	n_{test}	D
Lymphoma	96 (34,62)	60	36	4026
Monk	432	Non renseigné	Non renseigné	7
Pima Indians	768	Non renseigné	Non renseigné	9

Tableau 14 : Caractéristiques des bases employées pour l'évaluation

Dans un premier temps nous avons comparé seulement les résultats des différents algorithmes avec la base de données du problème de Monk et la base de données Pima Indians. Les taux de bonnes classifications sont présentés dans deux cas :

- Sélection des descripteurs en utilisant l'EMMD.
- Sans sélection des descripteurs

Et pour deux classifieurs différents :

- le classifieur naïf de Bayes [100], appelé BQ,
- les SVM.

Les résultats sont présentés dans le Tableau 15.

Data	Dimension	BQ	SVM
Monk	D=6	65.51%	84.68%
	R=3	69.44%	100%
Pima Indians	D=8	77.02%	78.73%
	R=4	77.81%	79.29%

Tableau 15: Taux de bonnes classifications avec le classificateur naïf de Bayes et les SVM avec et sans sélection des descripteurs

Nous remarquons que dans les deux cas notre critère améliore la performance du taux de bonnes classifications indépendamment du classifieur utilisé.

Maintenant nous comparons notre critère à deux algorithmes de type filtres qui sont l'algorithme de Fisher et l'algorithme MRMR ainsi qu'avec un algorithme de type enrouleur couplé à un algorithme de parcours SFFS et enfin sans sélection de descripteurs. Le classifieur utilisé est un classifieur de type SVM. Les résultats sont consignés dans le Tableau 16.

Data	Sans Sélection	EMMD	MRMR	Fisher	Enrouleurs avec SFFS
Monk	89.12%	100%	76.61%	87.9%	100%
Pima Indians	79.29%	78.73%	78.51%	78.65%	81,12%
Lymphoma	91.7%	86.5%	90.5%	88.7%	95.6%

Tableau 16: Taux de bonnes classifications avec les SVM, comparaison entre différents algorithmes de sélection des descripteurs

On voit que sur la base de Monk l'EMMD ainsi que l'enrouleur trouvent la combinaison de descripteurs qui engendre un taux de réussite de 100%, alors que les algorithmes de type filtre dégradent la performance de classification sur cet exemple. Cependant sur la base Pima Indians seul l'algorithme de type enrouleur permet d'augmenter la performance des SVM.

Pour la base Lymphoma, les résultats montrent les défauts de certains algorithmes car cette base contient un grand nombre de descripteurs non pertinents, dont l'élimination entraîne une amélioration des performances d'identification. On constate sur cette base la pertinence de la sélection des descripteurs. Cette amélioration de la performance s'explique par le fait que l'information apportée par les descripteurs pertinents se trouve noyée dans la part dominante du bruit de classification portée par le reste des descripteurs. La dégradation de la performance des algorithmes de type filtre montre clairement l'interdépendance des descripteurs pertinents dans l'optimisation du classifieur, qui ne peut être mesurée indépendamment sur chacun d'entre eux. La mauvaise performance de l'algorithme EMMD s'explique par la grande dimension de la base Lymphoma. En effet, l'algorithme EMMD est sensible au fléau de la dimension au niveau de l'estimation de la probabilité jointe lorsque la dimension du sous-ensemble à évaluer devient grande. Seul l'algorithme de type enrouleur avec SFFS obtient de bonnes performances, et ce malgré un temps de calcul plus conséquent que pour les autres algorithmes.

Ainsi ces différents tests ont confirmé les éléments exposés lors de la présentation de chacun de ces algorithmes. Sachant que le calcul de cette combinaison optimal ne se fait qu'une seule fois, nous utiliserons un algorithme de type enrouleur couplé à un algorithme de recherche de type SFFS, malgré le fait que le temps de calcul soit conséquent.

Nous verrons au sein de la prochaine partie comment évaluer le taux de bonnes classifications qui sera le critère de mesure de notre algorithme enrouleur. Nous présenterons ainsi le système de reconnaissance des signaux acoustiques que nous avons conçu.

Chapitre 5: Système automatique de reconnaissance des signaux acoustiques sous-marins

I. Performance d'un modèle

L'apprentissage supervisé effectué par la méthode SVM utilise une partie des exemples pour calculer un modèle de décision qui sera généralisé. Il faut alors avoir des mesures permettant de qualifier le comportement du modèle appris sur les exemples qui ne sont pas utilisés lors de l'apprentissage. Ces métriques sont calculées soit sur les exemples d'apprentissage eux-mêmes ou sur des exemples réservés à l'avance pour les tests.

A. Métrique de performance

Dans cette section, nous allons présenter des métriques servant pour l'évaluation de performance.

1. Taux de bonnes classifications

Il représente le rapport entre le nombre d'exemples bien classés et le nombre total d'exemples :

$$\text{Perf} = \frac{1}{N} \sum_{i=1}^N Q(y_i, \hat{f}_i) \quad \text{I-1}$$

avec

$$Q = \begin{cases} 1 & \text{si } y_i = \hat{f}_i \\ 0 & \text{sinon} \end{cases}$$

On peut multiplier Perf par 100 pour avoir ce résultat sous forme de pourcentage.

2. Matrice de confusion

Le taux de bonnes classifications nous donne un taux global de bonnes classifications qui ne permet pas de connaître la nature des erreurs du système de classification. Or il est intéressant de connaître la nature de nos erreurs, afin de savoir si une classe est souvent confondue avec une autre afin de réaliser certaines opérations sur nos données. Dans le cas d'une classification binaire 4 cas sont possibles :

- $\hat{f}(x_i) = 1$ et $y_i = 1$, correcte positive (CP) ;
- $\hat{f}(x_i) = 1$ et $y_i = -1$, fausse positive (FP) ;
- $\hat{f}(x_i) = -1$ et $y_i = -1$, correcte négative (CN) ;
- $\hat{f}(x_i) = -1$ et $y_i = 1$, fausse négative (FN) ;

Ainsi dans le cas d'un problème de classification binaire, la matrice de confusion C est définie ainsi :

$$C = \begin{pmatrix} CP & FN \\ FP & CN \end{pmatrix} \quad \text{I-2}$$

Dans le cas d'une classification parfaite $FN = FP = 0$ et ainsi les résultats seront concentrés dans la diagonal. On peut retrouver le taux de bonnes classifications à partir de la matrice de confusion, grâce à la formule suivante :

$$P = \frac{CP + CN}{CP + CN + FP + FN} \quad \text{I-3}$$

Pour un problème multi-classes la matrice sera de ce type mais sera de la taille $K \times K$. Le résultat I-3 peut être généralisé pour un problème multiclassé, dans ce cas la formule devient :

$$P = \frac{\sum_{k=1}^K C(k, k)}{\sum_{k=1}^K \sum_{l=1}^K C(k, l)} \quad \text{I-4}$$

Dans la littérature [101], on trouve quelques autres métriques de performances qui sont :

- La Sensitivité :

$$S_v = \frac{CP}{CP + FP} \quad \text{I-5}$$

- La Spécificité :

$$S_p = \frac{CN}{CN + FN} \quad \text{I-6}$$

- La moyenne harmonique :

$$Mh = \frac{2S_v S_p}{S_v + S_p} \quad \text{I-7}$$

B. Evaluation des performances

Nous avons vu précédemment que les frontières obtenues grâce au modèle SVM dépendent de plusieurs facteurs, à savoir:

- le paramètre de contrainte sur les multiplicateurs de Lagrange C ;
- le noyau utilisé K ;
- les hyperparamètres engendrés par le choix du noyau, plus précisément σ dans le cas RBF gaussien;
- la base d'apprentissage, plus précisément les vecteurs supports;
- les descripteurs de manière indirecte mais qui influent énormément sur la qualité, car plus la séparabilité est grande dans l'espace de départ, plus facile sera le problème SVM.

Le choix de ces valeurs se fait à travers plusieurs essais afin d'obtenir le modèle qui atteindra les meilleures performances. Les paramètres idéaux seraient ceux qui nous permettent d'avoir un taux de bonnes classifications égal à 100%. Cette situation serait idéale si la base d'apprentissage que nous avons choisi, était parfaitement représentative de la réalité, or nous pouvons obtenir un taux de 100% sur la base d'apprentissage et avoir un mauvais taux de bonnes classifications en généralisation, c'est le sur-apprentissage. Ainsi il faut pouvoir quantifier la qualité de notre modèle autrement qu'en se basant sur le taux de bonnes classifications lors de l'apprentissage. Nous présentons alors dans les prochaines sections différentes méthodes d'évaluation.

1. Méthode HoldOut

Cette méthode consiste à séparer l'ensemble des données disponibles en deux parties, une partie pour l'apprentissage du modèle et une partie pour le test du modèle. Le test du modèle obtenu sur la partie de test permet de se donner une idée du modèle en généralisation car nous utilisons des exemples de tests qui n'ont pas servi lors de l'apprentissage. Ainsi à l'issue de cette méthode nous sélectionnons le modèle qui maximise le taux de bonnes classifications sur la partie des données réservées aux tests.

La question importante du choix de la partie de test et de la partie de l'apprentissage lors de l'utilisation de cette technique a une forte influence sur la qualité du modèle.

2. Validation croisée

Cette méthode a été conçue pour minimiser l'influence du choix du partitionnement des données. Cette méthode consiste à diviser nos données en k parties disjointes de taille à peu près égale. Ainsi une phase d'apprentissage est effectuée sur $k - 1$ parties et la phase de test sur la partie restante, cette opération est réalisée de manière circulaire en changeant à chaque fois la partie à tester. On obtient donc k taux de bonnes classifications. La précision du modèle sera égale alors à la moyenne des k taux de bonnes classifications.

La méthode Leave-One-Out (LOO) est un cas particulier de la validation croisée pour laquelle $k = N$, c'est-à-dire que nous divisons notre ensemble de départ par le nombre d'individus le composant et on apprend à chaque fois sur $N - 1$ exemples et on teste l'exemple restant. Cette méthode permet de simuler plus précisément le cas de la généralisation, malheureusement il est coûteux en temps de calcul dès que notre base devient conséquente.

3. Bootstrap

La méthode Bootstrap, appelée aussi échantillonnage par remplacement, entraîne le modèle sur un exemple de N exemples choisis aléatoirement de l'ensemble des exemples, des exemples peuvent être choisis plusieurs fois tandis que d'autres peuvent ne pas être choisis. Les exemples non choisis pour l'entraînement sont choisis pour le test. Cette opération est répétée plusieurs fois et nous obtenons finalement une précision du modèle en effectuant la moyenne des précisions. Le Bootstrap est basé donc sur la méthode Monte-Carlo.

Parmi les différentes méthodes Bootstrap l'une des plus utilisées est la méthode ".632". Elle tire son nom du fait que 63.2% des exemples contribuent à l'entraînement et ceux restant participent aux tests.

En effet à chaque prélèvement, un exemple a la probabilité $\frac{1}{N}$ d'être sélectionné et $\left(1 - \frac{1}{N}\right)$ de ne pas l'être, et puisqu'on l'on répète l'opération N fois, chaque exemple aura une probabilité $\left(1 - \frac{1}{N}\right)^N$ de ne pas être sélectionné du tout dans l'ensemble d'apprentissage. Si N est grand on a :

$$\lim_{N \rightarrow \infty} \left(1 - \frac{1}{N}\right)^N = e^{-1} = 0.368$$

I-8

La méthode répète le processus k fois et le taux de bonnes classifications du modèle est donnée par :

$$P = \sum_{i=1}^k (0.632 \times P_{i_{test}} + 0.368 \times P_{i_{app}}) \quad \text{I-9}$$

Avec $P_{i_{test}}$ le taux de bonnes classifications du modèle appris à l'itération i et $P_{i_{app}}$ le taux de bonnes classifications des données de test à l'itération i sur le modèle appris à cette même itération.

Après avoir introduit les mesures de performance, un système de reconnaissance des signaux acoustiques sous-marins va être présenté.

II. Description du système automatique d'identification des signaux acoustiques sous-marins

Nous allons présenter le système automatique d'identification des signaux acoustiques sous-marins que nous avons décidé d'implémenter. La Figure 72, présente ce dernier, ainsi chaque bloc représente une partie du système qui sera décrit dans les parties suivantes.

A. Segmentation temps-fréquence

Dans cette thèse nous considérons cette étape comme un système boîte noire. L'opération de segmentation peut être vue comme la donnée d'un sous-ensemble du plan temps-fréquence appelé pavé.

B. Classification manuelle et création des classes

A l'aide d'un expert nous devons d'abord créer différents groupes. Cette étape est délicate car nous devons choisir le bon niveau de granularité, sachant qu'un système de classification automatique ne pourra pas atteindre le niveau de détail atteint par l'être humain. Il faut donc créer des classes en regroupant ce qui se ressemble, les classes doivent être homogènes.

Une fois les classes définies, nous devons étiqueter les différents pavés temps-fréquences issus de l'étape de segmentation. Cette étape est très importante car elle nous servira par la suite pour l'apprentissage de notre système ainsi que pour l'évaluation des performances. C'est pour cela qu'elle nécessite l'aide d'un expert afin de minimiser la probabilité d'erreur, car nous étiquetons un grand nombre de pavés. Il faut veiller aussi à la répartition des effectifs qui composent chaque classe.

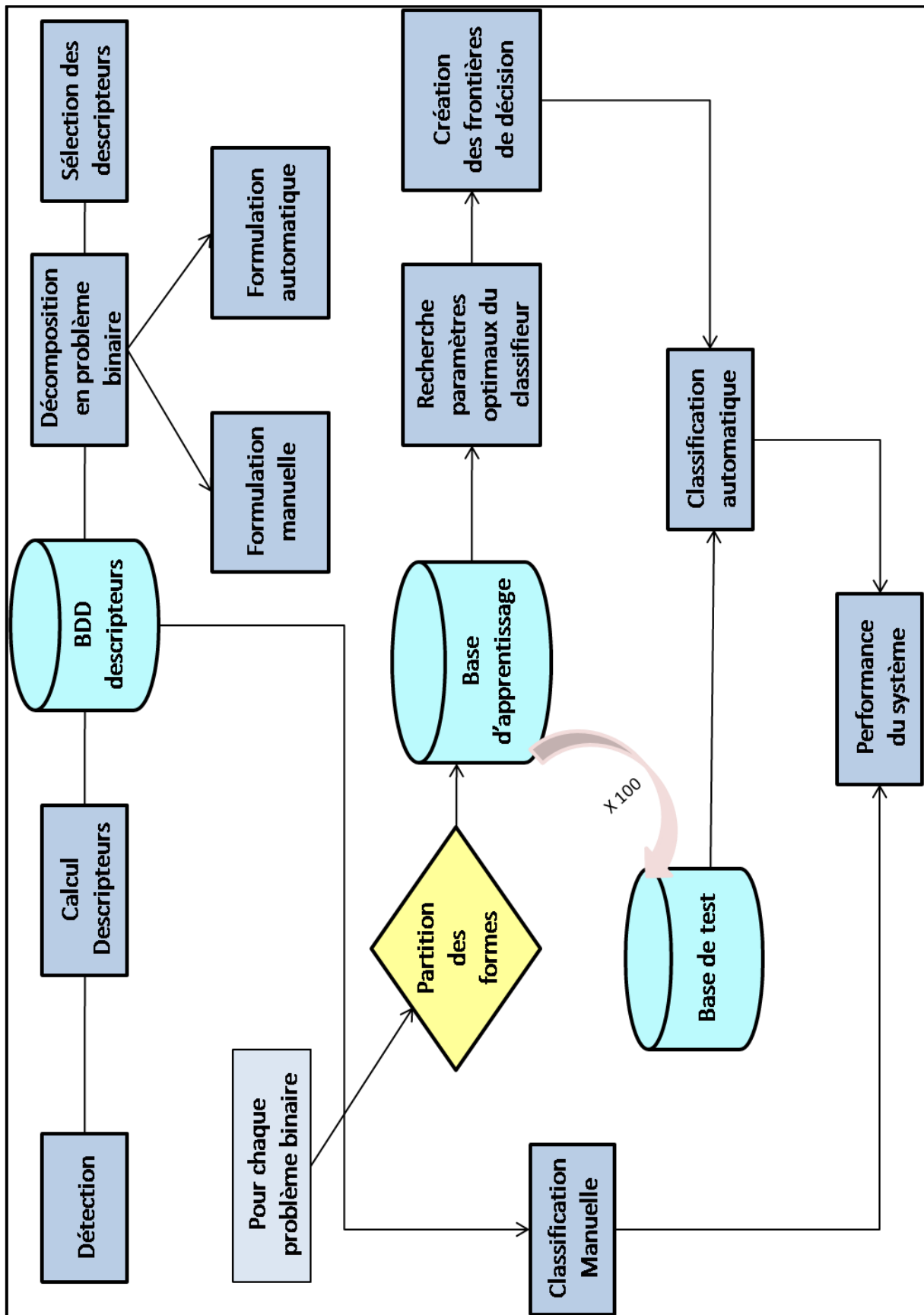


Figure 72: Système de reconnaissance automatique

En effet, dans le cas d'une base d'apprentissage optimale nous voudrions que les effectifs de la base soient représentatifs de la réalité.

C. Calcul des descripteurs

C'est lors de cette étape que les descripteurs sont calculés. Nous utilisons ainsi des descripteurs morphologiques se basant à la fois sur la représentation temporelle, la représentation fréquentielle et la représentation temps-fréquence du signal en utilisant le *Denoised Hearingogram*. De plus, sont ajoutés à ces derniers des descripteurs cepstraux qui proviennent des techniques de reconnaissance vocale et enfin sont utilisés également des descripteurs perceptuels, basés sur de la physiologie humaine.

D. Décomposition du problème en problème binaire

Le problème d'identification présenté dans ce manuscrit est de type multiclasses. Nous avons vu précédemment qu'en termes de résultats et de temps de calculs la meilleure option était de décomposer notre problème en plusieurs problèmes à deux classes. Les techniques les plus connues sont le « un contre un » ainsi que le « un contre tous », cependant il existe d'autres techniques permettant de réaliser une classification de manière hiérarchique sous forme d'arbre. En d'autres termes nous regroupons les différentes classes en nous basant sur un critère de ressemblance. Deux choix s'offrent à nous:

- La formulation manuelle d'un arbre à l'aide d'un expert, l'idée est tout d'abord de fusionner les classes les plus similaires et ensuite les séparer petit à petit. L'avantage de cette approche est que l'arbre obtenu colle à la réalité physique des signaux et le principe de décomposition des experts.
- Nous réalisons un dendrogramme de manière automatique afin de fusionner les classes, en nous basant sur une classification ascendante hiérarchique comme vu dans la partie 3, l'arbre sera ainsi créé de manière automatique.

L'objectif est que l'approche manuelle et automatique convergent vers un arbre unique.

E. Sélection des descripteurs

Maintenant que le problème a été décomposé en plusieurs problèmes binaires, il faut réaliser une sélection des descripteurs propres à chaque problème d'identification binaire. Ceci est dû au fait que par essence les SVM ont été créés pour résoudre un problème binaire. De plus, il est raisonnable de penser que prendre un seul et même ensemble de descripteurs pour discriminer toutes les classes n'est pas la configuration optimale. Ceci contraint en effet à faire intervenir à chaque fois tous les descripteurs pour la reconnaissance de chaque classe vis-à-vis des autres. On préférera donc la configuration qui consiste à rechercher pour chaque nœud de décision le meilleur jeu de descripteurs et ainsi optimiser la classification localement plutôt que globalement.

Nous avons vu dans la partie 4 qu'il y a un jeu de descripteurs optimal.

Sachant que pour cette étape le temps de calcul n'est pas primordial, dans la limite du raisonnable, car nous faisons cette étape en amont de l'utilisation du système une méthode de type enveloppeur est utilisée. Malheureusement la puissance de calcul actuelle de nos ordinateurs ne permet pas encore de faire une recherche exhaustive des différentes combinaisons. Cependant il existe des algorithmes de recherche intelligents permettant d'explorer astucieusement l'espace des combinaisons de descripteurs possibles. Ainsi l'algorithme SFFS est sélectionné pour réaliser cette tâche de recherche. A l'issue de cet algorithme la configuration retenue sera celle qui aura engendré le taux de bonnes classifications le plus élevé estimé par la méthode leave-one-out.

Nous allons donc pour chaque problème d'identification binaire utiliser cette stratégie de sélection des descripteurs. Il faut attirer l'attention sur un point crucial : le choix de la base de données et l'étiquetage des exemples ont un impact important sur les résultats et cette base doit donc avoir été réalisée avec le plus grand soin et par des experts du domaine.

F. Paramétrage du classifieur SVM

Cette étape est aussi réalisée pour chaque problème bi-classe. Ceci est dû au fait que chaque problème est différent nous devons donc optimiser chaque problème bi-classes afin que la classification multi-classe puisse engendrer des bonnes performances.

Le choix d'un noyau RBF Gaussien a été fait, car il permet de projeter les données dans un espace de dimension infinie, ainsi la probabilité de trouver un hyperplan séparateur augmente dans l'espace de transformation, de plus ce noyau n'engendre qu'un extra paramètre qui est l'écart-type σ de la gaussienne.

Une stratégie de recherche par maillage est utilisée, car dans le cas de deux paramètres la complexité reste raisonnable. Il est essentiel de rechercher simultanément le couple optimal car nous ne pouvons pas chercher l'un et l'autre indépendamment (nous avons vu dans la partie 3 que ces deux paramètres étaient fortement liés). Au niveau du maillage, nous choisirons pour C des valeurs réparties logarithmiquement entre 10^{-6} et 10^6 et pour la valeur σ des valeurs comprises entre 0.1 et 20 réparties aussi logarithmiquement. Le couple qui engendrera le taux de bonnes classifications, maximal, estimé à l'aide du leave-one-out, sera sélectionné.

G. Création des frontières

A ce stade les descripteurs et les paramètres servant à la création des frontières SVM ont été calculés. Nous allons donc utiliser les valeurs obtenues et résoudre le problème d'optimisation défini dans la partie 3. Pour cela nous résolvons le problème dual à l'aide d'un algorithme d'optimisation nommé SMO¹⁶ qui a été proposé premièrement par [102]. De nos jours, cet algorithme est le plus utilisé dans la littérature pour les problèmes de grande dimension. Il consiste à optimiser à chaque itération, deux multiplicateurs de Lagrange conjointement.

¹⁶ Sequential Minimal Optimization

H. Mesure de performance

Pour qualifier le système il faut être capable de mesurer sa performance, de plus nous avons vu que l'estimation du taux de bonnes classifications servait au paramétrage des SVM et à la sélection des descripteurs.

Deux choix s'offrent à nous suivant la taille de notre base d'apprentissage:

- utiliser l'estimation leave-one-out

- utiliser la validation croisée 10-*fold*, c'est-à-dire que l'on utilise $\frac{9}{10}$ de la base pour l'apprentissage et on teste sur les $\frac{1}{10}$ et on effectue une permutation des paquets, il est à noter que le nombre de fold optimal à utiliser durant la validation croisée est dépendant des données. Sachant que la partition des signaux en paquets est aléatoire il est nécessaire de répéter l'opération un grand nombre de fois afin de ne pas biaiser l'estimation, c'est une approche de type Monte-Carlo, où l'expérience est répétée 100 fois. Dans la première approche chaque élément de la base de test a reçu un tag de classification de manière automatique, et on le compare avec l'étiquette qui a été donnée manuellement par l'expert, en faisant ainsi on peut estimer la performance du système d'identification automatique.

Ce système a été implémenté en MATLAB[®] et ensuite porté en C++, il est le fruit de ces 3 années de thèse. Cependant des évolutions sont envisageables, nous les exposerons lors des perspectives.

Les résultats sur signaux réels sont prometteurs, et vont être testés en situation opérationnelle. Malheureusement, pour des raisons évidentes dues à la confidentialité des signaux réels, les résultats ne peuvent pas être exposés en détail dans le manuscrit. Nous pouvons néanmoins dire qu'une amélioration des résultats a été observée par rapport à l'ancien système d'identification.

Conclusion générale

Nous avons traité dans cette thèse la question de la représentation et de la reconnaissance des signaux acoustiques sous-marins. Le travail mené au cours de cette thèse a permis d'obtenir un système de reconnaissance automatique des signaux acoustiques sous-marins.

L'architecture de notre système final exploite un schéma de classification hiérarchique qui repose sur une taxonomie définie à l'aide d'experts en reconnaissance acoustique. Ce système est principalement constitué de trois grands modules :

- Représentation du signal à identifier ;
- Description du signal d'après la représentation précédente ;
- Reconnaissance du signal.

Le premier module concerne la représentation des signaux acoustiques sous-marins, nous avons réalisé un état de l'art des techniques des représentations temps-fréquence qui sont adaptées à la non-stationnarité des signaux réels. Ensuite partant du postulat que l'humain est le meilleur des classifieurs, nous avons construit une représentation, l'*Hearingogram*, basée sur la physiologie humaine en utilisant les filtres de Mel. Les résultats présentés ont montré une amélioration du spectrogramme dans les différentes expérimentations, pouvant ainsi faciliter l'identification automatique de certains phénomènes.

La seconde partie de ce module concerne la réduction du bruit au sein des signaux acoustiques sous-marins, nous avons donc comparé différentes techniques de l'état de l'art et confronté les résultats obtenus à un algorithme de réduction du bruit de l'*Hearingogram* : le *Denoised Hearingogram*. Les résultats de cet algorithme sont très intéressants. Bien qu'ils restent néanmoins proches de ceux obtenus par certaines approches de l'état de l'art, cette méthode nécessite peu, voir pas, de réglages de paramètres contrairement aux autres techniques. Cela est un avantage non négligeable pour l'implantation dans un système automatique.

Ces différents travaux sur l'*Hearingogram* et le *Denoised Hearingogram* ont mené à trois actes de conférences [103] [104] [105].

Le second module du système concerne la description du signal. Afin de produire un ensemble de descripteurs efficace, nous avons expérimenté plusieurs descripteurs de l'état de l'art de plusieurs types tel que morphologiques, statistiques, cepstraux et perceptuels. Les plus efficaces de ces descripteurs ont été retenus au moyen d'un algorithme de type enrouleur avec un algorithme d'exploration SFFS, qui reste la méthode la plus efficace malgré un temps de calcul conséquent. L'emploi de méthodes automatiques de sélection des descripteurs se justifie par le fait que la notion de pertinence est très complexe et ne peut être jugée indépendamment sur chaque descripteur.

Ensuite, un algorithme de sélection des descripteurs a été développé, ce dernier est basé sur une extension sur plusieurs dimensions du critère MMD, il s'agit de l'EMMD. Cependant, malgré des

résultats prometteurs sur certaines bases de données, cet algorithme a un défaut majeur lorsque l'on applique sur des données en grande dimension, il s'agit du fléau de la dimension.

Il est à noter que nous effectuons cette opération de sélection de façon binaire en recherchant un sous-ensemble d'attributs optimal pour la discrimination de chaque paire de classes possibles. En plus d'être performante, cette méthode offre la possibilité d'acquérir une meilleure compréhension du problème d'identification et de suggérer des voies d'amélioration du système.

Les travaux sur l'algorithme EMMD ont donné lieu à un acte de conférence [106].

Par la suite, nous nous sommes penchés sur le module d'identification. Nous avons choisi d'utiliser les machines à vecteur support. Cependant, les SVM s'appuient sur des hypothèses contraignantes qui nous ont obligés à étudier les méthodes d'extension à plus de deux classes. Elles ont été utilisées à base de décision binaire, qui s'appuie sur un arbre de classification, chaque nœud de l'arbre est une décision binaire à prendre à l'aide des SVM. De plus, une étude de la sélection de paramètres efficaces a été réalisée et nous avons donc mis en place une procédure de sélection par maillage.

Enfin un effort important a été consacré à la constitution d'une base de données de signaux acoustiques sous-marins et sur la création de classes permettant l'évaluation des systèmes proposés. Malheureusement, pour des raisons de confidentialité nous ne pouvons pas communiquer à propos de cette base de données.

Les différents choix du système d'identification sont exposés dans le dernier chapitre de ce manuscrit, avec la justification de chaque choix. Ainsi les différentes mesures de performance ont montré une amélioration des résultats par rapport à ce qui était fait précédemment au sein de l'entreprise.

ANNEXE A : Reconstruction du signal temporel à partir de la transformée de Fourier à court terme

A. Signal

Soit x_n un signal à temps discret $\forall n = 1 \dots N$ dont les échantillons sont contenus dans le vecteur x défini par

$$x = [x_1, \dots, x_N]^T \quad 0-1$$

B. Fenêtre d'observation du signal

L'information contenue dans x est observable sur N_h échantillons, avec $1 \leq N_h \leq N$, de sorte que toute manifestation lorsqu'observée parmi les N échantillons est considérée stationnaire lorsqu'observée parmi les N_h échantillons des fenêtres auxquels il est rattaché.

Ainsi, le signal est observé avec une fenêtre de pondération h_n , définie $\forall n = 1 \dots N_h$, de puissance

$$P_h = \frac{1}{N_h} \sum_{n=1}^{N_h} h_n^2 \quad 0-2$$

Les éléments de la fenêtre sont non-nuls,

$$h_n \neq 0 \quad \forall n = 1 \dots N_h \quad 0-3$$

Le cas échéant, soit h_n^0 une fenêtre contenant des éléments nuls. Soit I_h^0 l'ensemble des indices des éléments nuls de la fenêtre,

$$I_h^0 = \{n / h_n^0 = 0, \quad \forall n = 1 \dots N_h\} \quad 0-4$$

avec

$$\text{card}[I_h^0] = N_h^0 < N_h, \quad 0-5$$

le nombre d'éléments nuls de h_n^0 .

De façon complémentaire, l'ensemble I_h^{0C} contient les éléments non-nuls de la fenêtre. Afin de construire une fenêtre h_n sans éléments nuls et conservant la puissance P_{h^0} , les éléments nuls de h_n^0 sont relevés d'une faible quantité $\varepsilon_h \ll 1$ et le coefficient de pénalisation $\alpha_h > 0$ est appliqué aux éléments de I_h^{0C} , si bien que h_n est définie par

0-6

$$h_n = \begin{cases} \varepsilon_h, & \text{si } n \in I_h^0 \\ \alpha_h h_n^0, & \text{si } n \in I_h^{0C} \end{cases} \quad \forall n = 1 \dots N_h,$$

Sa puissance a pour expression

0-7

$$\begin{aligned} P_h &= \frac{1}{N_h} \sum_{n=1}^{N_h} h_n^2 \\ &= \frac{1}{N_h} \sum_{n=1}^{I_h^0} \varepsilon_h^2 + \frac{1}{N_h} \sum_{n=1}^{I_h^{0C}} (\alpha_h h_n^0)^2 \\ &= \frac{N_0^h}{N_h} \varepsilon_h^2 + \frac{\alpha_h^2}{N_h} \sum_{n=1}^{I_h^{0C}} (h_n^0)^2, \\ &= \frac{N_0^h}{N_h} \varepsilon_h^2 + \frac{\alpha_h^2}{N_h} \sum_{n=1}^{N_h} (h_n^0)^2 \\ &= \frac{N_0^h}{N_h} \varepsilon_h^2 + \alpha_h^2 P_{h_0} \end{aligned}$$

La conservation de la puissance permet de fixer α_h ,

0-8

$$\begin{aligned} P_{h_0} &= P_h && \Leftrightarrow \\ P_{h_0} &= \frac{N_0^h}{N_h} \varepsilon_h^2 + \alpha_h^2 P_{h_0} && \Leftrightarrow \\ \alpha_h^2 P_{h_0} &= P_{h_0} - \frac{N_0^h}{N_h} \varepsilon_h^2 && \Leftrightarrow, \\ \alpha_h^2 &= 1 - \frac{N_0^h}{N_h} \frac{\varepsilon_h^2}{P_{h_0}} && \Leftrightarrow \\ \alpha_h &= \sqrt{1 - \frac{(N_0^h / N_h) \varepsilon_h^2}{P_{h_0}}} \end{aligned}$$

C. Observation fenêtrée du signal

Deux fenêtres consécutives peuvent se recouvrir sur N_r échantillons, avec $0 \leq N_r \leq N_h - 1$.

Le nombre de fenêtres nécessaire pour observer les N échantillons du signal est le nombre de récurrences L défini par :

0-9

$$L = \left\lfloor \frac{N - N_r}{N_h - N_r} \right\rfloor,$$

où $\lfloor . \rfloor$ représente l'arrondi à l'inférieur.

Le signal x_n est prolongé afin de compléter la $L^{\text{ème}}$ récurrence, il est alors constitué de N_0 échantillons avec $N_0 = (L - 1)(N_h - N_r) + N_h$,

0-10

$$x = [x_1, \dots, x_N, \dots, x_{N_0}]^T$$

L'exemple ci-dessous illustre le procédé d'analyse fenêtré venant d'être décrit.

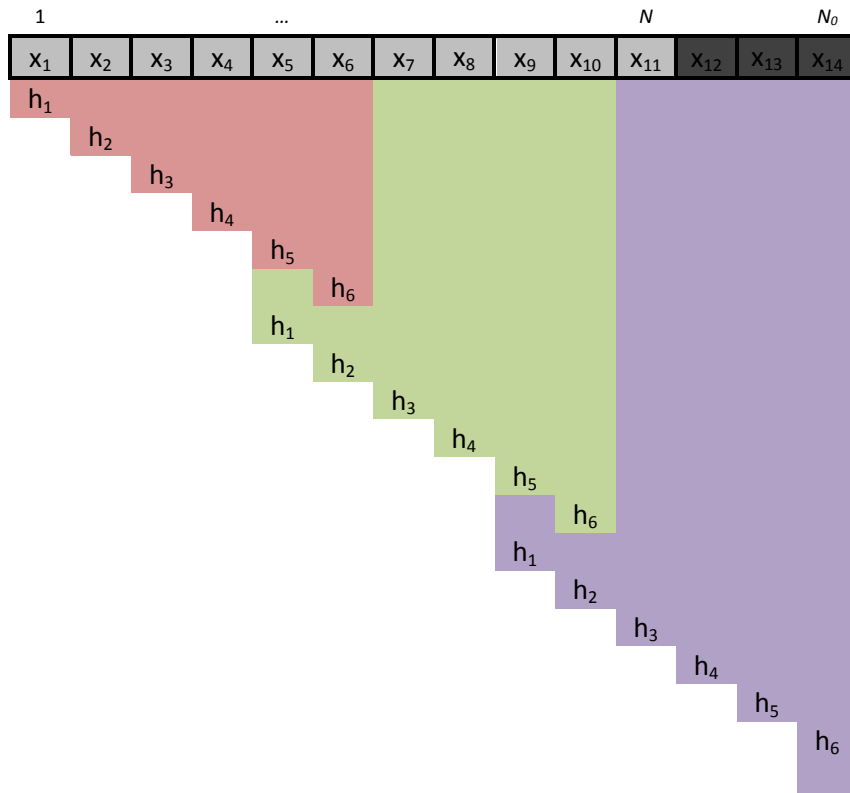


Figure 73 : Analyse temporelle fenêtrée du signal, avec $N = 11$, $N_h = 6$, $N_r = 2$, $N_0 = 14$, $L = 3$.

D. Analyse harmonique du signal

Les échantillons contenus dans chaque fenêtre sont analysés par transformée de Fourier discrète calculée sur N_{FFT} points, avec $N_{FFT} \geq N_h$ et $\exists q_{FFT} / N_{FFT} = 2^{q_{FFT}}$.

La transformée de Fourier à court terme du signal est définie pour chaque fenêtre par

$$X_{kl} = \sum_{n=1}^{N_h} x_{(l-1)(N_h-N_r)+n} h_n e^{-2i\pi nk/N_{FFT}} \quad \begin{matrix} \forall l = 1 \dots L \\ \forall k = 1 \dots N_{FFT} \end{matrix} \quad \text{0-11}$$

E. Expression algébrique de l'analyse harmonique fenêtrée du signal

La partie suivante a pour but de décrire de façon algébrique l'enchaînement des opérations de fenêtrage et d'analyse harmonique décrites précédemment, conduisant à une analyse harmonique fenêtrée du signal.

F. Expression matricielle des opérateurs

○ Matrice de fenêtrage R

Soit R la matrice binaire de dimension (LN_h, N_0) qui, lorsque appliquée au vecteur x , ordonne ses échantillons en concaténant le contenu de chacune des L récurrences. Les indices des éléments non-nuls de R sont définis par :

$$\begin{cases} I(n, l) = (l-1)N_h + n \\ J(n, l) = (l-1)(N_h - N_r) + n \end{cases} \quad \forall l = 1 \dots L, \quad \forall n = 1 \dots N_h \quad \text{0-12}$$

La matrice R a pour terme général

$$R_{ij} = \delta[i - I(n, l); j - J(n, l)], \quad \begin{matrix} \forall i = 1 \dots LN_h \\ \forall j = 1 \dots N_0 \end{matrix} \quad \text{0-13}$$

Comme $(n, l) \geq J(n, l)$, la matrice R est triangulaire inférieure. Comme de plus $R_{ij} \in \{0; 1\}$, la somme des lignes de la matrice R est unitaire,

$$\sum_{j=1}^{N_0} R_{ij} = 1, \quad i = 1 \dots LN_h \quad \text{0-14}$$

Par exemple, si $N = 11$, $N_h = 6$, $N_r = 2$, $N_0 = 14$, $L = 3$, alors R est de dimension $(18, 14)$ et a pour expression :

0-15

$$R_{(LN_h, N_0)} = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

○ **Matrice de pondération H**

Soit h le vecteur défini par

$$h = [h_1, \dots, h_{N_h}]^T \quad \text{0-16}$$

alors H est la matrice diagonale de dimension (N_h, N_h) de terme général

$$H_{nm} = h_n \delta[n - m] \quad \begin{matrix} \forall n = 1 \dots N_h \\ \forall m = 1 \dots N_h \end{matrix} \quad \text{0-17}$$

soit

$$H = \begin{bmatrix} h_1 & & 0 \\ & \ddots & \\ 0 & & h_{N_h} \end{bmatrix} \quad \text{0-18}$$

○ **Matrice de 0-padding P**

Soit P la matrice de 0-padding définie par

$$\underset{(N_{FFT}, N_h)}{P} = \begin{bmatrix} \text{Id}_{N_h} \\ 0 \end{bmatrix}_{(N_{FFT}-N_h, N_h)} \quad \text{0-19}$$

Elle vérifie

$$P^T P = \text{Id}_{N_h} \quad \text{0-20}$$

○ Matrice de Fourier W

Soit W la matrice de Fourier de terme général

$$W_{km} = e^{-\frac{2i\pi(m-1)(k-1)}{N_{FFT}}} \quad \text{0-21}$$

Elle vérifie

$$W^H W = N_{FFT} \text{Id}_{N_{FFT} N_{FFT}} \quad \text{0-22}$$

G. Expression algébrique de l'analyse et de la synthèse harmonique fenêtrée

○ Analyse

A partir du formalisme établi précédemment, l'expression matricielle de la transformée de Fourier à court terme, conduisant à l'analyse harmonique fenêtrée du signal, est donnée par

$$\underset{(LN_{FFT})}{X} = (\text{Id}_L \otimes WPH)(R x) \quad \text{0-23}$$

Le vecteur X est de type colonne constitué des L transformées de Fourier du signal x_n fenêtré, calculées sur N_{FFT} points.

Le calcul de l'expression 0-23 est décomposé en deux étapes :

- le changement de variable

$$\underset{(LN_h)}{\hat{X}} = R x \quad \text{0-24}$$

- le calcul de la matrice A définie par

$$\underset{(LN_{FFT}, LN_h)}{A} = \text{Id}_L \otimes WPH \quad \text{0-25}$$

Ainsi, l'analyse du signal revient à effectuer

$$\begin{cases} \hat{x} = Rx \\ X = A\hat{x} \end{cases} \quad \text{0-26}$$

L'expression 0-26 montre que les opérations de fenêtrage (i.e. R) et de traitement (i.e. A) sont découplées. De façon condensée, 0-26 s'écrit :

$$\underset{(LN_{FFT})}{X} = \underset{(LN_h)}{A} (Rx) \quad \text{0-27}$$

○ Synthèse

Réciproquement, l'opération de synthèse harmonique fenêtrée du signal revient à inverser la relation 0-27. Pour ce faire, soit $A^\#$ et $R^\#$ les matrices pseudo-inverses de Moore-Penrose de A et R définies respectivement par

$$\underset{(LN_h, LN_{FFT})}{A^\#} = \left(A^H A \right)^{-1} A^H \quad \text{0-28}$$

et

$$\underset{(N_0, LN_h)}{R^\#} = \left(R^H R \right)^{-1} R^H \quad \text{0-29}$$

Ces matrices existent si

$$\begin{cases} \det(A^H A) > 0 \\ \det(R^H R) > 0 \end{cases} \quad \text{0-30}$$

○ Expression algébrique condensée

La vérification des deux conditions 0-30 est donc nécessaire pour réaliser l'opération de synthèse.

Dans ce cas, la synthèse harmonique fenêtrée du signal est donnée par

$$\begin{cases} \hat{x} = A^\# X \\ x = R^\# \hat{x} \end{cases} \quad \text{0-31}$$

Ou de façon condensée,

$$\underset{(N_0)}{x} = \underset{(LN_h)}{R^\#} \left(\underset{(LN_h)}{A^\#} X \right) \quad \text{0-32}$$

L'expression détaillée est obtenue après avoir explicité $A^\#$ et $R^\#$ en donnant

- leur condition d'existence, selon 0-30,
- leur expression.

◦ **Expression de $A^\#$**

A partir de la définition 0-28 de $A^\#$, de la définition 0-25 de A , et des propriétés 0-18, 0-20 et 0-22 des matrices W, P, H ,

$$\begin{aligned} A_{(LN_h, LN_h)}^H A &= (\text{Id}_L \otimes WPH)^H (\text{Id}_L \otimes WPH) \\ &= (\text{Id}_L \otimes HP^T W^H) (\text{Id}_L \otimes WPH) \\ &= \text{Id}_L \otimes (HP^T W^H WPH) \\ &= N_{FFT} \text{Id}_L \otimes H^2 \end{aligned} \quad 0-33$$

Cette matrice est diagonale, de déterminant

$$\det(A^H A) = N_{FFT}^{LN_h} \left(\prod_{n=1}^{N_h} h_n^2 \right)^L \quad 0-34$$

Compte tenu de la propriété 0-3 de non-nullité des éléments h_n , la condition d'inversibilité de $A^H A$ est toujours vérifiée :

$$\det(A^H A) \neq 0 \Leftrightarrow h_n \neq 0 \quad \forall n = 1 \dots N_h \quad 0-35$$

si bien que

$$(A^H A)^{-1} = \frac{1}{N_{FFT}} \text{Id}_L \otimes H^{-2} \quad 0-36$$

Finalement,

$$\begin{aligned} A^\# &= (A^H A)^{-1} A^H \\ &= \left(\frac{1}{N_{FFT}} \text{Id}_L \otimes H^{-2} \right) (\text{Id}_L \otimes HP^T W^H) \\ &= \frac{1}{N_{FFT}} \text{Id}_L \otimes (H^{-2} HP^T W^H) \\ A_{(LN_h, LN_{FFT})}^\# &= \frac{1}{N_{FFT}} \text{Id}_L \otimes H^{-1} P^T W^H \end{aligned} \quad 0-37$$

Expression de $R^\#$

A partir de la définition 0-29 de $R^\#$, de la définition 0-13 de R , et de la propriété 0-14, la matrice $R^T R$ de dimension $(N_0; N_0)$ a pour terme général

$$\begin{aligned} (R^T R)_{ij} &= \sum_{k=1}^{LN_h} R_{ik}^T R_{kj} \\ &= \sum_{k=1}^{LN_h} R_{ki} R_{kj} \\ &= \left(\sum_{k=1}^{LN_h} R_{ki}^2 \right) \delta[i-j] \\ &= \left(\sum_{k=1}^{LN_h} R_{ki} \right) \delta[i-j] \end{aligned} \quad \text{0-38}$$

Cette matrice est diagonale, chaque terme est la somme de la colonne de R correspondante, qui est par construction strictement positive. Donc,

$$\begin{aligned} \det(R^T R) &= \prod_{i=1}^{N_0} (R^T R)_{ii} \\ &= \prod_{j=1}^{N_0} \underbrace{\left(\sum_{i=1}^{LN_h} R_{ij} \right)}_{>0} \\ \det(R^T R) &> 0 \end{aligned} \quad \text{0-39}$$

Ainsi, la condition d'inversibilité de $R^T R$ est toujours vérifiée, et son inverse a pour terme général

$$(R^T R)_{ij}^{-1} = \frac{\delta[i-j]}{\sum_{i=1}^{LN_h} R_{ij}} \quad \text{0-40}$$

Sachant que R^T a pour terme général

$$R_{ij}^T = R_{ji} = \delta[j - I(n, l); i - J(n, l)] \quad , \quad \begin{aligned} \forall j &= 1 \dots LN_h \\ \forall i &= 1 \dots N_0 \end{aligned} \quad \text{0-41}$$

alors $R^\# = (R^T R)^{-1} R^T$ a pour terme général

0-42

$$\begin{aligned}
R_{ij}^{\#} &= \sum_{k=1}^{LN_h} \left(R^T R \right)_{ik}^{-1} R_{kj}^T \\
&= \sum_{k=1}^{LN_h} \left(\frac{\delta[i-k]}{\sum_{m=1}^{LN_h} R_{mk}} \right) \delta[j-I(n,l); k-J(n,l)] \\
&= \sum_{k=1}^{LN_h} \delta[i-k] \frac{1}{\sum_{m=1}^{LN_h} R_{mk}} \delta[j-I(n,l); k-J(n,l)] \\
R_{ij}^{\#} &= \frac{1}{\sum_{m=1}^{LN_h} R_{mi}} \delta[j-I(n,l); i-J(n,l)]
\end{aligned}$$

Par exemple, si $N = 11, N_h = 6, N_r = 2, N_0 = 14, L = 3$, alors d'après 0-12

0-43

$$\begin{cases} I(n,l) \in \{1; 2; 3; 4; 5; 6; 7; 8; 9; 10; 11; 12; 13; 14; 15; 16; 17; 18\} \\ J(n,l) \in \{1; 2; 3; 4; 5; 6; 5; 6; 7; 8; 9; 10; 9; 10; 11; 12; 13; 14\} \end{cases}$$

et d'autre part

0-44

$$\sum_{m=1}^{18} R_{mi} = \begin{cases} 1 & \text{si } i \in \{1; 2; 3; 4; 7; 8; 11; 12; 13; 14\} \\ 2 & \text{si } i \in \{5; 6; 9; 10\} \end{cases}$$

$R^{\#}$ est de dimension (14,18) et a pour expression :

$$R^\# = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & \frac{1}{2} & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \frac{1}{2} & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \frac{1}{2} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \frac{1}{2} & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix} \quad \text{0-45}$$

H. Expression détaillée

A partir des expressions définissant la synthèse harmonique fenêtrée du signal 0-31 et 0-32.

respectivement $\begin{cases} \hat{x} = A^\# X \\ x = R^\# \hat{x} \end{cases}$ et $x = R^\# (A^\# X)$, et des expressions 0-37 de $A^\#$ et 0-42 de $R^\#$,

l'opération se déroule en deux étapes. D'une part,

$$\hat{x} = \frac{1}{N_{FFT}} (\text{Id}_L \otimes H^{-1} P^T W^H) X \quad \text{0-46}$$

et d'autre part

$$\begin{aligned} x_i &= \frac{1}{\sum_{m=1}^{LN_h} R_{mi}} \sum_{j=1}^{LN_h} \delta[j - I(n, l); i - J(n, l)] \hat{x}_j \\ &= \frac{1}{\sum_{m=1}^{LN_h} R_{mi}} \delta[i - J(n, l)] \hat{x}_{I(n, l)} \\ x_i &= \frac{\hat{x}_{I(n, l)}}{\sum_{m=1}^{LN_h} R_{mJ(n, l)}} \delta[i - J(n, l)] \quad \forall i = 1 \dots N_0 \end{aligned} \quad \text{0-47}$$

Ainsi les N premiers échantillons de x sont les échantillons du signal original

Bibliographie

- [1] W. Burdic, Underwater acoustic system analysis, Prentice Hall, 1984.
- [2] W. C. Knight, R. G. Pridham, et S. M. Kay, «Digital signal processing for sonar,» Proceedings of the IEEE, vol. 69, no. 11, pp. 1451-1506, November 1981
- [3] T. Glisson, C. Black et A. Sage, «On sonar signal analysis,» *IEEE Transactions on AES*, 1970.
- [4] DCNS, «Rapport interne,» 2006.
- [5] M. Bouvet, Traitement des signaux pour les systèmes SONAR, MASSON, 1992, pp. 316-321.
- [6] L. Kopp, Traitement d'antenne, 2003.
- [7] C. Gasquet et P. Witmoski, Analyse de Fourier et ses applications, Université de Grenoble, 1996.
- [8] S. Marcos, Les methodes à haute résolution, Hermès, 1998.
- [9] B. Boashash, Time-Frequency Signal Analysis and Processing: A Comprehensive Reference, Elsevier, Oxford, 2003, Oxford, 2003.
- [10] D. Gabor, «Theory of communications,» *Journal of institute of Electrical Engineers IEE*, 1946.
- [11] Y. Grenier, «Thèse de Doctorat: Modélisation des signaux non-stationnaires,» 1984.
- [12] L. Cohen, Time-frequency analysis, Prentice Hall, Englewoods Cliffs, 1995, pp. 429-457.
- [13] J. Cooley, W. Louis, A. Peter, Welch et D. Peter, «An algorithm for the machine calculation of complex series,» *Math-computation*, pp. 297-301, 1967.
- [14] M. Slanley, «Auditory Toolbox, version 2,» 1998.
- [15] J. Ville, «Théorie et applications de la notion de signal analytique,» *Câbles et transmissions*, pp. 61-74, 1948.
- [16] E. Wigner, «On the quantum corection for thermodynamic equilibrium,» *Physiques Review*, pp. 749-759, 1932.
- [17] T. Claassen et W. Mecklenbauer, «The aliasing problem in discrete-time Wigner distributions,» *IEEE transactions on acoustics, speech and signal processing*, pp. 1067-1072, 1984.
- [18] P. Flandrin et B. Escudié, «Principe et mise en oeuvre de l'analyse temps-fréquence par transformation de Wigner-Ville,» *Traitement du signal*, 1985.

- [19] J. Cexus, «Thèse de Doctorat: Analyse des signaux non-stationnaires par transformation de Huang, opérateur de Teager-Kaiser, et transformation de Huang-Teager,» 2005.
- [20] B. Boashash et P. Flandrin, «Wigner Ville analysis of time-varying signals,» *Proceedings ICASSP*, pp. 1329-1331, 1982.
- [21] F. Hlawatsch, «Interferences terms in the Wigner distribution,» *Digital signal processing*, pp. 363-367, 1984.
- [22] T. Claasen et W. Mecklenbauer, «The Wigner distribution- a tool time-frequency signal analysis,» *Philips*, pp. 217-250, 1980.
- [23] K. Kodera, C. D. Villedary et R. Gendrin, «A new method for the numerical analysis of nonstationary signals,» *Phys. Earth and Plan*, pp. 142-150, 2005.
- [24] F. Auger et P. Flandrin, «Improving the readability of time frequency and time scale representations by reassignment method,» *IEEE transactions on signal processing*, pp. 1068-1089, 1995.
- [25] E. Chassande-Mottin, «Thèse de Doctorat, Méthodes de réallocation dans le paln temps-fréquence pour l'analyse et le traitement des signaux non stationnaires,» 1998.
- [26] E. Chassande-Mottin, F. Auger et P. Flandrin, «Temps-fréquence: concepts et outils (chapitre 9),» 2003.
- [27] A. Grossman et J. Morlet, «Decomposition of hardy functions into square integrable wavelets of constant shape,» *SIAM*, pp. 723-726, 1984.
- [28] S. Mallat, «A theory for multiresolution signal decomposition: the wavelets representation,» *IEEE Transactions on pattern analysis and machine intellignece*, pp. 674-693, 1989.
- [29] M. Holdschneider, R. Kronland-Martinet, J. Morlet et P. Tchamitchian, «A real time algorithm for the signal analysis with the help of wavelet transform,» *Inverse problem and theoretical imaging*, pp. 286-297, 1990.
- [30] R. Mill et G. Brown, «Auditory-based time-frequency detection of chirps and feature extraction techniques for Sonar processing,» *Speech and Hearing Reasearch Group*, 2005.
- [31] X. Vuylsteke, «Cours d'acoustique et de mécanique ondulatoire,» 2012.
- [32] F. Zheng, G. Zhang et Z. Song, «Comparison of different implementations of MFCC,» *Journal of computer, science and technology*, pp. 582-589, 2001.
- [33] S. Davis et P. Mermelstein, «Comparison of pramaetric representations of monsyllabic word recognition in continuously spoken sentences,» *IEEE Transaction on Acoustics, Speech and signal processing*, pp. 357-366, 1980.

- [34] R. Hodges, «Underwater acoustics: analysis, design and performance of sonar,» *Wiley and sons*, 2011.
- [35] Y. Ephraim et D. Malah, «Speech enhancement using a minimum mean square error short-time spectral amplitude estimator,» *IEEE. Trans. Acoust., Speech, Signal Process*, pp. 1109-1121, 1984.
- [36] C.Plapous, C.Marro et P.Scalart, «Improved Signal-to-Noise Ratio Estimation for Speech Enhancement,» *IEEE Trans. Speech, Audio Process*, 2006.
- [37] X. Zhang, H. Jiang et J. Zhang, «Improved Priori SNR Estimation for Sound Enhancement with Gaussian Statistical Model,» *IEEE*, 2012.
- [38] I. Cohen, « Speech enhancement using a noncausal a priori SNR estimator,» *IEEE Signal Process. Lett.*, 2004.
- [39] E.Kalman, «A New Approach to Linear Filtering and Prediction Problems,» *Transactions of the ASME - Journal of Basic Engineering*, pp. 35-45, 1960.
- [40] Y. Malah et D. Ephraim, «Speech enhancement using a minimum mean square error log-spectral amplitude estimator,» *IEEE Trans. Acoust., Speech, Signal Process*, p. 443–445, 1985.
- [41] P. Wove et S. Godsill, «Simple alternatives to the Ephraim and Malah suppression rule for speech enhancement,» *IEEE*, 2001.
- [42] T. Lotter et P.Vary, «Speech enhancement using a super gaussian speech model,» *EURASIP journal on applied signal processing*, pp. 1110-1126, 2005.
- [43] D. Ruppert, *Statistics and data analysis for financial engineering*, Springer, 2011.
- [44] D. Donoho et J. Johnstone, «Ideal spatial adaptation by wavelet shrinkage,» *Biometrika*, pp. 422-455, 2005.
- [45] S. Quackenbush, T. Barnwell et M. Clements, «Objectives measures of speech quality,» *Prentice-Hall*, 1988.
- [46] X. Lurton, *Acoustique sous-marine: présentation et applications*, 2001, pp. 27-29.
- [47] S. Tollari, *Indexation et recherche d'images par fusion d'informations textuelles et visuelles*, 2006.
- [48] J. MacQueen, «Some methods for classification and analysis of multivariate observations,» *Proceedings of the Fifth symposium on mathematical statistics and probability*, pp. 281-297, 1967.
- [49] U. Luxburg, «A tutorial on spectral clustering,» *Statistics and computing*, pp. 395-416, 2007.

Bibliographie

- [50] M. Ester, H. Kriegel, J. Sander et X. Xiu, «A density-based algorithm for discovering clusters in large spatial,» *AAAI*, 1996.
- [51] J. Bezdek, «Pattern recognition with fuzzy objective function algorithms,» *Plenum Press*, 1981.
- [52] A. Bensaid, «Validity guided with applications to image clustering,» *IEEE transactions on fuzzy systems*, pp. 112-113, 1996.
- [53] X.L.Xie et G. Beni, «Validity measure for fuzzy clustering,» *IEEE transactions on PAMI*, pp. 841-846, 1991.
- [54] K.Nishiguchi, «Time-period analysis for pulse train deinterleaving,» *Computers of the society of instrument and control engineers*, pp. 68-78, 2005.
- [55] O. L. Bot, C. gervaise, J.-I. Mars et J. Bonnel, «Séparation d'impulsions bio-acoustique par analyse du rythme,» *Gretsi*, 2013.
- [56] P. G. Ciarlet, *Introduction à l'Analyse Numérique Matricielle et à l'Optimisation*, Paris: Masson, 1982.
- [57] J. Suykens et J. Vandewalle, «Multiclass least squares support vector machines,» *World Scientific*, 1999.
- [58] W. Karush, Master's thesis: Minima of functions of several variables with inequalities as side constraints, Chicago, 1939.
- [59] H. Kuhn et A. Tucker, «Nonlinear programming,» *Mathematical Statistics and probabilistics*, pp. 481-492, 1951.
- [60] J. Burges, «A tutorial on support vector machines for pattern recognition,» *Journal of data mining and knowledge discovery*, pp. 1-43, 1998.
- [61] B. Sholkopf et A. Smola, «Learning with kernels support vector machines, regularization, optimization, and beyond,» *The MIT Press*, 2002.
- [62] J. Mercer, «Functions of positive and negative type, and their connection with the theory of integral equations,» *Philosophical transactions of the royal society of London. Series A, containing papers of a mathematical or physical character*, n° 1209, pp. 415-446, 1909.
- [63] K. Crammer et Y. Singer, «On the algorithmic implementation of multiclass kernel-based vector machines,» *Journal of machine learning research*, pp. 143-160, 2002.
- [64] G. Fung et O. Mangasarian, «Multicategory proximal support vector machine classifiers,» *Machine learning*, pp. 77-97, 2005.
- [65] B. Scholkopf, C. Burges et V. Vapnik, «Extracting support data for given task,» *KDD'95*, pp. 252-

257, 1995.

- [66] V. Vapnik, «The nature of statistical learning theory,» *Springer-Verlag*, 1995.
- [67] S. Knerr, L. Personnaz et G. Dreyfus, «Single-layer learning revisited: a stepwise procedure for building and training a neural network,» *Neurocomputing: algorithms, architectures and applications*, pp. 41-50, 1990.
- [68] J. Friedman, «Another approach to polychotomus classification,» *Technical report, departement of statistics*, 1996.
- [69] J. Platt, N. Cristianini et J. Shawe-Taylor, «Large margine DAGs for multiclass classification,» *NIPS*, pp. 547-553, 2012.
- [70] K. Benabdeslem et Y. Bennani, «Dendrogram-based SVM for multi-class classification,» *Journal of computing and information technology-CIT*, pp. 283-289, 2006.
- [71] J. Weston et C. Watkins, «Support vector machines for multi-class pattern recognition,» *ESANN*, pp. 77-78, 1999.
- [72] C. Watkins et J. Weston, «Multi-class support machines,» *Technical report, Department of computer science, royal holloway, university of London*, 1998.
- [73] Y. Singer et K. Crammer, «On the learnability and design output codes for multiclass problems,» *Machine learning*, pp. 201-233, 2002.
- [74] Y. Lee, «Multicategory support vector machines, theory, and application to the classification of microarray data and satellite radiance data,» *Technical report 1063, university of Wisconsin*, 2002.
- [75] C. Hsu et C.-J. Lin, «A comparison of methods for multiclass support vector machines,» *IEEE transactions on neural network*, pp. 69-78, 2002.
- [76] S. Theodoridis et K. Koutroumbas, «Pattern recognition,» *Academic Press*, p. 86, 2008.
- [77] S. Aksoy et M. Haralick, «Feature normalization and likelihood-based similarity measures for image retrieval,» *Pattern recognition letters*, p. 87, 2001.
- [78] B. Kedem, «Spectral analysis and discrimination by zero-crossings,» *Proceedings IEEE*, p. 88, 1986.
- [79] S. F. Chang, T. Sikora et A. Puri, «Overview of the mpeg-7 standard.,» *IEEE transactions on circuits and systems*, pp. 688-695, 2001.
- [80] E. Scheirer et M. Slaney, «Construction and evaluation of a robust multifeature speech/music discrimination,» *IEEE International conference on acoustics, speech and signal processing*, pp.

1331-1334, 1997.

- [81] C.D'Alessandro, *Analyse synthèse et codage de la parole*, Hermes, Lavoisier, 2002.
- [82] G. Peeters, «A large set of audio features for sound description (similarity and classification) in the cuidado project,» *IRCAM*, 2004.
- [83] R. Belleman, «Adaptative control processes: a guided tour,» *Princeton University Press*, 1961.
- [84] I. Guyon et A. Elisseeff, «An introduction to variable and feature selection,» *Journal of machine learning research*, pp. 93-94, 2003.
- [85] G. Toussain, «Note on optimal selection of independant binary-valued features of pattern recognition,» *IEEE transaction on information theory*, p. 93, 1971.
- [86] G. John, R. Kohavi et K. Pfleger, «Irrelevant features and the subset slection problem,» *machine learning*, pp. 121-129, 1994.
- [87] J. Huang, D. Yang et Y. Chuang, «Application of wrapper approach and composite classifier to the stock trend prediction,» *Expert system applocation*, pp. 2870-2878, 2008.
- [88] A. Whitney, «A direct method of nonparametric measurement selection,» *IEEE transaction on computation*, 1971.
- [89] P. Pudil, J. Novovicova et J. Kittler, «Floating search methods, in feature selection,» *Pattern recognitions letters*, 1994.
- [90] A. Jain et D. Zongker, «Feature selection: Evaluation, application and small performance,» *IEEE transacations on PAMI*, pp. 153-158, 1997.
- [91] Y. Li et L. Guo, «Tcm-knn scheme for network anomaly detection using feature based optimizations,» *Proceedings of the 2008 ACM symposium on applied computing*, pp. 2103-2109, 2008.
- [92] R. Duda, P. Hart et D. Stock, «Pattern classification,» *Wiley-interscience*, 2000.
- [93] T. Furey, N. Cristianini, N. Duffy, D. Bednarski, M. Schummer et D. Haussler, «Support vector machine classification and validation of cancer tissue samples using microarray expression data,» *Bioinformatics*, pp. 906-914, 2000.
- [94] P. Hanchuan, L. Fuhui et C. Ding, «Feature selection based on mutual information: criteria of max dependancy, max relevance, and min-redundancy,» *IEEE on pattern analysis and machine intelligence*, pp. 1126-1238, 2005.
- [95] S. Kullback et A. Liebler, «On information and sufficiency,» *Annals of mathematical statistics*, pp. 79-86, 1951.

- [96] R. Moddemeijer, «On estimation of entropy and mutual information of continuous distributions,» *Signal processing*, pp. 233-246, 1989.
- [97] N. Vasconcelos, «Feature selection by maximum maximal diversity: optimality and implications for visual recognition,» *Proceeding of IEEE international conference on image processing (ICIP)*, pp. 762-769, 2003.
- [98] UCI, «<http://archive.ics.uci.edu/ml/>,» [En ligne].
- [99] S. Thrun, «A performance comparison of different learning algorithms,» *Technical report, Carnegie Mellon university*, 1991.
- [100] I. Rish, «An empirical study of the naive Bayes classifier,» *IJCAI Workshop on Empirical Methods in Artificial Intelligence*, 2001.
- [101] K. Hoff, M. Tech, T. Lingner, R. Daniel, B. Morgenstern et P. Meinicke., «Gene prediction in metagenomic fragments: a large scale machine learning approach,» *BMC bioinformatics*, 2008.
- [102] J. Platt, «Sequential minimal optimization: a fast algorithm for training support vector machines,» *Advances in Kernel Methods-Support Vector Learning*, pp. 98-112, 1999.
- [103] P. Courmontagne, S. Ouelha, U. Moreaud et F. Chaillan, «A blind denoising process with applications to underwater acoustic signals,» chez *IEEE Oceans*, San Diego, 2013.
- [104] P. Courmontagne, S. Ouelha et F. Chaillan, «A new time-frequency representation for underwater acoustic signals: The denoised hearingogram,» chez *MTS/IEEE OCEANS*, Bergen, 2013.
- [105] P. Courmontagne, S. Ouelha et F. Chaillan, «On time-frequency representations for underwater acoustic signal,» chez *Oceans*, Hampton Roads, 2012.
- [106] S. Ouelha, J.-R. Mesquida, F. Chaillan et P. Courmontagne, «Extension of maximal marginal diversity based feature selection applied to underwater acoustic data,» chez *IEEE/MTS Oceans*, San Diego, 2013.
- [107] P. Flandrin, «Some features of time-frequency representations of multicomponents signals,» *IEEE on acoustics, speech and signal processing, ICASSP*, 1984.
- [108] P. Flandrin, W. Martin et M. Zakharia, «On a hardware implementation of the Wigner Ville transform,» *Digital signal processing*, pp. 262-266, 1984.

Représentation et reconnaissance des signaux acoustiques sous-marins

Résumé en français

Cette thèse a pour but de définir et concevoir de nouvelles techniques de représentation des signaux acoustiques sous-marins. Notre objectif est d'interpréter, reconnaître et identifier de façon automatique les signaux sous-marins émanant du système sonar. L'idée ici n'est pas de substituer la machine à l'officier marinier, dont l'expérience et la finesse d'ouïe le rendent indispensable à ce poste, mais d'automatiser certains traitements de l'information pour soulager l'analyste et lui offrir une aide à la décision.

Dans cette thèse, nous nous inspirons de ce qui se fait de mieux dans ce domaine : l'humain. A bord d'un sous-marin, ce sont des experts de l'analyse des sons à qui l'on confie la tâche d'écoute des signaux afin de repérer les sons suspects. Ce qui nous intéresse, c'est cette capacité de l'humain à déterminer la classe d'un signal sonore sur la base de son acuité auditive. En effet, l'oreille humaine a le pouvoir de différencier deux sons distincts à travers des critères perceptuels psycho-acoustiques tels que le timbre, la hauteur, l'intensité. L'opérateur est également aidé par des représentations du signal sonore dans le plan temps-fréquence qui viennent s'afficher sur son poste de travail. Ainsi nous avons conçu une représentation qui se rapproche de la physiologie de l'oreille humaine, autrement dit de la façon dont l'homme entend et perçoit les fréquences. Pour construire cet espace de représentation, nous utiliserons un algorithme que nous avons appelé l'*Hearingogram* et sa version débruitée le *Denoised Hearingogram*. Toutes ces représentations seront en entrée d'un système d'identification automatique, qui a été conçu durant cette thèse et qui est basé sur l'utilisation des SVM.

Mot clés : Représentation temps-fréquence, Reconnaissance, Descripteurs.

Representation and recognition of underwater acoustic signals

English abstract

This thesis aims to identify and develop new representation methods of the underwater acoustic signals. Our goal is to interpret, recognize and automatically identify underwater signals from sonar system. The idea here is not to replace the machine petty officer, whose experience and hearing finesse make it indispensable for this position, but to automate certain processing information to relieve the analyst and offer support to the decision.

In this thesis, we are inspired by what is best in this area: the human. On board a submarine, they are experts in the analysis of sounds that are entrusted to the listening task signals to identify suspicious sounds. What interests us is the ability of the human to determine the class of a sound signal on the basis of his hearing. Indeed, the human ear has the power to differentiate two distinct sounds through psychoacoustic perceptual criteria such as tone, pitch, intensity. The operator is also helped by representations of the sound signal in the time-frequency plane coming displayed on the workstation. So we designed a representation that approximates the physiology of the human ear, i.e how humans hear and perceive frequencies. To construct this representation space, we will use an algorithm that we called the *Hearingogram* and a denoised version the *Denoised Hearingogram*. All these representations will input an automatic identification system, which was designed during this thesis and is based on the use of SVM.

Keywords : Time-frequency representation, Identification, features.