

ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
UNIVERSITÉ DU QUÉBEC

MÉMOIRE PRÉSENTÉ À
L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

COMME EXIGENCE PARTIELLE
À L'OBTENTION DE LA
MAÎTRISE EN GÉNIE DE L'ENVIRONNEMENT
M. Sc. A.

PAR
Isabelle PRÉVOST

APPLICATION DE LA VISION ARTIFICIELLE À L'IDENTIFICATION DE GROUPES
BENTHIQUES DANS UNE OPTIQUE DE SUIVI ENVIRONNEMENTAL DES RÉCIFS
CORALLIENS

MONTRÉAL, LE 26 OCTOBRE 2015



Isabelle Prévost, 2015



Cette licence [Creative Commons](https://creativecommons.org/licenses/by-nc-nd/4.0/) signifie qu'il est permis de diffuser, d'imprimer ou de sauvegarder sur un autre support une partie ou la totalité de cette œuvre à condition de mentionner l'auteur, que ces utilisations soient faites à des fins non commerciales et que le contenu de l'œuvre n'ait pas été modifié.

PRÉSENTATION DU JURY
CE MÉMOIRE A ÉTÉ ÉVALUÉ
PAR UN JURY COMPOSÉ DE :

M. Jacques-André Landry, directeur de mémoire
Département de génie de la production automatisée à l'École de technologie supérieure

M. Robert Hausler, président du jury
Département de génie de la construction à l'École de technologie supérieure

M. Richard Lepage, membre du jury
Département de génie de la production automatisée à l'École de technologie supérieure

IL A FAIT L'OBJET D'UNE SOUTENANCE DEVANT JURY ET PUBLIC

LE 19 OCTOBRE 2015

À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

REMERCIEMENTS

Je tiens à remercier mon directeur de recherche, Jacques-André Landry, pour ses conseils, son support moral et ses apports financiers au projet. Il a su guider mon travail et son expertise m'a été précieuse. Je n'aurais également pas pu réaliser ce mémoire sans l'*Australian Institute of Marine Science*, qui a fourni la base de données, et le laboratoire du LIVIA dans lequel j'ai réalisé mes recherches.

Je remercie également Jean-Nicola Blanchet pour les idées qu'il a amenées, ainsi que pour les longues discussions au sujet de la problématique. Ces dernières m'ont grandement aidée à développer ma méthodologie. Je veux aussi exprimer ma gratitude à Moslem Ouled Sghaier pour ses critiques constructives et à Christopher Pagano pour son aide avec les serveurs du laboratoire du LIVIA.

Je remercie finalement ma famille de m'avoir épaulée au cours de mes études et d'avoir aidé à corriger mon mémoire. Vous, mes parents Louis et Mireille, mon frère Philippe Édouard, mon fiancé Patrick, puis Sylvie et Christine, avez été essentiels à la réalisation de ce travail.

APPLICATION DE LA VISION ARTIFICIELLE À L'IDENTIFICATION DE GROUPES BENTHIQUES DANS UNE OPTIQUE DE SUIVI ENVIRONNEMENTAL DES RÉCIFS CORALLIENS

Isabelle PRÉVOST

RÉSUMÉ

La composition des récifs coralliens est un excellent indicateur de la santé de la faune marine. Pour cette raison, les biologistes de l'*Australian Institute of Marine Science* (AIMS) en effectuent un suivi constant par l'analyse de photos acquises chaque année à travers la Grande Barrière de corail. Pour accélérer l'identification de leur contenu, on développe des algorithmes automatisés de reconnaissance de forme qui sont basés sur l'intelligence et la vision artificielles.

Nous avons optimisé chaque étape d'un de ces algorithmes pour les caractéristiques de la base de données de l'AIMS. Pour débiter, nous avons évalué divers prétraitements pour compenser l'effet de l'imagerie sous-marine. Ensuite, nous avons itéré sur la taille de la fenêtre d'analyse pour former une segmentation simple et facile d'application. Puis, nous avons extrait des descripteurs à plusieurs échelles et sur plusieurs canaux de couleur de manière à exploiter adéquatement la richesse en information visuelle des images de coraux et nous avons réduit la dimensionnalité de l'espace de descripteurs. Enfin, nous avons défini une plage de valeurs idéales pour les paramètres des classificateurs. Pour compléter le tout, nous avons comparé la performance des étapes optimisées à celle d'algorithmes correspondant à la fine pointe de la technologie.

Postérieurement, nous avons généralisé l'application à toute la base de données par des validations croisées qui ont permis de définir les limites de performance du système. Nous avons aussi développé des stratégies réalistes d'exploitation de la base de données. Pour ce faire, nous avons évalué la compatibilité de divers groupes d'entraînement et de test, appartenant à la même période temporelle, mais à des emplacements spatiaux distincts et vice-versa, puis à des groupes petits, grands, homogènes et diversifiés. La stratégie ainsi développée a permis d'atteindre les objectifs fixés et de compléter un outil efficient pour les biologistes marins de l'AIMS.

Mots clés : récifs coralliens, groupes benthiques, intelligence artificielle, vision artificielle, reconnaissance de formes

APPLICATION OF ARTIFICIAL VISION TO BENTHIC GROUPS IDENTIFICATION TOWARDS THE ENVIRONMENTAL MONITORING OF CORAL REEFS

Isabelle PRÉVOST

ABSTRACT

The content of coral reefs is an excellent indicator of the health of the sea fauna. Because of this, biologists at the Australian Institute of Marine Science (AIMS) constantly monitor it through the analysis of pictures, which are acquired every year in the Great Barrier Reef. To accelerate the identification of their content, automated coral recognition algorithms that are based on artificial intelligence and artificial vision have been developed.

We optimized each step of a pattern recognition algorithm for the AIMS' database's particularities. To begin with, we tested the pre-processing techniques in order to compensate the effect of underwater imaging. Afterwards, we iterated on the size of the region of interest to create a simple and easy to apply segmentation method. Then, we extracted features through different scales and color channels to use as much of the available visual information as possible and we reduced the dimensionality of the feature space. Finally, we defined a range of ideal values for the classifier's parameters. We completed this series of step by comparing the performance of the optimized pattern recognition algorithm to that of the state of the art.

Subsequently, we generalized the application of this algorithm to the entire database by the use of cross-validations. The latter helped define the performance's limits of the system. We also developed operational strategies for AIMS' database. To do this, we assessed the compatibility of various training and testing groups belonging to the same time period, but to different locations and vice versa, then to small, large, homogeneous and diverse groups. The resulting strategy achieved the set objectives and yielded an efficient tool for the marine biologists of the AIMS.

Keywords: coral reefs, benthic groups, artificial intelligence, artificial vision, pattern recognition

TABLE DES MATIÈRES

	Page
INTRODUCTION	1
CHAPITRE 1 REVUE DE LITTÉRATURE	5
1.1 Les récifs coralliens	5
1.2 Collecte des données sous-marines.....	6
1.3 Reconnaissance de forme.....	9
1.3.1 Bases de données	14
1.3.2 Débalancement des classes	16
1.3.3 Prétraitements.....	18
1.3.4 Segmentation.....	23
1.3.5 Extraction des descripteurs de texture	24
1.3.5.1 Espaces de représentation	25
1.3.5.2 Analyse statistique de la distribution des pixels	27
1.3.5.3 Analyse de la distribution des pixels par filtrage.....	31
1.3.6 Extraction des descripteurs de couleur.....	32
1.3.7 Dimension de l'espace de représentation.....	34
1.3.7.1 Analyse en composantes principales (ACP)	34
1.3.7.2 Sélection d'attributs	35
1.3.8 Classificateurs	36
1.3.9 Analyse statistique des résultats.....	40
1.3.10 Application à la classification de groupes benthiques	42
CHAPITRE 2 OPTIMISATION D'UN ALGORITHME DE CLASSIFICATION POUR LA BASE DE DONNÉES DE L'AIMS	45
2.1 Introduction.....	45
2.2 Méthodologie d'optimisation.....	45
2.3 Sélection des prétraitements	47
2.3.1 Caractéristiques des images	47
2.3.2 Méthode de comparaison des prétraitements	47
2.3.3 Résultats de la comparaison des prétraitements.....	48
2.3.4 Interprétation des conséquences du prétraitement	49
2.4 Segmentation.....	50
2.4.1 Choix du type de segmentation	50
2.4.2 Méthode de sélection de la taille de l'imagette.....	50
2.4.3 Résultats de la sélection de la taille de l'imagette.....	50
2.4.4 Impact du choix de taille d'imagettes	51
2.5 Sélection des descripteurs	52
2.5.1 Définition de l'impact du choix des descripteurs.....	52
2.5.2 Méthodologie de sélection des descripteurs.....	52
2.5.3 Distribution des descripteurs sélectionnés à travers les récifs	54
2.5.4 Apport de chaque descripteur à l'ensemble	57

2.5.5	Comparaison de la performance des descripteurs à la fine pointe de la technique	59
2.6	Optimisation des paramètres des SVM.....	60
2.6.1	Utilisation d'un SVM.....	60
2.6.2	Méthodologie d'optimisation pour la base de données de l'AIMS	61
2.6.3	Résultats de l'optimisation des paramètres sur plusieurs récifs.....	62
2.6.4	Interprétation des valeurs prises par les paramètres.....	64
2.6.5	Discussion sur le processus d'optimisation	64
2.7	Configuration finale	65
CHAPITRE 3	GÉNÉRALISATION SUR LA BASE DE DONNÉES DE L'AIMS	67
3.1	Introduction.....	67
3.2	Seuils de performance du système	67
3.2.1	Méthodologie d'évaluation des limites du processus de classification.....	68
3.2.2	Résultats des validations croisées	68
3.2.2.1	Interprétation des limites de performance du système de classification	73
3.3	Scénarios de classification des images de l'AIMS	75
3.3.1	Méthodologie d'évaluation des couples entraînement-test.....	76
3.3.2	Résultats d'entraînements et tests sur deux récifs distincts, au cours d'une période d'échantillonnage.....	77
3.3.3	Résultats d'entraînements sur plusieurs récifs et de test sur un récif distinct, au cours d'une période.....	79
3.3.4	Résultats d'entraînements sur plusieurs périodes d'échantillonnage pour un seul récif.....	80
3.3.5	Résultats d'entraînements sur d'autres années de tous les récifs.....	82
3.3.6	Interprétation des résultats des différents scénarios.....	85
CONCLUSION		89
RECOMMANDATIONS		93
ANNEXE I	STATISTIQUES APPLIQUÉES À LA MATRICE DE COOCCURRENCE	97
ANNEXE II	STATISTIQUES APPLIQUÉES AUX HISTOGRAMMES DE TONS DE GRIS.....	101
ANNEXE III	BASE DE DONNÉES DE L'AIMS	103
ANNEXE IV	RÉSULTATS DE L'OPTIMISATION DES PARAMÈTRES DU NOYAU RBF DU SVM.....	107
ANNEXE V	RÉSULTATS DES VALIDATIONS CROISÉES	113

ANNEXE VI	RÉSULTATS DE LA GÉNÉRALISATION DE L'APPLICATION DE L'ALGORITHME	117
LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES.....		130

LISTE DES TABLEAUX

	Page
Tableau 1.1	Combinaisons optimales de techniques pour chaque étape de la classification selon (Shihavuddin et al., 2013)43
Tableau 2.1	Paramètres de chaque étape de la classification pour le scénario de base46
Tableau 2.2	Ensemble des descripteurs de texture et de couleur considérés.....53
Tableau 2.4	Exemples de différence entre le taux de classification obtenu sur chaque récif selon que les descripteurs aient été sélectionnés sur eux-mêmes ou sur Martin 2012-201356
Tableau 2.5	Taux de classification du système en fonction du nombre de valeurs utilisées et de leur méthode de sélection.....60
Tableau 2.6	Plage des valeurs des exposants pour $\gamma = 2^x$ et $C = 2^y$ lors de la recherche en grille grossière.....61
Tableau 2.7	Plage de valeurs des exposants pour $\gamma = 2^x$ et $C = 2^y$ lors de la recherche en grille fine centrée sur $\gamma = 2^{x_0}$ et $C_0 = 2^{y_0}$62
Tableau 2.8	Configuration finale du système de classification65
Tableau 3.1	Taux de rappel, précision et population des groupes benthiques, moyennés sur toutes les validations croisées pour chaque couple de récif et de période, à l'exception de ceux éliminés71
Tableau 3.2	Taux de classification, précision et population des groupes benthiques fusionnés75
Tableau 3.3	Comparaison des écarts, avec l'étalon, des taux de classification pour des couples entraînement-test intervertis.....78
Tableau 3.4	Taux de classification (TC) de tous les récifs d'une période lors de l'entraînement sur tous les récifs de toutes les autres périodes83
Tableau 3.5	Taux de classification, précision et population des groupes benthiques fusionnés pour un entraînement sur les périodes 2006-2007, 2008-2009 et 2010-2011 et un test sur 2012-2013, sans les récifs No Name, Lizard et Carter83

LISTE DES FIGURES

	Page
Figure 1.1	Niveaux taxonomiques étiquetés7
Figure 1.2	Décomposition d'une image (O) dans les canaux gris (G), puis rouge (R), vert (V) et bleu (B)10
Figure 1.3	A) Représentation de pois verts et B) de barres rouges10
Figure 1.4	Processus de classification d'une image de récif corallien11
Figure 1.5	Représentation d'une matrice de confusion12
Figure 1.6	Nombre d'échantillons dans les récifs en fonction de la période évaluée .15
Figure 1.7	Nombre d'échantillons de chaque groupe benthique selon les récifs16
Figure 1.8	Illustration des transformations d'un histogramme A) original, B) avec seuillage, C) avec allongement D) avec égalisation19
Figure 1.9	Traitement d'une image par ICM-R23
Figure 1.10	Variation de la texture à travers les échelles26
Figure 1.11	Illustration d'une transformée de Fourier en 1D et en 2D27
Figure 1.12	Zone de l'intégration pour une plage de fréquence.28
Figure 1.13	Exemple de génération de matrices de cooccurrence29
Figure 1.14	Représentation d'un A) pixel et de son voisinage pour un rayon de 2 pixels, de la B) différence de l'intensité relativement au pixel central, du C) LBP équivalent, exprimé par le signe de la différence et de D) la magnitude de la différence pour compléter CLBP30
Figure 1.15	Filtrage d'une image par un filtre de Gabor31
Figure 1.16	Algorithme des k plus proches voisins37
Figure 1.17	Représentation d'un réseau de neurones38
Figure 1.18	Représentation de A) l'entraînement d'un SVM et de B) la classification d'un élément à partir de ce SVM38
Figure 1.19	Représentation d'une boîte à moustache41

Figure 2.1	Application de la balance des blancs (BB), de CLAHS, d'un filtre médian (M), de ICM-R et de CCN à une image originale (O)48
Figure 2.2	Taux de classification associés à divers prétraitements49
Figure 2.3	Taux de classification associé à la taille des imageries testées51
Figure 2.4	Taux de classification associé à chaque combinaison testée d'imageries de tailles différentes51
Figure 2.5	Fréquence relative d'occurrence de chaque technique d'extraction dans le groupe de descripteurs sélectionnés.....55
Figure 2.6	Fréquence relative d'occurrence de chaque échelle des pyramides gaussiennes et laplaciennes dans les descripteurs sélectionnés55
Figure 2.7	Fréquence relative d'occurrence de chaque dimension de couleur dans les attributs sélectionnés56
Figure 2.8	Gain moyen d'information associé à chaque modalité d'extraction d'attributs58
Figure 2.9	Différence de performance entre l'application des descripteurs sélectionnés et certains de leurs sous-ensembles.59
Figure 2.10	Écarts entre les taux de classifications lors de l'optimisation des paramètres γ et C sur le groupe d'entraînement et lorsque nous conservons ces paramètres à travers les groupes d'entraînement63
Figure 2.11	Écarts entre les taux de classifications des cas d'optimisation des paramètres sur les deux récifs d'entraînement et sur un seul des récifs d'entraînement64
Figure 3.1	Matrice de confusion moyenne pour les validations croisées menées sur chaque récif et chaque période69
Figure 3.2	(A) Distribution des taux de classification à travers les années pour les validations croisées sur chaque couple de récif et de période et (B) distribution des écart-types correspondants liés à la robustesse du système.....69
Figure 3.3	Taux de classification selon les récifs pour les validations croisées sur chaque couple de récif et de période70
Figure 3.4	Matrice de confusion obtenue pour la validation croisée touchant le récif Linnet pour la période 2012-2013.....71

Figure 3.5	Matrice de confusion moyenne pour les validations croisées menées sur chaque récif et chaque période, à l'exception des récifs éliminés72
Figure 3.6	Comparaison des taux de classification de validations croisées (VC) effectuées sur tous les récifs d'une période aux moyennes des résultats des validations croisées des récifs de cette période pris isolément72
Figure 3.7	Comparaison des taux de classification de validations croisées (VC) effectuées sur toutes les périodes d'un récif aux moyennes des résultats des validations croisées de chaque période de ce récif prise isolément73
Figure 3.8	Matrice de confusion moyenne pour les groupes benthiques fusionnés75
Figure 3.9	Écart du taux de classification par rapport à l'étalon lors de l'entraînement sur un seul autre récif de la même année pour dix cas77
Figure 3.10	Écart entre le taux de classification de chaque classe benthique et l'étalon lors d'un entraînement sur un seul autre récif77
Figure 3.11	Corrélation entre C_{comp} et l'écart à l'étalon lors d'un entraînement sur un autre récif de la même année que le récif test78
Figure 3.12	Écart entre l'étalon et le taux de classification d'un entraînement sur un, deux, trois ou sept récifs d'une même année que le récif test79
Figure 3.13	Écart entre le taux de classification et l'étalon lors de l'entraînement sur les sept autres récifs de la même année79
Figure 3.15	Variation par rapport à l'étalon du taux de rappel de chaque groupe benthique lors de l'utilisation d'une autre année d'un même récif comme groupe d'entraînement sur 36 cas81
Figure 3.16	Corrélation entre C_{comp} et l'écart entre le taux de classification d'un récif test par un entraînement sur le même récif à une autre année81
Figure 3.17	Écarts entre les taux de classification et l'étalon selon que l'entraînement ait été fait avec le même récif dont l'échantillonnage s'est fait sur une, deux ou trois périodes dans le temps sur 36 cas82
Figure 3.18	Évolution du taux de classification d'un récif selon le nombre de périodes utilisées lors de l'entraînement82
Figure 3.19	Écart à l'étalon des taux de classification de périodes de test lors d'entraînements sur tous les récifs (1) de toutes les périodes

précédentes, (2) de toutes les périodes suivantes et (3) de toutes les périodes autres84

Figure 3.21 Matrice de confusion consolidée résultante de la classification de la période 2012-2013 à partir d'un SVM entraîné sur les périodes 2006-2007, 2008-2009 et 2010-2011, sans les récifs No Name, Lizard et Carter85

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

ACP	Analyse en composantes principales ou <i>Principal Component Analysis</i>
AIMS	<i>Australian Institute of Marine Science</i>
CCN	Normalisation de couleur complète ou <i>Comprehensive Color Normalization</i>
CLAHS	Spécification d’histogramme adaptative limitée en contraste ou <i>Contrast Limited Adaptive Histogram Specification</i>
CLBP	Motif linéaire binaire complet ou <i>Complete Linear Binary Pattern</i>
C-SVC	<i>C-Support Vector Classification</i>
CURTeT	Base de données de Columbia-Utrecht
ICM-R	Intégration de couleurs à partir de la distribution de Rayleigh
KNN	k plus proches voisins ou <i>k-nearest neighbors</i>
LBP	Motifs linéaires binaires ou <i>Linear Binary Pattern</i>
MCR-LTER	<i>Moorea Coral Reef-Long Term Ecological Research</i>
MLC	<i>Moorea-labeled corals</i>
MR8	Algorithme des huit réponses maximales ou <i>Maximum Reponse 8</i>
RBF	Fonction en base radiale ou <i>Radial Basis Function</i>

RVB	Rouge, vert, bleu ou <i>Red, Green, Blue</i>
SVM	Machine à vecteurs de support, Séparateur à vaste marge ou <i>Support Vector Machine</i>
TC	Taux de classification
VC	Validation croisée

Alphabet romain majuscule

$\mathcal{C}$	Paramètre de coût dans le noyau RBF d'un SVM
$\mathcal{C}_{comp}$	Coefficient de corrélation de la composition de deux groupes de récifs

Alphabet grec minuscule

α	Erreur associée au test de Student
γ	Paramètre de taille du noyau RBF d'un SVM

INTRODUCTION

0.1 Contexte et problématique

La situation actuelle des coraux est un enjeu fortement étudié à cause de leur important déclin. L'Agence spatiale canadienne (2012), qui les étudie à l'aide de satellites, attribue ce déclin entre autres au réchauffement climatique, à l'acidification des océans et aux interventions humaines dans leur milieu de vie. Cette situation inquiète les scientifiques, qui effectuent donc un suivi constant des variations des populations. Pour ce faire, on crée de volumineuses bases de données contenant des images de récifs coralliens provenant de toute la surface du globe. En intégrant systématiquement les informations qui s'y trouvent, les écologistes et les biologistes marins répertorient le taux de survie des espèces, leur couverture des fonds marins et l'effet de divers phénomènes sur leur survie.

L'analyse de ces images permet d'effectuer un suivi de l'évolution des espèces en fonction des conditions environnementales auxquelles chacune fait face. Les résultats de cette méthode sont utiles et nécessaires à la sauvegarde des habitats marins. Néanmoins, l'identification des entités présentes dans les images par des experts est coûteuse en temps. En effet, de la totalité des dizaines de milliers d'images recueillies par année, seules entre 1 % et 2 % sont traitées (Beijbom et al., 2012) et à peine 2 % de leur surface (CHAPITRES 1 et 2).

L'utilisation de nouveaux outils devient conséquemment une nécessité. C'est pourquoi on automatise le processus par l'application d'intelligence et de vision artificielles aux clichés et aux vidéos de ces récifs coralliens. Cette approche permet de réaffecter des ressources humaines chez les experts en écologie tout en ajoutant de nouvelles perspectives aux techniques informatiques déjà utilisées, puisque les recherches menées dans ce cadre posent de nouvelles contraintes et poussent à améliorer le domaine.

0.2 Objectifs

Ce mémoire vise à fournir à l'*Australian Institute of Marine Science* (AIMS) un outil de classification d'images de récifs coralliens. Pour ce faire, le principal objectif consiste à adapter un algorithme de reconnaissance de coraux à la base de données fournie par l'AIMS. Plus précisément, cette approche est divisée en deux sous-objectifs :

- 1) Optimiser les paramètres de chaque étape du système de reconnaissance de forme.
- 2) Évaluer la limite de performance de l'approche sur un sous-ensemble de la base de données.

0.3 Méthodologie

Afin de répondre à ces objectifs, les opérations suivantes sont entreprises :

- 1) Faire une revue de littérature qui recense les connaissances sur la reconnaissance de forme et sur l'analyse des images de récifs de coraux.
- 2) Évaluer les caractéristiques spécifiques de la base de données et leurs conséquences sur le processus de reconnaissance de forme.
- 3) Sélectionner des techniques de prétraitement, de segmentation, d'extraction de descripteurs de couleur et de texture, de réduction de la dimensionnalité et de classification ajustées à la problématique.
- 4) Optimiser les techniques expérimentalement pour la base de données de l'AIMS.
- 5) Développer une stratégie de sélection de données d'entraînement pour le système de classification optimisé.

- 6) Généraliser l'application du système de classification optimisé et de la stratégie de sélection de données d'entraînement à la totalité du sous-ensemble de la base de données de l'AIMS disponible.

0.4 Structure du mémoire

Le présent document contient trois chapitres qui explicitent des connaissances de base et les démarches entreprises au cours de ce projet. Tout d'abord, le premier chapitre contient une revue de littérature liée à la problématique de classification d'images de récifs coralliens. On y trouve donc une explication du cycle de vie des coraux et une description des initiatives en place pour en effectuer un suivi. Nous abordons aussi les étapes de la reconnaissance de forme et les particularités de son application aux images de récifs coralliens.

Le second chapitre établit une méthodologie d'optimisation de la reconnaissance de forme à cette problématique. Nous la détaillons ainsi pour les prétraitements, la segmentation, l'extraction des descripteurs, la diminution de la dimensionnalité des descripteurs et l'entraînement du classificateur. Nous appliquons également cette méthodologie à la base de données de l'AIMS et nous expliquons les résultats qui en sont issus.

Enfin, le troisième chapitre présente le développement d'une stratégie de sélection des données d'entraînement d'un classificateur en fonction des propriétés des images à cataloguer. Il explore les diverses possibilités de couplage des données d'entraînement et de test à l'intérieur de la base de données de l'AIMS. Nous en extrayons donc plusieurs recommandations pour la sélection des ensembles.

CHAPITRE 1

REVUE DE LITTÉRATURE

1.1 Les récifs coralliens

Les récifs coralliens sont des agglomérats d'organismes très diversifiés en taille, en forme et en couleur, qui couvrent le fond des mers. Les coraux constituent la base de l'habitat de près de 25 % des espèces marines de la planète et ils protègent les côtes des catastrophes naturelles provenant de l'océan. Ils influencent leur écosystème de manière notoire et servent aussi efficacement de baromètre pour déterminer la santé de leur environnement. En effet, les coraux dépendent fortement de leur milieu, qu'ils protègent et alimentent en retour (Chaudhury et Ajai, 2014). Les écosystèmes marins abritent deux grandes familles de coraux : les coraux solitaires et les coraux coloniaux. Les coraux solitaires survivent en se fixant de manière individuelle sur diverses surfaces. Les coraux coloniaux sont, pour leur part, constitués d'une structure de calcaire et d'une colonie de polypes dont chaque individu est un clone des autres. Ce sont alors ces animaux qui génèrent la structure mère de calcaire et qui se multiplient pour la faire grandir. Lorsque plusieurs de ces colonies sont réunies à un même endroit et qu'elles croissent dans un arrangement serré, un récif est créé.

Pour survivre, les coraux coloniaux ont besoin de la présence des algues qui sont la source de leurs couleurs éclatantes. Ils entretiennent avec elles une relation symbiotique. Ces algues produisent, par photosynthèse, des nutriments que les coraux absorbent. Les déchets que ces derniers dégagent sont postérieurement réabsorbés par ces mêmes algues. Cette relation influence la forme et l'emplacement des coraux, puisqu'ils se développent de manière à intercepter les nutriments et l'oxygène transportés par les courants marins, puis à offrir une grande surface disponible pour l'absorption de la lumière par les algues. Ainsi, la forme des coraux ne dépend pas uniquement de l'espèce des polypes, mais également du milieu dans lequel ils vivent.

Dans le cas où les coraux subissent du stress, leur réaction primaire est de rejeter les algues et cela cause leur mort. Conséquemment, ces animaux ne peuvent généralement vivre que dans des conditions physicochimiques très précises. Leur survie est influencée par la qualité de leur eau (salinité, acidité, turbidité, transparence), ainsi que par leur emplacement (quantité de lumière, température, profondeur, substrat) et par les cycles océaniques qu'ils subissent (vagues, nutriments, circulation)(Chaudhury et Ajai, 2014).

Puisque les récifs coralliens sont très sensibles à leur environnement, ils réagissent rapidement à tout changement qui s'y produit. Pour cette raison, des scientifiques effectuent un suivi constant de leur santé de manière à évaluer les impacts de phénomènes ponctuels, comme les ouragans, et les répercussions de variations continues, telles que le réchauffement climatique. Ces études sont menées tout autour du globe, comme dans le regroupement de récifs coralliens le plus important au monde, la Grande barrière de corail d'Australie, dont l'âge est estimé à 500 000 ans et dont la longueur dépasse 2 000 km.

1.2 Collecte des données sous-marines

L'*Australian Institute of Marine Science* (AIMS) est un organisme de recherche qui maintient un programme de surveillance à long terme des récifs des mers adjacentes à l'Australie. Il s'intéresse particulièrement à leurs couverts benthiques, soit aux structures qui recouvrent leurs sols, ainsi qu'à leurs populations de poissons et d'étoiles de mer. Il assemble donc une banque de données dont l'analyse permet d'obtenir une vue globale de l'évolution spatiale et temporelle des espèces marines. Pour ce faire, l'AIMS utilise la technique de photographie sous-marine dont (Jonker, Johns et Osborne, 2008) détaille le protocole. Elle consiste à prendre une quarantaine de photographies sur des parcours linéaires présélectionnés d'une cinquantaine de mètres de long. Chacun des 48 récifs surveillés contient cinq de ces *transects*.

Pour la cohérence de la recherche et pour que les clichés soient uniformes, le photographe se positionne à une distance constante de 50 cm par rapport au substrat marin et il prend les photos à l'intérieur d'une plage horaire permettant une illumination suffisante des sujets. Par

la suite, on marque chaque cliché de cinq points de référence dont la position est fixe. Quatre experts identifient alors les structures situées sous chacun des points et leur associent une étiquette telle que celles de la figure 1.1. Il est à noter que les biologistes sélectionnés pour faire ce travail obtiennent des écarts d'à peine 10 % pour ce qui est de l'identification de la famille des spécimens observés et qu'ils identifient tous le même nombre de clichés pour chaque récif. La précision des identifications dépend toutefois fortement du niveau de détail perceptible sur la photo, ce qui est lié à la distance focale (Ninio et al., 2003).

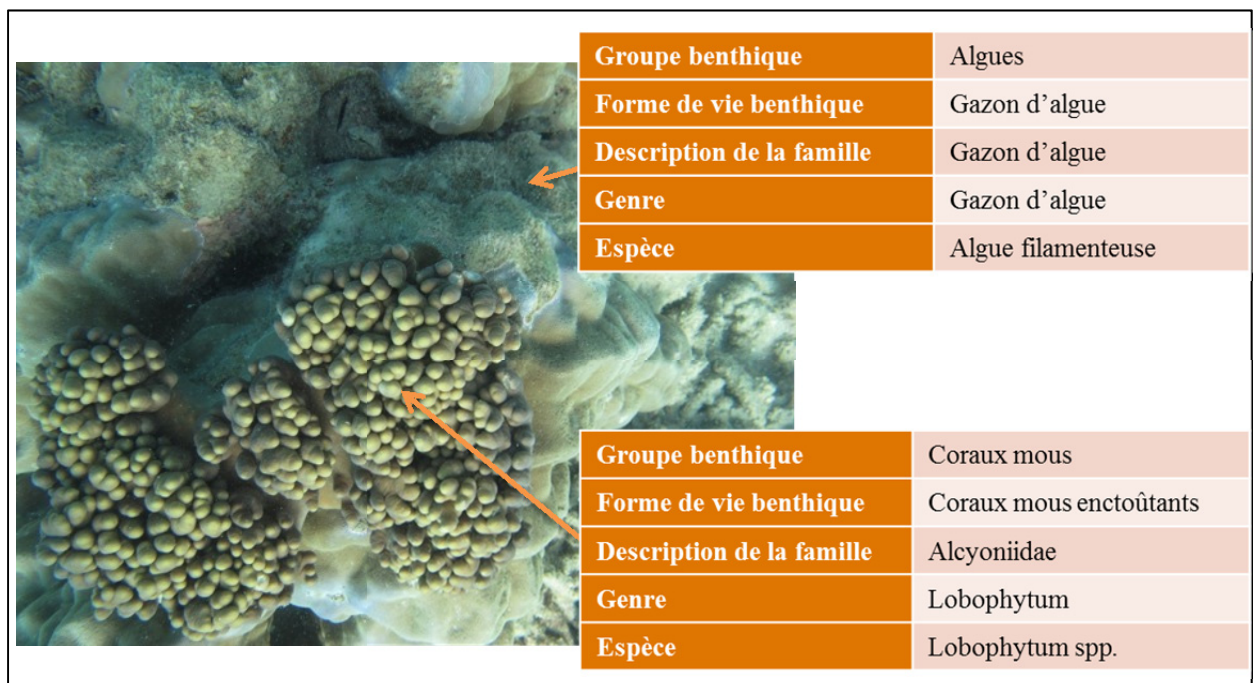


Figure 1.1 Niveaux taxonomiques étiquetés
Photo reproduite et adaptée avec l'autorisation de l'AIMS (2013)

Le niveau d'identification le plus fréquemment utilisé pour la surveillance de récifs est le groupe benthique (Jonker, Johns et Osborne, 2008). La connaissance détaillée des formes de vie benthiques, des familles, des genres et des espèces présents dans ces groupes benthiques est pratique pour aborder plusieurs questions écologiques. Cependant, l'exactitude de l'étiquetage y est plus faible. L'ajout d'information sur les clichés, comme le passage d'une image en 2D à une image en 3D, peut augmenter de manière importante la précision des identifications faites par la suite (Ninio et al., 2003).

Plusieurs autres techniques sont mises en œuvre pour obtenir les images de récifs. Ainsi, l'utilisation de satellites permet d'évaluer rapidement les changements survenus dans les récifs coralliens. Entre autres, elle repère la mort massive de coraux et le passage d'un environnement où les coraux prédominent vers un environnement principalement composé d'algues. Les satellites peuvent aussi fournir une carte de la surface marine occupée par des coraux. De surcroît, ils offrent autant des images optiques que multi-spectrales (Chaudhury et Ajai, 2014). Les images prises sous l'eau sont toutefois plus précises. Elles peuvent être récoltées par des plongeurs, par des caméras tirées par des embarcations ou encore par des véhicules robotisés, sous-marins et autonomes. En outre, on peut recueillir les images à l'aide de modalités d'imagerie supplémentaires qui permettent d'obtenir de l'information en 3D (ex : ultrasons) et spectrale (ex : infrarouges et ultraviolets). On peut également faire un suivi des images au fil du temps. Par exemple, (Bryson et al., 2013) élaborent une technique qui compare les images acquises à divers moments, en les recalant spatialement. Ainsi, on regroupe des photos prises sur une période qui s'étend de 24 heures à quelques années de manière à observer l'évolution des coraux de cet endroit précis.

L'identification manuelle de toutes ces images prend cependant aux scientifiques un temps immense. Les bases de données contiennent des dizaines de milliers d'images pour chaque année et seules de 1 % à 2 % d'entre elles sont analysées (Beijbom et al., 2012). Pour ces raisons, elles sont de plus en plus souvent annotées automatiquement, par des algorithmes de reconnaissance de forme. Cette démarche diminue significativement le temps de traitement et les biologistes peuvent se concentrer sur les problèmes des récifs coralliens, leurs causes et leurs solutions.

1.3 Reconnaissance de forme

La reconnaissance de forme est un procédé par lequel une image est observée et ses différentes régions sont associées à une étiquette d'identification selon ses caractéristiques et les connaissances préalables de l'observateur. Cette suite d'actions est triviale au niveau humain, puisqu'elle est effectuée en continu par le cerveau suite à la perception des informations par les cinq sens. Pour l'automatiser informatiquement, il est toutefois nécessaire de faire plus d'analyses. En effet, la reconnaissance de forme ne se limite pas à identifier des formes géométriques. Par exemple, pour évaluer le contenu d'une image, la vision artificielle exploite régulièrement la texture et la couleur.

Ainsi, la représentation numérique d'une image est constituée d'une suite de pixels d'intensité plus ou moins grande, associés à certains canaux de couleur. Nous pouvons par exemple l'illustrer par des teintes de gris ou encore par le rouge (R), le vert (V) et le bleu (B) comme le montre la figure 1.2. Il est possible d'identifier les structures représentées sur une image grâce, entre autres, à la texture, qui s'exprime à travers des motifs de pixels répétés (Jähne, 2004; 2005). Ainsi, les patrons de pixels dans les sous-images 1.3A et 1.3B sont différents et leur analyse permettrait de différencier deux textures, donc deux structures différentes et distinctes.

Il ne serait cependant pas possible de savoir que 1.3A représente des pois et 1.3B, des barres, sans avoir appris de quoi ont l'air des pois et des barres. Pour y arriver, il est nécessaire de constituer une base de données qui associe certains motifs à une étiquette. Cette base de données est donc une sorte de dictionnaire contenant les mots «pois» et «barre» dont la définition décrit leur texture et leur couleur respective. En comparant la texture de l'image 1.3A au dictionnaire, il serait alors possible de conclure à des pois verts. Toutefois, les images à analyser sont généralement plus complexes que celles présentées à la figure 1.3 et le processus de leur description et de leur identification n'est alors plus si simple.

Le protocole typique d'identification des étiquettes d'objets utilise des statistiques, la vision artificielle et l'intelligence artificielle. Le tout commence par la constitution d'un dictionnaire dont la description de chaque classe d'objet contient des modèles : ces modèles sont créés à partir d'objets dont la classe était identifiée au préalable et qui forment un groupe d'entraînement. Le dictionnaire sert par la suite de référence pour la classification de nouveaux objets qui appartiennent à un groupe de test.

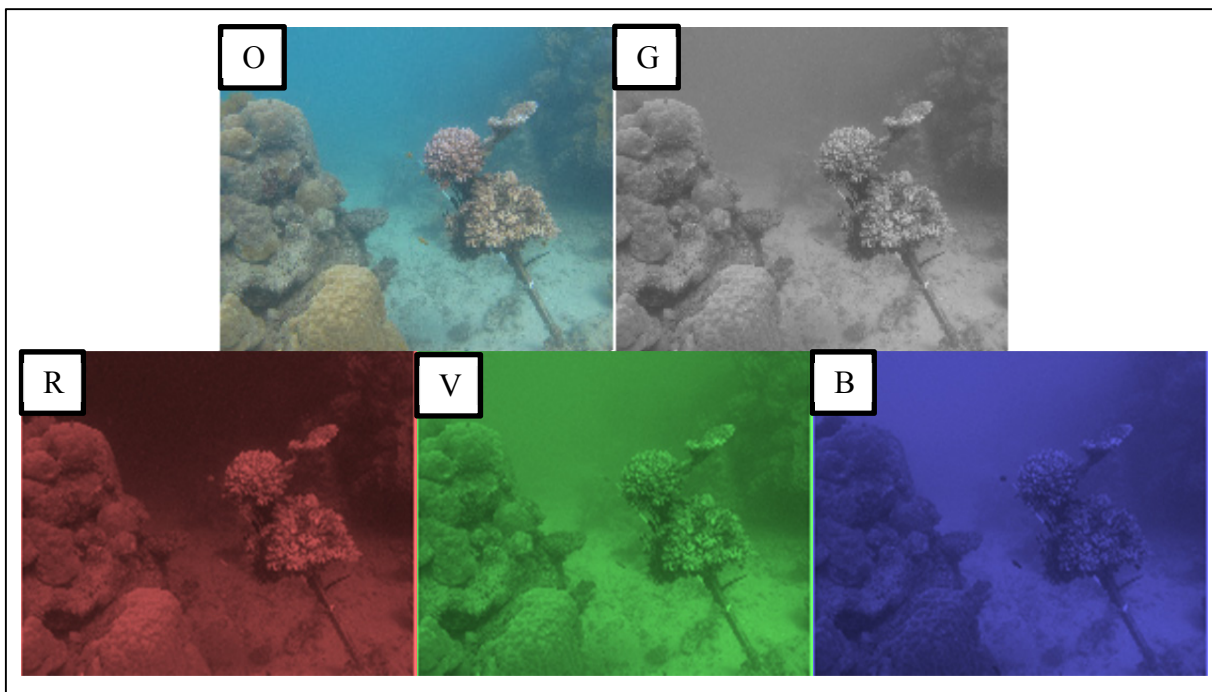


Figure 1.2 Décomposition d'une image (O) dans les canaux gris (G), puis rouge (R), vert (V) et bleu (B)

Photo reproduite et adaptée de l'AIMS (2013)

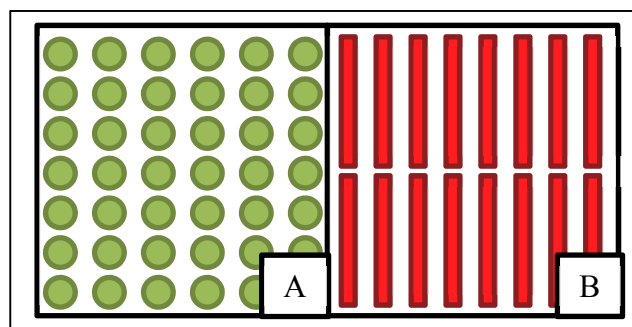


Figure 1.3 A) Représentation de pois verts et B) de barres rouges

Pour construire les modèles des groupes d'entraînement et de test, la première étape consiste à appliquer un prétraitement aux images de manière à améliorer leur qualité et à optimiser la justesse des informations qui en seront tirées. Nous découpons ensuite une région d'intérêt par segmentation autour des objets, puis nous extrayons les caractéristiques, ou descripteurs, de ceux-ci. Enfin, nous développons un classificateur par le traitement du groupe d'entraînement et nous y comparons les modèles du groupe de test pour leur assigner une étiquette identifiant leur classe. L'ensemble du processus est présenté à la figure 1.4.

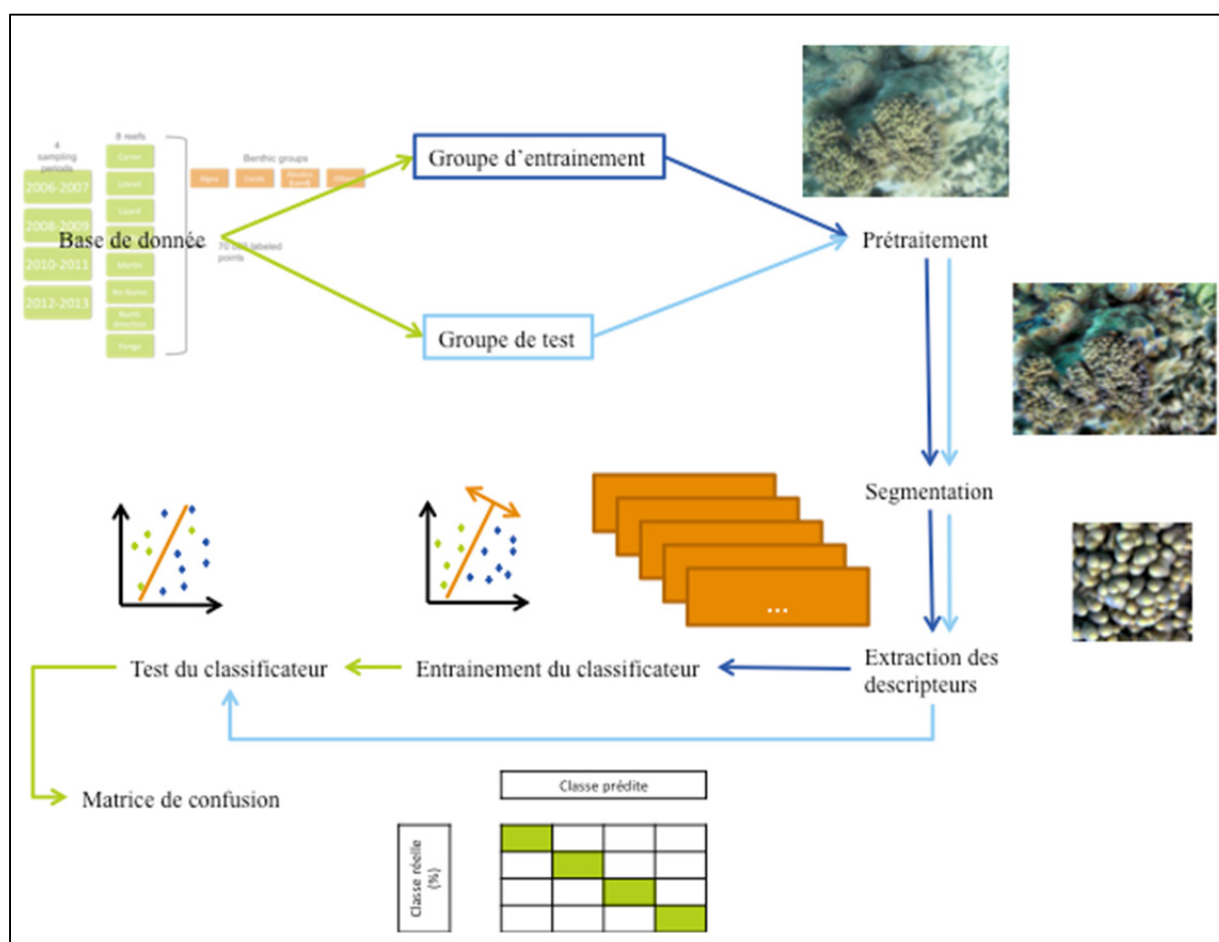


Figure 1.4 Processus de classification d'une image de récif corallien
Photo reproduite et adaptée de l'AIMS (2013)

Nous illustrons enfin les résultats dans une matrice de confusion telle que celle de la figure 1.5. Les lignes y correspondent à la classe d'appartenance réelle d'un échantillon, alors que

les colonnes correspondent à la classe prédite par le classificateur. Conséquemment, un échantillon qui a été associé à la bonne classe est situé dans la diagonale de la matrice de confusion. Les indicateurs de performance de base sont l'exactitude (éq. 1.1), le taux de rappel (éq. 1.2) et la précision (éq. 1.3), qui représentent respectivement le taux d'échantillons bien classés dans l'ensemble du groupe, la proportion d'une classe qui a bien été classée et la proportion d'échantillons bien étiquetés à travers toutes les étiquettes d'une classe prédite.

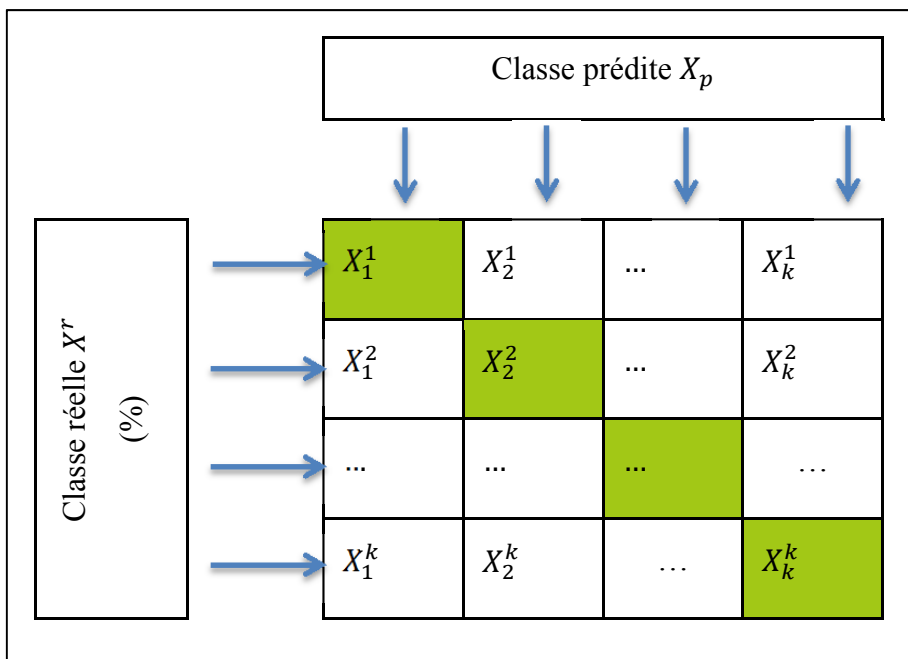


Figure 1.5 Représentation d'une matrice de confusion

L'exactitude correspond ainsi à la somme des éléments de la diagonale par rapport au nombre total d'objets. Le taux de rappel est, pour sa part, calculé le long d'une ligne de la matrice de confusion, donc par rapport à la classe réelle des échantillons. Il équivaut au nombre d'objets dans la cellule de la diagonale de cette classe en comparaison avec le nombre d'éléments sur la ligne. Le taux de rappel exprime le taux d'échantillons correctement classifiés pour une classe spécifique. La précision est enfin mesurée le long d'une colonne de la matrice de confusion, soit avec tous les échantillons étiquetés selon une certaine classe. Elle est équivalente au nombre d'échantillons dans la cellule de la diagonale par rap-

port au nombre d'échantillons dans la colonne qui lui correspond. Elle mesure donc la quantité d'échantillons associés à une classe qui auraient effectivement dû y être associés.

$$exactitude = \frac{\sum_{i=1}^k X_i^i}{\sum_{j=1}^k \sum_{i=1}^k X_i^j} \quad (1.1)$$

$$rappel_{classe\ i} = \frac{X_i^i}{\sum_{j=1}^k X_j^i} \quad (1.2)$$

$$précision_{classe\ i} = \frac{X_i^i}{\sum_{j=1}^k X_i^j} \quad (1.3)$$

Dans le cas de la reconnaissance de coraux, nous devons adapter chaque étape aux contraintes propres à ce champ d'étude. Par exemple, les prétraitements doivent tenir compte du fait que les images sont sous-marines. Elles peuvent donc être bruitées, floues et faiblement contrastées. Elles sont également bleuâtres et verdâtres, car l'eau absorbe le rouge, et l'éclairage peut manquer d'uniformité (Abdul Ghani et Mat Isa, 2015). La segmentation doit quant à elle délimiter une structure qui correspond à l'étiquette et nous devons ajuster la classification à la grande variabilité des coraux à l'intérieur d'une même espèce selon l'endroit où ils se sont développés. Le dictionnaire contient pour sa part des classes telles que les groupes benthiques (algues, coraux durs, coraux mous, abiotiques, millépores, éponges ou autres) et leurs caractéristiques.

Donc, l'identification d'objets benthiques pose des contraintes importantes. Le peu d'uniformité des images à traiter en est particulièrement la cause. Contrairement à plusieurs autres situations d'imagerie, les sujets observés varient constamment en taille, en forme et en couleur, même à l'intérieur d'une seule classe d'objets. Les clichés utilisés dans le cadre de ces études n'illustrent pas, non plus, un objet unique, mais plutôt un amalgame de divers organismes sous-marins qui se superposent et s'entremêlent, compliquant d'autant plus leur distinction. Pour ces raisons, le choix des techniques d'analyse doit être fait de manière judi-

cieuse : elles doivent être en mesure de discerner des objets naturels ayant une faible variabilité interclasse, ainsi qu'une grande variabilité intraclasse et inter-site dans la morphologie des espèces (Shihavuddin et al., 2013). De plus, puisque les structures identifiées sont complexes, l'étalon de référence est biaisé et chaque base de données possède ses particularités propres.

1.3.1 Bases de données

La reconnaissance de forme peut être développée sur des bases de données de textures, dont la classification représente une simplification du problème de classification d'espèces benthiques. Celle qu'on utilise le plus souvent est la Columbia-Utrecht (CURTeT), qui est exploitée par (Liu, Lin et Yu, 2009; Shihavuddin et al., 2013; Varma et Zisserman, 2005) pour leurs expérimentations. Les banques de données de textures contiennent typiquement des images de surfaces d'objets dont l'illumination, l'angle d'observation et l'échelle varient. Généralement uniformes, elles ne contiennent qu'un objet et sa texture unique. L'analyse est donc significativement plus difficile dans le cas d'images de récifs coralliens, où les scènes sont complexes.

Pour la classification automatique d'espèces benthiques, les bases de données sont constituées soit de photographies simples soit de vidéos, en tons de gris ou en couleurs rouge, vert et bleu (RVB), avec ou sans cadre de référence. Ces cadres fournissent des indicateurs de couleurs et assurent la conformité des images entre elles. Toutefois, ils peuvent aussi constituer des artéfacts dans les clichés. Le sous-ensemble de bases de données le plus souvent utilisé est le *Moorea-labeled corals* (MLC), qui est un ensemble de 2055 clichés comptant chacun 25 points identifiés répartis selon neuf classes. Ces images ont la particularité de contenir un cadre sur lequel se retrouvent les couleurs de référence RVB. Le MLC fait partie de la base de données du *Moorea Coral Reef-Long Term Ecological Research* (MCR-LTER).

Le sous-ensemble de la base de données de l'AIMS utilisé pour ce projet a, pour sa part, été élaboré selon le protocole décrit dans la section 1.2. Il contient près de 77 000 points identi-

fiés sur environ 15 400 images de 2448 par 3264 pixels, soit cinq étiquettes par image. Grâce à la constitution de l'outil automatisé présenté dans ce mémoire, leur potentiel pourrait cependant s'élever à 970 200 points en exploitant des régions carrées de 350 pixels de côté sur toute la surface des images, permettant 63 étiquettes par image (CHAPITRE 2).

Les points sont étiquetés selon trois méthodes. La première implique une description taxonomique complète (famille, genre, espèce), alors que la seconde exploite les formes de vie benthiques. La troisième est quant à elle une description sommaire de ces formes, qu'elle divise en sept groupes. Ce sont ces derniers qui sont associés aux classes dans notre cas et qu'on nomme les groupes benthiques. La constitution du sous-ensemble s'est faite entre 2006 et 2013, sur huit récifs (*Carter reef*, *Lizard Island*, *Macgillivray reef*, *Martin reef*, *No Name reef*, *Yonge reef*, *North direction reef*, *Linnet reef*). Les échantillons ont été recueillis à travers les années et les récifs selon la figure 1.6. Leur distribution à travers les groupes benthiques est donnée à la figure 1.7, qui montre une fréquence d'occurrence beaucoup plus haute pour les algues que pour les autres classes. En effet, elles comptent pour plus de 50 % des échantillons. Les classes de cette base de données ne sont donc pas équilibrées, ce qui représente une contrainte pour la classification.

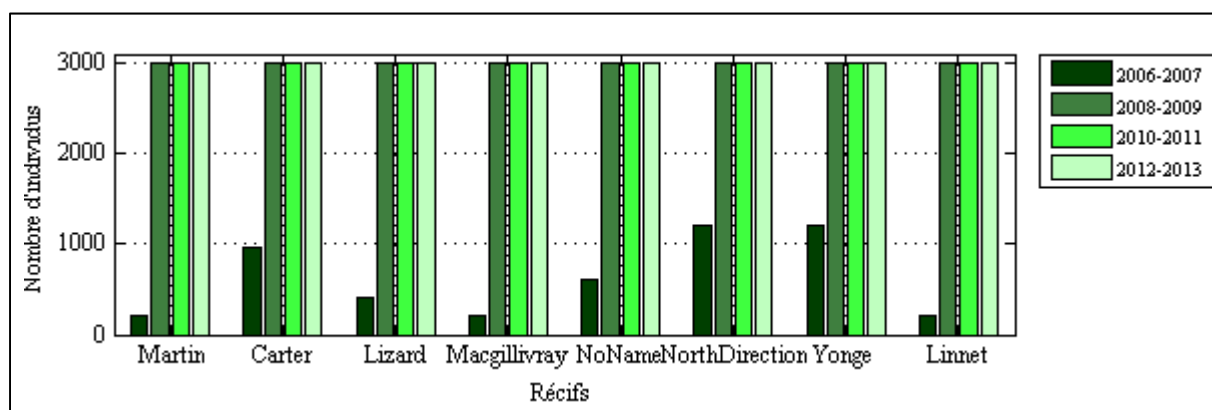


Figure 1.6 Nombre d'échantillons dans les récifs en fonction de la période évaluée (*Voir ANNEXE III, p. 103*)

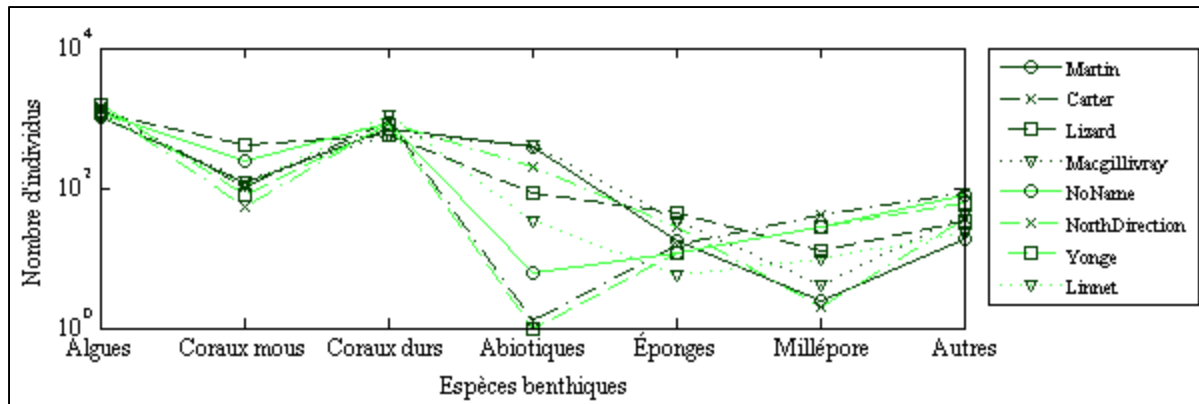


Figure 1.7 Nombre d'échantillons de chaque groupe benthique selon les récifs (*Voir ANNEXE III, p.103*)

1.3.2 Débalancement des classes

La présence d'une classe significativement plus importante que les autres dans un groupe d'entraînement peut avoir pour effet de favoriser l'attribution de cette étiquette lors du test, car la probabilité de se tromper en y plaçant un échantillon est moindre. Ce débalancement influence donc la capacité de catalogage des classificateurs : il y a généralement un surapprentissage en faveur des classes populeuses.

Le débalancement en tant que tel n'est pas la cause du problème, qui résulte plutôt d'autres aspects des bases de données amplifiant les difficultés (Sun et al., 2007). Tout d'abord, les échantillons très petits contiennent peu d'information. Augmenter le nombre d'individus dans chaque classe compense alors le phénomène, même si les proportions sont conservées. Ensuite, une faible séparabilité entre les classes est également problématique, puisque, si les espaces de descripteurs de deux catégories sont superposés, il est difficile de les distinguer. En dernier lieu, le fait que les grandes classes soient composées de sous-catégories complique aussi la délimitation de leurs frontières. Puisque les sous-catégories ont des caractéristiques diversifiées, les descripteurs sont dispersés, ce qui augmente la complexité d'apprentissage. Ce problème peut être difficile à surmonter parce qu'il est généralement intrinsèque à la base de données.

Plusieurs techniques permettent de diminuer ces erreurs. Les premières agissent directement sur la base de données, en sous-échantillonnant les grandes classes ou en suréchantillonnant les petites de manière à équilibrer le nombre d'échantillons. Cet équilibre varie, néanmoins, pour chaque base de données. En effet, il n'est pas toujours souhaitable que le ratio des populations soit égal (Sun et al., 2007). De plus, le sous-échantillonnage détruit de l'information qui pourrait être pertinente, alors que le suréchantillonnage, en créant de l'information, peut introduire des erreurs de modélisation (Anand et al., 2010).

Les autres techniques s'appliquent plutôt au niveau de l'algorithme de classification. Par exemple, le coût de l'erreur de classification des petites classes peut être posé comme étant plus considérable que celui des grandes classes. L'algorithme favorise donc la bonne classification des petites classes. Il est également possible de déplacer les frontières qui séparent les classes de manière à englober une plus grande part de la petite classe, quitte à ce qu'il y ait d'avantage de fausses prédictions (Yu et al., 2015).

Par ailleurs, nous pouvons modifier les critères de performance lors de l'optimisation des paramètres des classificateurs. En effet, utiliser l'exactitude mène à un surapprentissage, soit à l'intégration de bruit aléatoire au modèle plutôt que de liens significatifs entre les classes et les descripteurs (Jiang, Missoum et Chen, 2014). D'autres options sont aussi disponibles comme l'exactitude balancée, qui est une moyenne des taux de rappel de chaque classe, l'aire sous la courbe *Receiver Operating Characteristic* (Jiang, Missoum et Chen, 2014) ou les critères *F-measure* (eq. 1.4) et *G-mean* (eq. 1.5) (Sun et al., 2007; Yu et al., 2015). Ainsi, *F-measure* reflète la capacité d'une classe en particulier à récolter tous les échantillons qui lui appartiennent, associée au rappel, tout en ne recevant pas les autres, soit la précision. *G-mean* représente, pour sa part, la capacité de l'ensemble des k classes à rapatrier leurs échantillons sans être biaisée par leur taille.

$$F - measure_i = \frac{2 * précision_i * rappel_i}{précision_i + rappel_i} \quad (1.4)$$

$$G - mean = \sqrt[k]{\prod_{i=1}^k rappel_i} \quad (1.5)$$

1.3.3 Prétraitements

L'acquisition d'images est, en soi, une opération qui dégrade leur qualité. Chaque étape de la transformation d'un objet en sa représentation numérique, par exemple le milieu à travers lequel la lumière doit passer pour se rendre à l'objectif, la lentille de la caméra, la numérisation et la compression, introduit des artéfacts et du bruit (Shih, 2010). Les prétraitements sont donc utiles pour extraire les informations pertinentes de l'image et pour rendre son analyse plus facile. Ils permettent d'éliminer les éléments indésirables et de mettre l'accent sur ceux qui sont primordiaux. Réduire le bruit dans l'image tout en conservant des contours nets est un défi important. Les prétraitements peuvent agir soit dans le domaine spatial, directement sur les pixels, soit dans le domaine fréquentiel, ou encore sur les deux en même temps selon l'utilisation subséquente de l'image. Pour arriver à l'amélioration recherchée, l'utilisation itérative de plusieurs techniques de prétraitement différentes peut être nécessaire.

L'opération la plus simple au niveau spatial est la modification des niveaux de gris dans une image, qui sert à augmenter son contraste et à permettre une meilleure visualisation de ses détails. Plusieurs de ces cas sont présentés à la figure 1.8. Pour y arriver, il est possible d'imposer un seuil à son intensité, ou encore de modifier la distribution de ses valeurs. Par exemple, calculer le logarithme de l'intensité permet de faciliter la perception humaine des niveaux de gris. Dans les cas où le contraste est faible, comme lorsque l'éclairage est déficient, la plage dynamique d'intensité utilisée dans l'image est petite. L'allongement du contraste étire alors cette plage pour atteindre un intervalle maximal, ce qui a pour effet d'augmenter la luminosité.

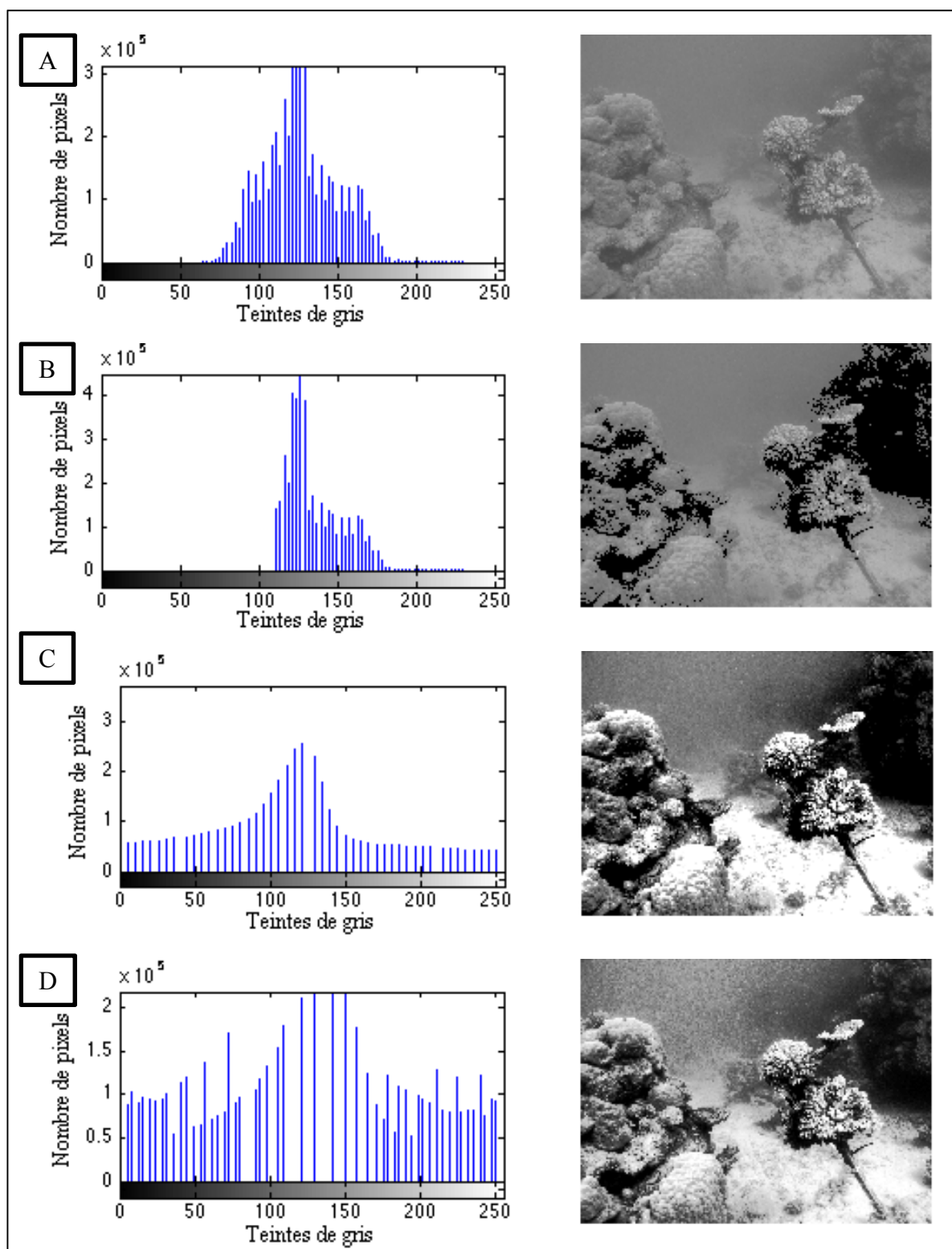


Figure 1.8 Illustration des transformations d'un histogramme A) original, B) avec seuillage, C) avec allongement D) avec égalisation
Photos reproduites et adaptées de l'AIMS (2013)

La spécification d'histogramme permet également de modifier la distribution de l'intensité. Cette méthode repose sur l'hypothèse que les niveaux de gris devraient être peuplés selon une certaine distribution. Ainsi, l'égalisation d'histogrammes selon une loi uniforme suggère que chacun des niveaux de gris est équiprobable. L'égalisation d'histogramme amplifie les détails des images, mais elle augmente aussi l'apparence du bruit proportionnellement.

L'égalisation de l'histogramme sur l'ensemble de l'image n'est cependant pas toujours souhaitable. En effet, si la composition de l'image n'est pas uniforme, l'égalisation favorise certaines régions au détriment des autres. Dans ces cas, la spécification d'histogramme adaptative limitée en contraste (CLAHS) de (Zuiderveld, 1994) propose une solution : plutôt que d'appliquer la spécification d'histogramme à l'ensemble de l'image, elle est appliquée à des sous-régions. Les sous-images sont ensuite jointes par interpolation. Cette méthode permet d'augmenter le contraste localement, même si l'image n'est pas uniforme.

Les prétraitements dans le domaine fréquentiel sont, pour leur part, effectués à partir de filtres qui pourront soit réduire les hautes fréquences de l'image, ce qui équivaut à une atténuation du bruit, soit augmenter leur contribution, ce qui met en valeur les contours d'objets (Shih, 2010). Pour réduire le bruit, les filtres les plus communs ont comme propriété de réduire les différences entre les pixels voisins d'une image. Le filtre moyennant effectue une simple moyenne sur un carré de pixels et l'associe au pixel central. Le filtre gaussien fait le même processus, mais il donne plus de poids aux pixels centraux. Il conserve donc mieux les caractéristiques de l'image initiale. Le filtre médian, pour sa part, ne conserve que la valeur médiane du dit carré de pixels. Ce filtre réduit autant le bruit que le filtre moyennant, mais il préserve généralement mieux les détails. Il est efficace lorsque du bruit en forme d'impulsions, qui donne un effet de poivre et sel, est présent dans l'image.

Les filtres d'affutage servent quant à eux à augmenter l'apparence des contours de formes. Lorsque le bruit a été atténué, les hautes fréquences correspondent aux endroits dans l'image où il y a des changements abrupts dans l'apparence. Une première méthode est donc d'éliminer les basses fréquences, qui forment le fond de l'image, par un filtre passe-haut.

Seuls les contours sont alors conservés. L'alternative est, au contraire, de garder les basses fréquences, mais en amplifiant les hautes fréquences. Ce faisant, le fond de l'image initiale reste visible, mais les contours sont marqués d'un fort accent.

Dans le cas spécifique des images sous-marines, la qualité est influencée par le milieu dans lequel elles sont récoltées. L'eau entraîne une atténuation de la couleur rouge et une diminution de l'intensité de l'éclairage par rapport à la propagation de la lumière à l'air libre. Bref, elle diminue l'étendue de la plage dynamique des couleurs. Elle provoque aussi une plus grande diffusion de la lumière, ce qui rend les images moins nettes (Abdul Ghani et Mat Isa, 2015). Ces phénomènes diminuent la capacité du système à les différencier.

La méthode la plus intuitive de rétablissement de la plage dynamique est l'utilisation d'étalons de couleurs à l'intérieur des clichés (Beijbom et al., 2012). En ayant un échantillon de la représentation des couleurs rouge, vert et bleu à travers les images, il est possible de rétablir les couleurs avant l'atténuation par l'eau. Ce procédé permet donc d'avoir un éclairage uniforme à travers les images.

Lorsqu'aucun étalon n'est disponible, la balance des blancs est une alternative fréquemment utilisée (Shihavuddin et al., 2013). Cette méthode prend pour acquis que, dans un monde idéal, toutes les couleurs sont représentées uniformément. Elle fait l'hypothèse que la moyenne du monde est gris. Conséquemment, la valeur moyenne de chacun des canaux de couleurs RVB est posée au même seuil, puis leurs plages dynamiques respectives subissent un allongement du contraste de manière à ce que 1,5 % des pixels soient saturés et que 1,5 % des pixels soient noirs pour chacun des canaux. Ce faisant, les variations dues à l'éclairage et à la turbidité de l'eau sont compensées (Beijbom et al., 2012). La normalisation de couleur complet (CCN), proposée par (Finlayson, Schiele et Crowley, 1998), utilise la même hypothèse. Cette normalisation vise, plus spécifiquement, à éliminer la variation de l'illumination à travers l'image et à compenser sa couleur. Pour ce faire, les couleurs sont normalisées tant à travers l'image que sur chaque pixel.

L'hypothèse d'un monde gris ne peut cependant pas être généralisée à tous les cas. Elle n'est valable que si une grande richesse de couleur est présente dans l'image (Ebner, 2009). Ainsi, prétraiter une image monochrome à l'aide de techniques qui se basent sur cette hypothèse n'améliorera pas la représentation. Plutôt que de simplement équilibrer les couleurs pertinentes, cette hypothèse augmentera aussi les couleurs absentes. Ces techniques ne sont donc pas parfaitement adaptées aux images sous-marines car celles-ci sont souvent déficientes en rouge.

Certaines méthodes de prétraitement ont été développées pour répondre à ce cas spécifique. Leur but est d'augmenter le contraste et de diminuer le bruit dans les images en utilisant des techniques adaptées d'amélioration. Le modèle d'intégration de couleurs à partir de la distribution de Rayleigh (ICM-R) est l'une d'elles (Abdul Ghani et Mat Isa, 2015). Elle exploite les propriétés d'absorption de la lumière pour éviter de surexposer et de sous-exposer les régions de l'image sous-marine.

Ainsi, ICM-R impose aux histogrammes de couleurs la distribution de Rayleigh, qui défavorise les intensités extrêmes et qui est reconnue comme la plus efficace pour les images sous-marines. Ensuite, les canaux de couleurs sont allongés sur la plage dynamique de 255 niveaux en tenant compte que le rouge est généralement sous-représenté dans les images (niveaux d'intensité faibles) et que le bleu est la couleur la plus fréquente (niveaux d'intensité élevés). Conséquemment, l'intensité du rouge est amplifiée pour s'étendre des niveaux 13 à 255, alors que celle du bleu est réduite pour couvrir les niveaux 0 à 242. Le vert est, pour sa part, étiré dans les deux sens, soit de 0 à 255. Enfin, la saturation et l'intensité globale de l'image sont eux aussi étirés sur l'ensemble de la plage dynamique. Un exemple de ce traitement est présenté à la figure 1.9.

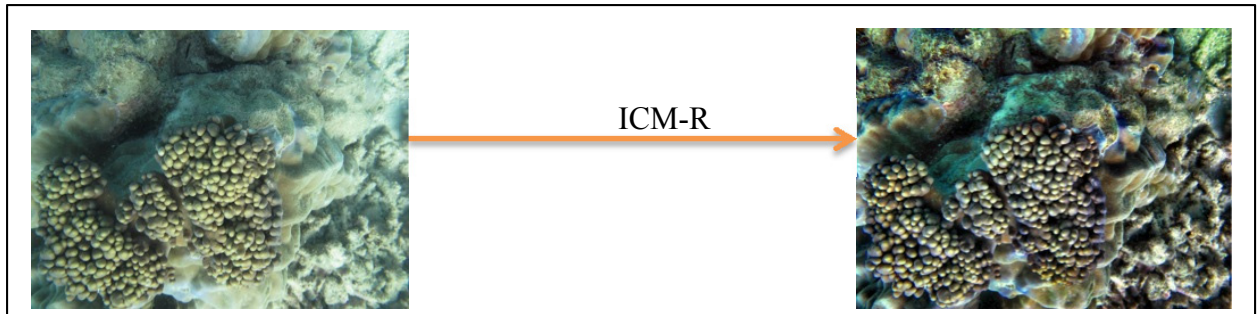


Figure 1.9 Traitement d'une image par ICM-R
Photo reproduite et adaptée de l'AIMS (2013)

1.3.4 Segmentation

La segmentation sert à découper les zones d'intérêt d'une image. Ces zones doivent contenir assez d'information sur la structure observée pour en faire un modèle fidèle. En contrepartie, elles ne doivent idéalement contenir que cette structure. Tout intrus dans la zone introduit des erreurs dans le modèle en y faisant varier la texture et la couleur. Ces intrusions ne sont cependant pas toujours évitables. En bref, la segmentation vise à réduire la taille de la région d'intérêt pour qu'elle soit plus pertinente à l'analyse et pour faciliter son traitement (Yogamangalam et Karthikeyan, 2013).

Sélectionner une région carrée autour d'un point d'intérêt est une version simplifiée de la segmentation. En choisissant judicieusement la taille du carré, il est possible de limiter le nombre d'intrus tout en conservant une part importante de la structure d'intérêt. La taille idéale dépend fortement du contenu de la base de données et de l'arrangement des structures dans les images. Cette méthode est rapide à mettre en place et, dans les images complexes, elle évite que l'on rejette les structures d'intérêt en ne conservant que les intrus. Le modèle des caractéristiques de l'image contient donc toujours, en statistique, un biais dû aux intrus qui ont été inclus dans le carré.

Autrement, la segmentation peut être axée sur le seuillage, sur les contours des formes ou sur les régions de l'image (Shih, 2010). Puisqu'aucune de ces techniques n'est très efficace, nous en faisons souvent des hybrides. Ainsi, le seuillage prend pour acquis que le fond de l'image

et l'objet d'intérêt ont des intensités différentes. Il effectue donc une découpe en se basant sur ce seul critère. Pour leur part, les techniques utilisant les contours des formes tentent de trouver des discontinuités abruptes dans l'intensité, la couleur ou la texture des images. Elles exploitent les variations des caractéristiques tout au long de la surface de l'image pour définir des frontières. Pour ce faire, une première méthode est d'utiliser un modèle de contour actif qui équivaut à entourer un contour fermé à l'aide d'un élastique. Ce dernier colle ainsi au contour lorsqu'on le laisse aller. Ce procédé n'est cependant pas bien adapté aux images complexes. Une autre méthode est de trouver les contours par des filtres. Ce faisant, les valeurs maximales de la dérivée de l'image identifie la position des discontinuités.

Les techniques de segmentation basées sur les régions recherchent au contraire des similitudes dans les caractéristiques de l'image. Dans ces cas, nous utilisons des points répartis dans l'image comme points de départs. Les régions s'étendent ensuite autour de ces points en fonction de la similitude des pixels. Nous pouvons ensuite fusionner les différentes surfaces pour réunir celles qui sont similaires et adjacentes. Il est aussi possible d'effectuer un partitionnement de l'image au niveau de ses descripteurs en utilisant un histogramme des caractéristiques de l'image pour former des groupes de pixels similaires.

1.3.5 Extraction des descripteurs de texture

Les méthodes disponibles pour caractériser la texture dans les images sont très diversifiées. Le degré de difficulté de leur traitement, leur temps de calcul et leur efficacité sont variables selon leur contexte d'application. Nous les choisissons en fonction des spécificités de la structure représentée. Une technique plus complexe n'est donc pas garante de meilleurs résultats.

Le choix d'une technique nécessite ainsi de tenir compte des particularités de l'image. Par exemple, il arrive que la texture d'une structure change d'orientation d'une image à l'autre et qu'elle soit perçue différemment de près et de loin. La méthode choisie doit donc bien représenter cette base de données qui a des contraintes des motifs en rotation et en échelle. Les

variations en illumination ont aussi des impacts significatifs sur la perception des textures par les algorithmes puisqu'elles sont les symptômes d'un éclairage non uniforme. Les techniques sélectionnées doivent pallier ces contraintes.

Les opérations de base pour analyser les textures comprises dans un vecteur x sont, tout d'abord, la moyenne (éq. 1.6), ensuite la variance (éq. 1.7) et finalement leurs dérivées (éq. 1.8 et 1.9) qui permettent d'évaluer les tendances statistiques des valeurs analysées. Elles peuvent s'appliquer directement aux pixels d'une image ou encore aux informations qui en ont été extraites (Jähne, 2004; 2005).

$$x = [x_1, \dots, x_i, \dots, x_N]$$

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i \quad (1.6)$$

$$var(x) = \sigma^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2 \quad (1.7)$$

$$asymétrie(x) = \frac{1}{\sigma^3} \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^3 \quad (1.8)$$

$$kurtosis(x) = \frac{1}{\sigma^4} \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^4 \quad (1.9)$$

1.3.5.1 Espaces de représentation

Une image ne montre pas toujours la même information selon la méthode utilisée pour l'observer. Par exemple, elle peut varier au travers des échelles, comme le montre la figure 1.10. La pyramide gaussienne et la pyramide laplacienne sont des outils qui permettent de prendre en compte ces différences et de faciliter leur analyse (Jähne, 2004; 2005). La pyra-

vide gaussienne filtre, puis sous-échantillonne itérativement les pixels de l'image initiale, de manière à obtenir une série de représentations de plus en plus grossières. La pyramide laplacienne est, pour sa part, une série d'images résultantes, issues de la différence entre les images consécutives dans la série de la pyramide gaussienne. Cette manœuvre permet d'isoler les structures fines qui appartiennent à chaque tranche d'échelle.

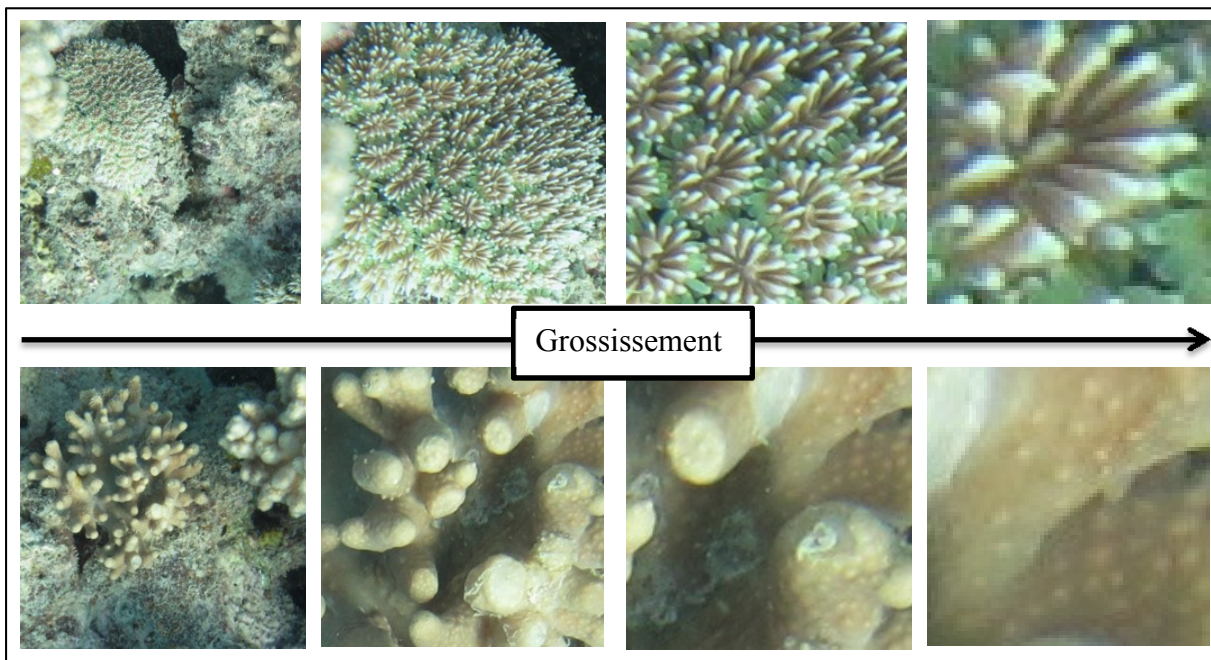


Figure 1.10 Variation de la texture à travers les échelles
Photos reproduites et adaptées de l'AIMS (2103)

L'image peut également être représentée dans l'espace de Fourier, comme le montre la figure 1.11. Cette transformée de l'image n'illustre pas la relation spatiale entre les pixels de manière traditionnelle. Elle présente plutôt leur relation en fréquence par une somme de sinusoides de différentes périodes. Ces sinusoides se caractérisent par leur amplitude et leur phase. Puisque les sinusoides sont périodiques, nous pouvons les utiliser pour détecter les textures d'une image, elles aussi périodiques.

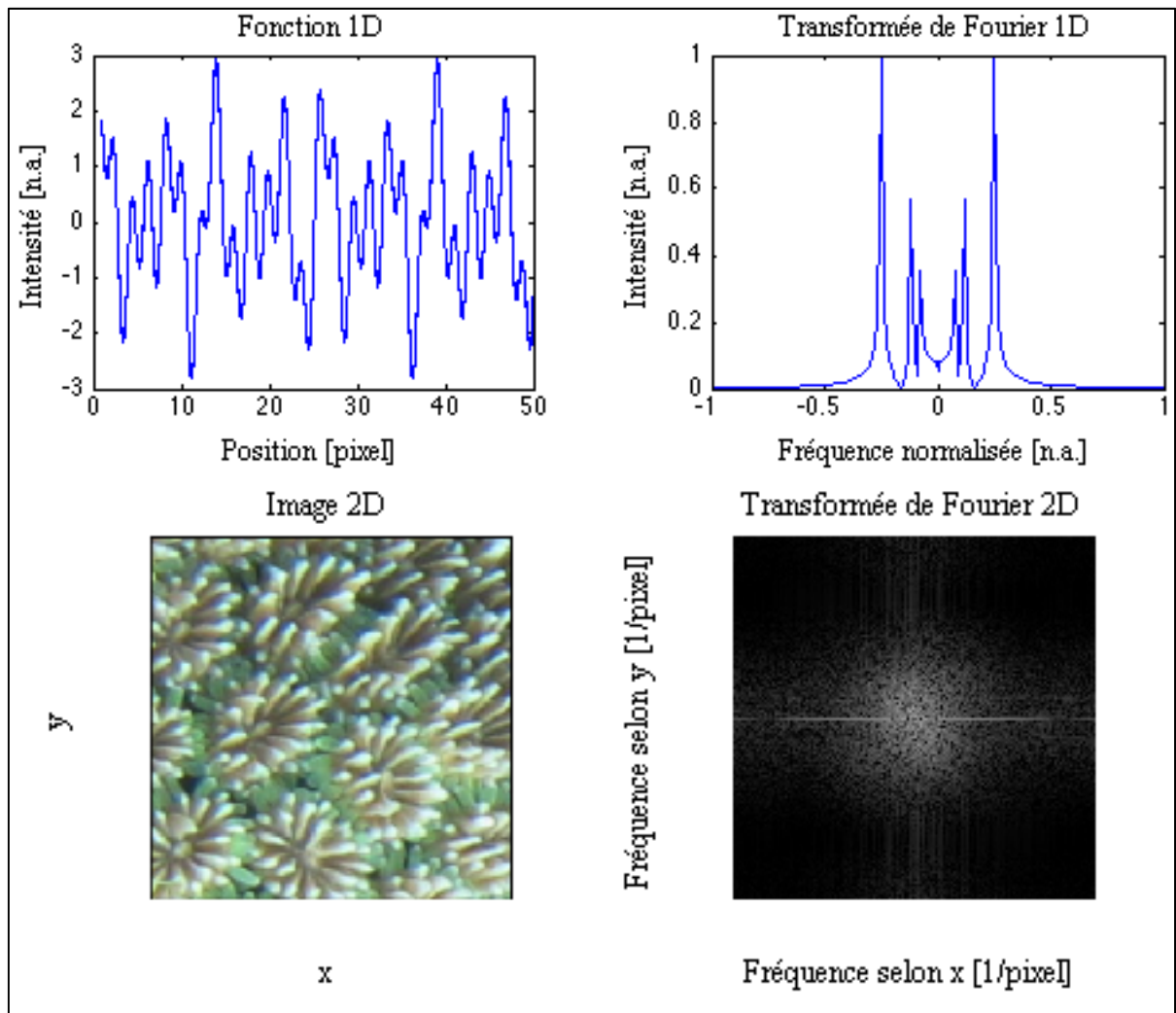


Figure 1.11 Illustration d'une transformée de Fourier en 1D et en 2D
Photo reproduite et adaptée de l'AIMS (2013)

1.3.5.2 Analyse statistique de la distribution des pixels

L'analyse de la texture met en lumière l'arrangement des pixels entre eux. Leurs relations en intensité et leurs dispositions spatiales sont autant d'aspects qui permettent de distinguer des structures. Les types de descripteurs suivants exploitent la répartition statistique des pixels dans la texture pour caractériser les images.

Histogrammes de tons de gris

Un histogramme distribue les pixels d'une image selon leur intensité. Ensuite, le nombre de pixels qui appartient à chaque groupe est comptabilisé, comme le présente la figure 1.8A. Nous pouvons extraire des statistiques (moyenne, variance, etc.) sur la distribution des pixels.

Intégration de la transformée de Fourier

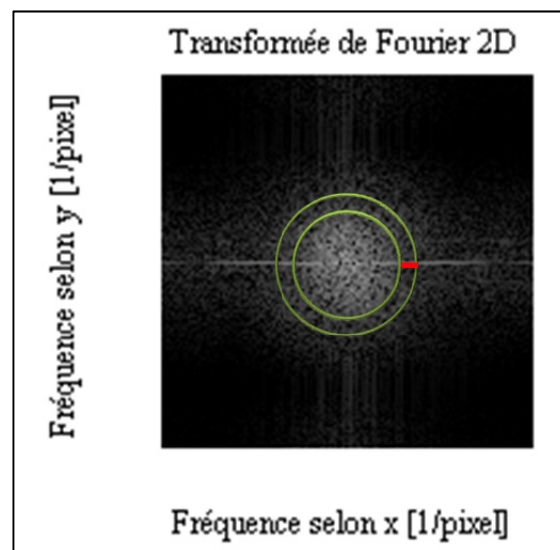


Figure 1.12 Zone de l'intégration pour une plage de fréquence. La zone entre les deux cercles est là où on intègre

L'intégration de la transformée de Fourier repose sur l'hypothèse que la texture a une forte représentation pour certaines fréquences précises qui lui sont spécifiques (Bouchard, 2011). Conséquemment, en divisant l'espace de Fourier en plusieurs plages de fréquences et en y intégrant les valeurs de la transformée de Fourier, il est possible de déduire quels intervalles de fréquences sont les mieux représentées dans l'image. La figure 1.12 montre la zone d'intégration sur une de ces plages.

Matrices de cooccurrence de teintes de gris

Les matrices de cooccurrence servent à évaluer la fréquence d'apparition de couples contigus de pixels. Ainsi, chaque combinaison de pixels de différente intensité et selon une orientation définie est enregistrée et des statistiques sur ces distributions sont calculées (Haralick, Shanmugam et Dinstein, 1973). La figure 1.13 présente un exemple d'extraction des couples pour les matrices de cooccurrence. L'ANNEXE I contient une liste de statistiques pertinentes pour l'évaluation des textures par cette technique.

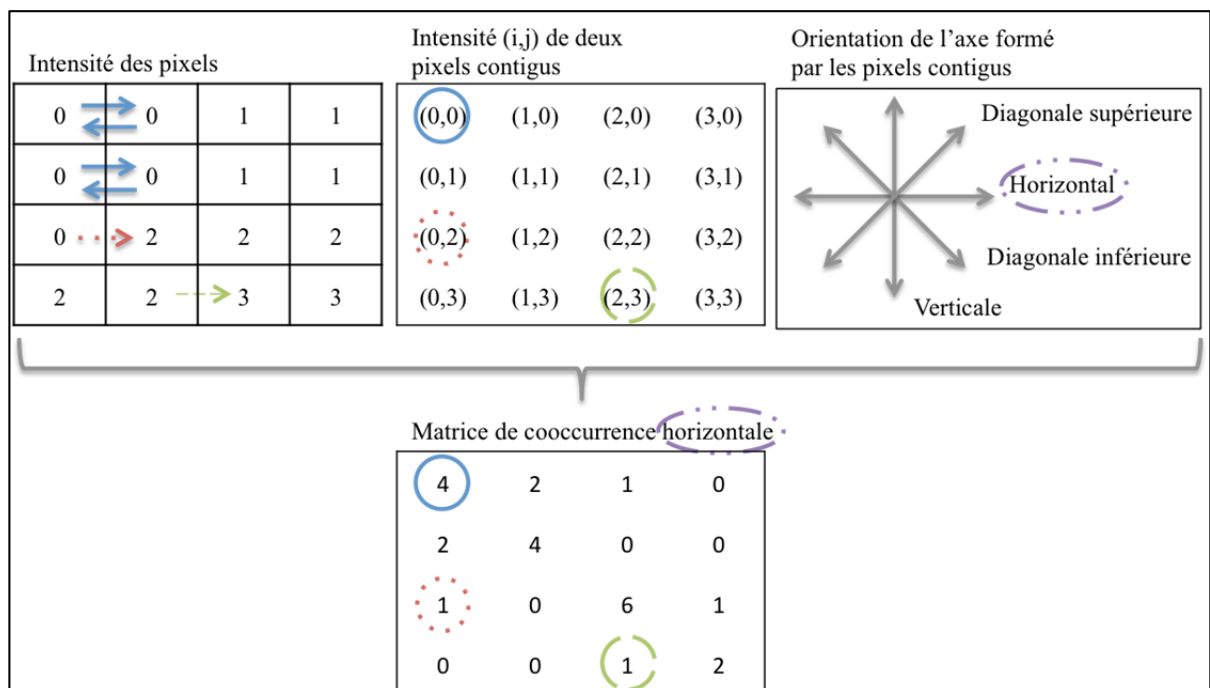


Figure 1.13 Exemple de génération de matrices de cooccurrence
Adaptée de Haralick et al. (1973)

Motifs linéaires binaires (LBP ou *Linear Binary Pattern*)

Les matrices de cooccurrence ont un voisin proche appelé le motif linéaire binaire. Plutôt que d'évaluer des couples de pixels, le LBP fait des statistiques sur le voisinage de chacun des pixels à un certain rayon de distance (voir figure 1.14). Il cherche ainsi à déterminer si les voisins du pixel central ont une intensité plus ou moins grande que celui-ci et à préciser leur positionnement. Chaque combinaison est représentée par une matrice dont la fréquence est décomptée sur toute l'image. Nous appliquons ensuite des statistiques à ces décomptes.

La différence de magnitude décrite entre l'intensité du pixel central et celle de ses voisins, à la figure 1.14D, peut également servir de descripteur et nous pouvons la recueillir de la même manière. Cette information contient moins d'aspects discriminants que l'analyse du signe de l'intensité par LBP, mais elle lui est complémentaire. Leur association est invariante en rotation et elle est appelée «motifs linéaires binaires complets» (CLBP ou *Complete Linear Binary Pattern*) (Guo et Zhang, 2010).

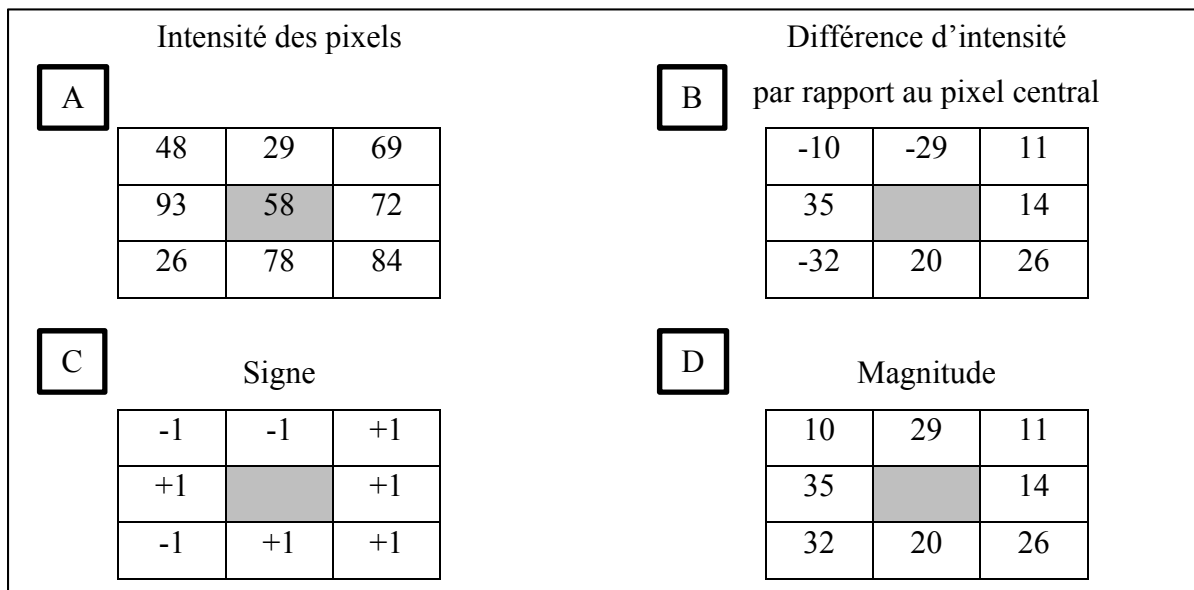


Figure 1.14 Représentation d'un A) pixel et de son voisinage pour un rayon de 2 pixels, de la B) différence de l'intensité relativement au pixel central, du C) LBP équivalent, exprimé par le signe de la différence et de D) la magnitude de la différence pour compléter CLBP

Adaptée de Guo et Zhang (2010)

1.3.5.3 Analyse de la distribution des pixels par filtrage

Bien que les méthodes statistiques puissent avoir un haut niveau d'efficacité, elles s'adaptent moins bien aux variations locales de l'image en texture et en illumination que les méthodes qui utilisent du filtrage (Whelan et Ghita, 2008). En conséquence, elles sont moins adéquates pour le traitement d'images naturelles, entre autres.

Les méthodes par filtrage utilisent la réponse $H(x,y)$ d'images $F(x,y)$ à des filtres $G(x,y)$ de différentes tailles et de différentes orientations pour déceler des textures. L'image est donc convoluée avec un filtre (éq. 1.10) comme l'illustre la figure 1.15. Les méthodes qui utilisent le filtrage diffèrent principalement par les types de filtres choisis et par le traitement subséquent des réponses au filtrage.

$$H(x,y) = (F * G)(x,y) = \sum_{y_1} \sum_{x_1} F(x,y)G^*(x - x_1, y - y_1) \quad (1.10)$$

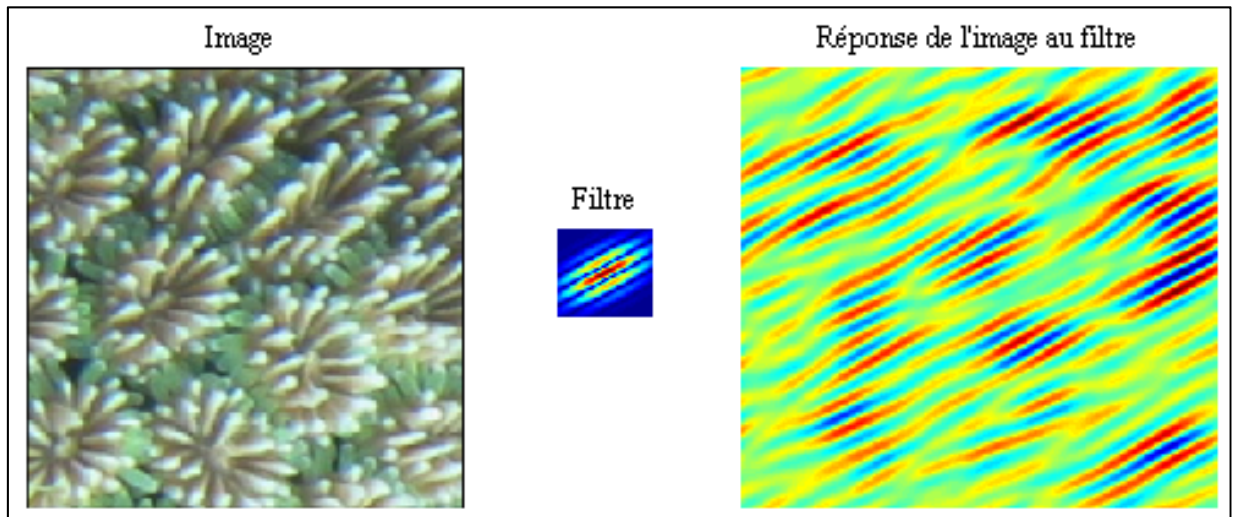


Figure 1.15 Filtrage d'une image par un filtre de Gabor
Photo reproduite et adaptée de l'AIMS (2013)

Réponse aux filtres de Gabor

Les filtres de Gabor sont parmi les plus connus et sont considérés comme robustes et efficaces (Naval et al., 2014). Ils sont formés d'une sinusoïde modulée par une fonction gaussienne. De manière à extraire de l'information sur les fréquences présentes dans la texture et sur l'orientation de celle-ci, chaque filtre se fait imposer plusieurs tailles et plusieurs orientations (Petrou et Sevilla, 2006). La convolution de chacun des filtres avec l'image construit donc une réponse qui révèle une information différente sur l'image. Pour l'extraire, la moyenne et l'écart-type des valeurs de chaque réponse sont utilisés comme étalon (Manjunath et Ma, 1996).

Autres méthodes

D'autres approches utilisent différentes banques de filtres et emploient leurs réponses autrement. Ainsi, l'algorithme des huit réponses maximales (MR8) agrège les réponses maximales de ses filtres à travers toutes les orientations qui leurs sont imposées (Varma et Zisserman, 2005). Cette démarche permet de réduire la dépendance du modèle à la rotation de la texture à travers les images. Le descripteur basé sur la représentation de Radon utilise, pour sa part, la transformée de Radon pour filtrer les textures (Liu, Lin et Yu, 2009). Ce faisant, en plus d'être robuste à la rotation et à l'échelle d'observation, la dépendance à l'illumination est réduite. Sur la base de données CURTeT, sa performance dépasse de 2 % celle de MR8. Toutefois, elle est inefficace sur les images dont le fond n'est pas texturé.

1.3.6 Extraction des descripteurs de couleur

Les images naturelles sont souvent riches en texture, mais aussi en couleur. Ces deux aspects sont intimement liés pour la reconnaissance de forme. On a montré à maintes reprises que d'exploiter les couleurs dans une image naturelle, en plus de la texture, améliore la performance des algorithmes de classification (Whelan et Ghita, 2008). Le temps supplémentaire associé à l'extraction de ces descripteurs est d'ailleurs minime en comparaison au temps pas-

sé sur les descripteurs de textures. Un des groupes de descripteurs les plus simples provient de l'extraction de statistiques à partir d'histogrammes formés sur les canaux de couleurs.

Une alternative performante est l'extraction de statistiques sur les histogrammes de teintes et d'angle opposé (Van De Weijer et Schmid, 2006). Cette démarche diminue l'impact des variations photométriques, telles que les ombres, sur la représentation des couleurs de l'image et elle ne dépend pas de la qualité de celle-ci. Par rapport à l'histogramme de couleurs, ces descripteurs sont donc invariants à l'illumination. Les histogrammes de teintes et d'angle opposé sont formés à partir des équations 1.15 et 1.16, suite à l'application des équations de 1.11 à 1.14 à l'ensemble des pixels. R , V et B représentent respectivement les canaux rouge, vert et bleu de chaque pixel de l'image. Les indices x et y indiquent, pour leur part, l'application d'un gradient à l'intensité de l'image selon les dimensions x et y .

$$O1 = \frac{1}{\sqrt{2}}(R - G) \quad (1.11)$$

$$O2 = \frac{1}{\sqrt{6}}(R + G - 2B) \quad (1.12)$$

$$O1_{x,y} = \frac{1}{\sqrt{2}}(R_{x,y} - G_{x,y}) \quad (1.13)$$

$$O2_{x,y} = \frac{1}{\sqrt{6}}(R_{x,y} + G_{x,y} - 2B_{x,y}) \quad (1.14)$$

$$Teinte = \arctan\left(\frac{O1}{O2}\right) \quad (1.15)$$

$$Angle\ opposé_{x,y} = \arctan\left(\frac{O1_{x,y}}{O2_{x,y}}\right) \quad (1.16)$$

Pour permettre une analyse différente, il est aussi possible de modifier l'espace de représentation de couleur de chaque pixel. La représentation RVB exploite un rapport d'intensité entre chacun de ses canaux pour illustrer une couleur. L'espace colorimétrique de teinte-saturation-intensité utilise une alternative : la teinte représente alors la couleur d'un pixel. La saturation, quant à elle, équivaut à la pureté de la couleur. Ainsi, une faible saturation apparaît grisâtre. Enfin, l'intensité est la quantité de lumière associée au pixel.

1.3.7 Dimension de l'espace de représentation

Selon la complexité d'un problème, le vecteur qui représente un échantillon peut contenir une grande quantité de descripteurs. Toutefois, ceux-ci ne sont pas tous utiles à une classification performante. Certains peuvent n'être efficaces que lorsqu'ils sont jumelés, ou d'autres, être source d'erreur dans la classification (Pal et Mitra, 2004). Ainsi, une bonne sélection des descripteurs peut créer un effet synergique sur la qualité de la classification.

De plus, un trop grand nombre de descripteurs, autrement dit de dimensions, a pour effet de diluer l'information pertinente : le fléau de la dimension provoque une stagnation, ou même une dégradation des performances lorsque l'espace des descripteurs atteint une certaine taille (Shih, 2010). Il faut donc réduire le nombre de descripteurs associés à chaque échantillon pour éviter ce phénomène. La diminution de la taille de l'ensemble permet également de réduire le temps de traitement des données. Deux approches sont alors possibles : regrouper les descripteurs ou en éliminer.

1.3.7.1 Analyse en composantes principales (ACP)

L'analyse en composantes principales réduit la dimension du vecteur en le projetant dans un nouvel espace. Ce dernier est supporté par des vecteurs décorrélés les uns par rapport aux autres (Shih, 2010). Ainsi, chacun des descripteurs est projeté dans un nouvel espace de manière à réduire la redondance des informations. Chacune des nouvelles valeurs contient donc une information maximalement différente de celle des autres.

1.3.7.2 Sélection d'attributs

La sélection d'attributs vise, pour sa part, à réduire la quantité de descripteurs en ne conservant que ceux qui contribuent le mieux à la performance de la classification. Les deux approches principales sont le filtrage et l'emballage (Pal et Mitra, 2004). Le filtrage est particulièrement prisé dans le cas où la classe associée aux attributs est inconnue. Son principe de base est de déterminer si un attribut contribue statistiquement à l'ensemble en fournissant suffisamment de renseignements nouveaux et pertinents pour départager les classes, ceci étant estimé par une fonction objective. Si ce n'est pas le cas, le descripteur est éliminé de l'ensemble.

L'emballage est plus efficace, mais il augmente significativement le temps de calcul. Contrairement au filtrage, qui n'utilise que la distribution générale des descripteurs, il optimise le choix des attributs en se basant sur les caractéristiques du classificateur qui sera utilisé et en utilisant directement les classes connues des échantillons.

L'algorithme de sélection de descripteurs basée sur la corrélation de (Hall, 1999) constitue un hybride de ces deux méthodes. En effet, il évalue les coefficients de corrélation de Pearson de sous-ensembles de descripteurs avec les classes, et aussi entre eux. Ainsi sont favorisés les sous-ensembles où l'information est peu redondante, mais fortement discriminante pour les classes. Sa performance est similaire à celle de l'emballage, mais il est plus rapide.

Nous pouvons d'ailleurs utiliser diverses méthodes pour sonder l'espace des descripteurs. La recherche exhaustive évalue toutes les combinaisons possibles d'attributs et sélectionne la plus performante. Sa complexité est toutefois exponentielle et le temps de calcul, prohibitif. Une autre possibilité est la recherche à l'aide d'un algorithme génétique, qui est basé sur la sélection naturelle. Dans ce cas, nous comparons itérativement la performance des ensembles de descripteurs. À chaque itération, les ensembles sont modifiés jusqu'à ce qu'un critère de performance soit atteint. Cette méthode est, encore une fois, coûteuse en temps. Pour sa part, l'algorithme glouton débute avec un ensemble de descripteurs de base et y fait des itérations

d'additions et de soustractions de descripteurs individuels, jusqu'à atteindre un maximum local de performance. Puisqu'il évalue moins de combinaisons, son temps de calcul est significativement plus faible que celui des autres algorithmes pour de grands ensembles de descripteurs, soit ceux de plus de cinquante instances (Kudo et Sklansky, 2000).

Si le but n'est toutefois pas de déterminer un sous-ensemble, l'algorithme d'évaluation du gain d'information permet de classer les descripteurs en ordre décroissant d'apport à la classification (Novakovic, 2009). Il est basé sur l'entropie (H) du système. Ainsi, le gain d'information d'un type de descripteur pour une classe précise dépend de l'entropie de la classe et de l'entropie de la classe sachant la valeur du descripteur, selon l'équation 1.17.

$$\begin{aligned} \text{Gain d'information}(\text{Classe}, \text{Descripteur}) \\ = H(\text{Classe}) - H(\text{Classe} | \text{Descripteur}) \end{aligned} \quad (1.17)$$

1.3.8 Classificateurs

Les classificateurs sont les techniques qui servent à regrouper les objets qui appartiennent à une même classe. Ils synthétisent les informations disponibles sur les objets et sur les classes pour effectuer leurs décisions. Un cas particulier est celui où les classes sont inconnues. Dans ces circonstances, les algorithmes servent donc à séparer les objets selon leur réponse à une fonction objective. L'algorithme K-moyennes utilise la moyenne des valeurs des objets appartenant à une classe pour les répartir (Shih, 2010). Un objet appartient à une classe si sa valeur se situe à l'intérieur d'une distance maximale de la moyenne des valeurs des objets de cette classe. Les classes sont donc formées selon le critère de ressemblance des objets.

Lorsque les classes sont connues, nous pouvons lier à chacune d'elle une banque de caractéristiques pour y associer les objets inconnus. Par exemple, les vecteurs de descripteurs des objets de cette classe peuvent tout simplement en constituer les caractéristiques. L'algorithme des k plus proches voisins (KNN ou *k-nearest neighbors*) assigne des étiquettes aux objets selon la similarité des descripteurs des classes, comme le montre la figure 1.16 (Duin et Tax,

2005). Cet algorithme nécessite l'évaluation de la distance entre l'objet recherché et chacun des éléments du groupe d'entraînement à chaque classification, ce qui peut s'avérer couteux en temps.

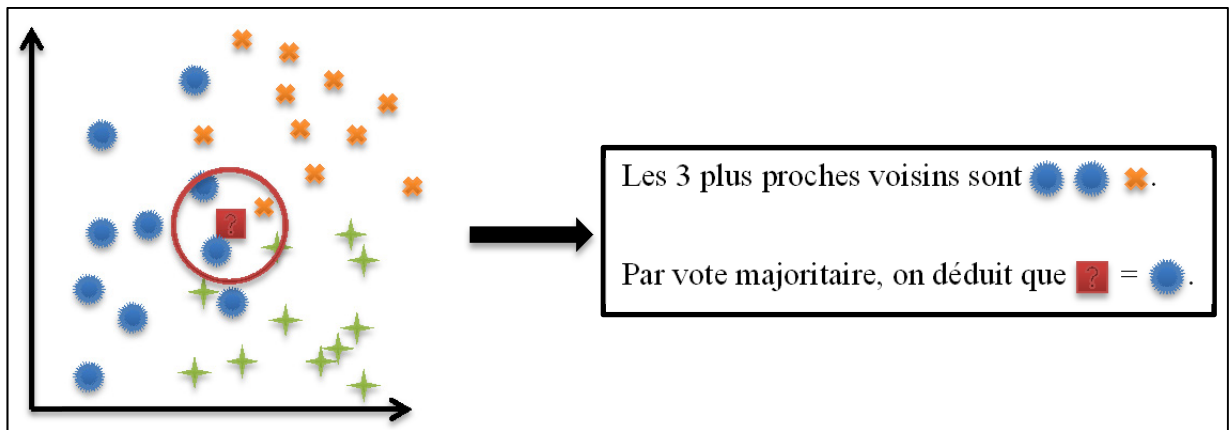


Figure 1.16 Algorithme des k plus proches voisins

Le classificateur de Bayes utilise la distribution des données pour évaluer la probabilité p d'appartenance d'un échantillon x à chaque classe w_i (Shih, 2010). L'échantillon est ensuite associé à la classe qui est la plus probable, selon l'équation 1.18, dans le but de minimiser les erreurs de classification.

$$p(w_i|x) = \frac{p(x|w_i)p(w_i)}{p(x)} \quad (1.18)$$

Les réseaux de neurones, pour leur part, utilisent une structure souvent cachée, constituée de neurones, pour faire des classifications. Ces neurones sont organisés en couches qui effectuent diverses opérations sur leurs intrants (*Voir* figure 1.17). La réponse de chacun des neurones dépend donc des informations qu'il a reçues de la couche précédente. L'entraînement d'un réseau de neurones a pour but d'assigner un poids aux liaisons entre les neurones. Le réseau opère ensuite simplement en mettant le groupe test dans les intrants. La réponse du réseau de neurones ressort ensuite par les neurones extrants.

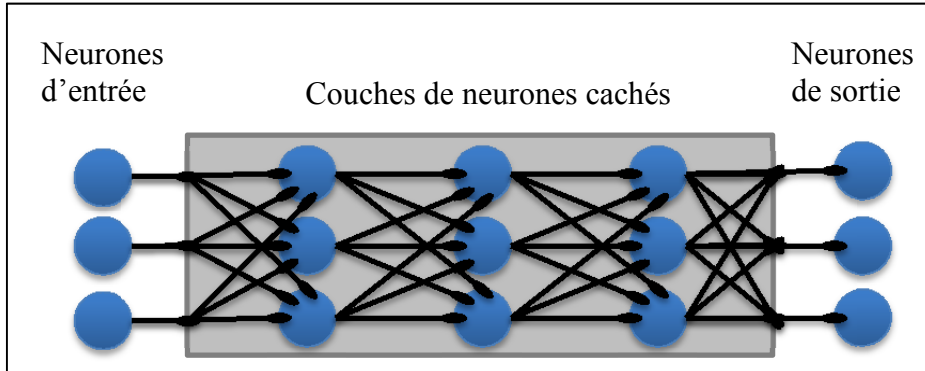


Figure 1.17 Représentation d'un réseau de neurones

Enfin, la machine à vecteurs de support (SVM) utilise le groupe d'entraînement pour créer des frontières entre les classes, ce qui est présenté à la figure 1.18. Ainsi, une droite est positionnée dans le plan de manière à séparer aussi efficacement que possible deux classes, en maximisant la distance entre la droite et les points pour éviter qu'elle ne soit trop proche d'une des deux distributions. La droite doit aussi minimiser l'impact des erreurs lorsqu'il est impossible de séparer parfaitement les classes.

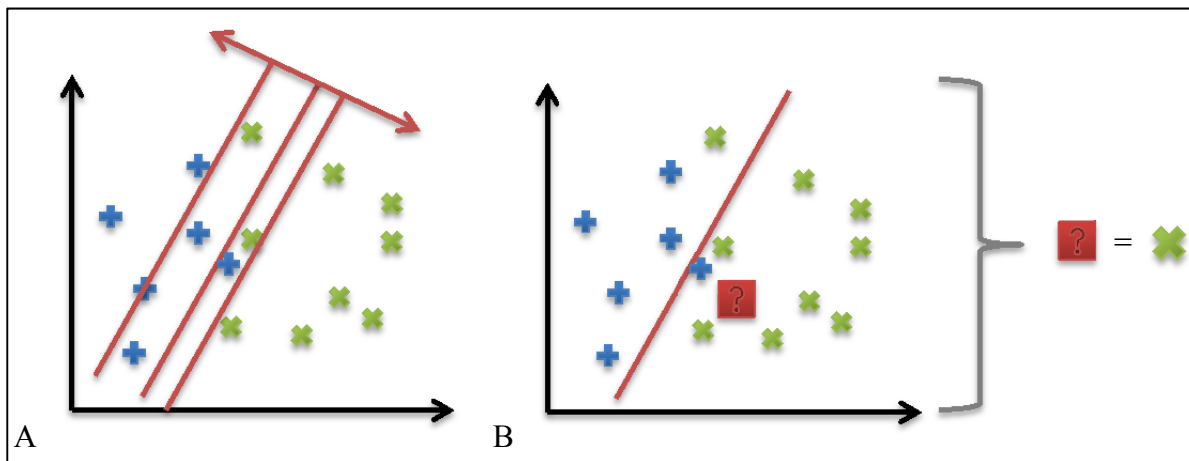


Figure 1.18 Représentation de A) l'entraînement d'un SVM et de B) la classification d'un élément à partir de ce SVM

À l'origine, les SVM séparaient les données à partir d'hyperplans linéaires. Cette méthode n'est cependant valable que pour les bases de données séparables linéairement. On a donc créé des alternatives pour séparer les bases de données dans lesquelles la dépendance entre

les attributs et les classes des objets est non-linéaire. Pour ce faire, les attributs sont projetés dans de nouveaux espaces, dans lesquels ils sont séparables linéairement, à l'aide de noyaux. La méthode classique pour ce type d'algorithme est le classificateur à vecteurs de support régularisé (C-SVC ou *C-Support Vector Classification*) (Novakovic et Veljovic, 2011).

Quatre noyaux sont régulièrement utilisés pour faire les projections (Novakovic et Veljovic, 2011). Le premier est le noyau linéaire, surtout performant dans les situations linéairement séparables, mais aussi dans les cas non-linéaires. Le second noyau est polynomial. Sa performance est similaire à celle du noyau linéaire et elle dépend du degré du polynôme exploité. Le troisième noyau est une fonction en base radiale (RBF). Sa conception fait que la projection des attributs est effectuée de manière non-linéaire. Il est donc performant dans ces cas. Toutefois, deux de ses paramètres, γ et C , doivent être optimisés, sans quoi le SVM final dépend fortement des données d'entraînement et donne lieu à un surapprentissage qui engendre des erreurs de classification. Il est donc beaucoup plus complexe à mettre en place que les noyaux linéaires et polynomiaux. Notons par ailleurs que le noyau linéaire est un cas particulier du RBF et que le modèle final de celui-ci est moins complexe à utiliser que celui du noyau polynomial. Le dernier noyau est le noyau sigmoïdal, souvent moins efficace que les autres puisque certains paramètres le rendent indéfini. Exigeant aussi l'optimisation de paramètres, il est peu utilisé.

Nous pouvons également généraliser cette démarche à plus de deux classes. Pour ce faire, nous devons générer un SVM pour chaque couple de classe, dans un format «un contre un» ou «un contre tous». Dans ce dernier cas, le format compare une classe à toutes les autres grâce à un hyperplan. Lorsque les SVM ont été entraînés, la classification des objets de test se fait en trouvant le SVM unique pour lequel la réponse d'appartenance à la classe est positive. Le format «un contre un» compare deux classes à la fois. Un SVM est donc entraîné pour chaque couple de classes. Lors du test, les SVM votent pour savoir laquelle des classes est la plus probable. Bien que performante, cette démarche est lente, car chaque SVM doit être évalué pour chacun des échantillons de test. Par contre, nous pouvons réduire le nombre de SVM en les organisant en arbre. Conséquemment, la réponse de chaque SVM détermine

quel doit être le prochain SVM évalué. Ces méthodes d'évaluation en série sont généralement performantes, cependant le résultat du SVM multiclasse dépend de chaque SVM individuel.

Nous pouvons également avoir recours à la notion de classification floue (Bankman, 2009). Elle repose sur l'idée qu'il n'est pas toujours nécessaire de choisir une classe en particulier, mais qu'il est utile d'évaluer la ressemblance à chacune d'elles. Si la similitude entre un échantillon et toutes les classes est trop faible, l'échantillon peut être rejeté, ou des tests supplémentaires ajoutés pour compléter la démarche de classification.

Le domaine médical, qui est similaire à celui de la classification de coraux, utilise souvent les SVM, car ils ont une grande capacité de généralisation. En effet, les hyperplans qui divisent les données sont choisis pour maximiser la distance entre eux et chaque classes. Plus cette distance est grande, plus le SVM pourra, par la suite, être efficace sur de nouvelles données. Il minimise donc les erreurs tant au niveau du groupe d'entraînement que du groupe test (Novakovic et Veljovic, 2011). Le SVM est également moins sensible au déséquilibre des classes que la plupart des autres algorithmes (Anand et al., 2010) et ses paramètres peuvent être adaptés aux particularités de la base de données (Jiang, Missoum et Chen, 2014; Sun et al., 2007; Yu et al., 2015). Dans le cas particulier de la classification de coraux, le noyau RBF serait le plus efficace (Yu et al., 2015).

1.3.9 Analyse statistique des résultats

Le processus de classification peut être évalué à l'aide de différents indicateurs, tels que l'exactitude, le taux de rappel, la précision, le *F-mesure* ou encore le *G-mean*. Toutefois, outre ces outils, il est important d'avoir une représentation statistique de la performance. Pour ce faire, il existe diverses méthodes d'exploitation, de représentation et de comparaison des données.

Examinons d'abord la validation croisée, démarche qui permet des évaluations statistiques de la qualité d'un système. Dans le cas particulier d'une classification d'objets, elle indique la

capacité du système à faire des classifications sur une base de données précise. Cette dernière est donc divisée aléatoirement en n groupes uniformes. L'un sert alors de groupe test, alors que les $(n - 1)$ autres forment le groupe d'entraînement pour la classification. Nous effectuons ensuite une rotation à l'intérieur des groupes jusqu'à ce que chacun d'entre eux ait été le groupe test. Une moyenne sur les résultats de classification de chacun des cas est enfin effectuée. Ces valeurs représentent la capacité du système à classer cette base de données. Typiquement, la validation croisée utilise dix groupes.

La boîte à moustache (*Voir* figure 1.19) illustre, pour sa part, une série de données en délimitant clairement la médiane, les quartiles et les valeurs extrêmes. Elle facilite la visualisation de distributions et permet donc d'en comparer plusieurs. Le test de Student, ou test t , sert quant à lui à déterminer mathématiquement si deux distributions sont significativement différentes l'une de l'autre. Pour ce faire, il compare la distance entre leurs moyennes respectives. L'erreur (α) typique de ce test est de 1 % ou 5 %. Les distributions peuvent également être comparées entre elles en évaluant leur corrélation. Cette démarche vise à déterminer si les valeurs que prend un premier groupe ont une influence sur la prise de valeur d'un second groupe. Le coefficient de corrélation de Pearson donne le degré de similitude de ce lien avec la relation linéaire. La plage de valeurs de ce coefficient s'étend entre -1 et +1, où 1 correspond à une corrélation parfaite et -1 à une relation négative. 0 indique qu'aucun lien linéaire n'existe entre les deux distributions. Des valeurs supérieures à 0,7 signifient qu'il y a une forte ressemblance entre les distributions.

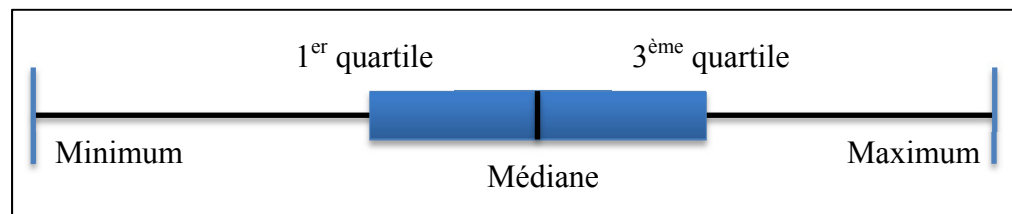


Figure 1.19 Représentation d'une boîte à moustache

1.3.10 Application à la classification de groupes benthiques

La recherche sur la classification d'images de récifs coralliens a permis d'identifier plusieurs facteurs supplémentaires qui influencent la performance des systèmes. Ainsi, la plupart des algorithmes sont basés sur des images sous-marines qui ont été prises selon les canaux rouges, verts et bleus. Ils exploitent donc tant la texture que la couleur des objets identifiés pour en faire la classification, en imposant parfois un poids significativement plus grand aux descripteurs de texture que de couleur pour les favoriser lors du traitement. Lors du choix des classes, il est aussi nécessaire de faire un compromis entre le nombre d'échantillons dans chaque classe étudiée et le nombre de classes. Effectivement, regrouper des classes permet d'obtenir plus d'échantillons par classe, mais cela élargit aussi la distribution des descripteurs à l'intérieur d'une même classe (Pizarro et al., 2008; Stokes et Deane, 2009).

L'impact de la construction du modèle est également déterminant. Ainsi, avoir des échantillons représentatifs de chaque type de benthos est un avantage. La qualité et l'uniformité des images sont aussi des facteurs qui influencent la classification de manière importante (Stokes et Deane, 2009). En outre, l'extraction de descripteurs sur les canaux de couleur pris isolément, selon diverses échelles de filtre ou diverses tailles d'images segmentées s'avère un choix judicieux pour améliorer la classification, bien que cela augmente significativement la taille des vecteurs de descripteurs (Beijbom et al., 2012).

Les descripteurs de texture tels que les filtres de Gabor (Naval et al., 2014; Pizarro et al., 2008), de MR8 (Beijbom et al., 2012; Litimco et al., 2013), LBP et CLBP (Marcos et al., 2008; Stokes et Deane, 2009) sont les plus fréquemment exploités et les plus performants. Pour la couleur, les histogrammes (Marcos et al., 2008; Stokes et Deane, 2009) sont souvent utilisés, tout comme, pour les classificateurs, le SVM (Beijbom et al., 2012), KNN (Stokes et Deane, 2009) et les réseaux de neurones (Marcos et al., 2008). Le taux de classification oscille entre 70 % et 80 % selon les bases de données et le nombre de classes utilisées. Par exemple, (Beijbom et al., 2012) atteint 74 % pour 9 classes sur MLC à partir de MR8 et

d'une analyse multidimensionnelle. Cependant, puisque chacune de ces techniques est testée sur des bases de données différentes, leur robustesse est indéfinie.

Une mise en commun de ces techniques sur six bases de données a été faite par (Shihavuddin et al., 2013), à des fins de comparaison. Ces optimisations ont ainsi atteint 85,5 % de taux de classification sur MLC. Les combinaisons les plus performantes sont présentées au tableau 1.1.

Tableau 1.1 Combinaisons optimales de techniques pour chaque étape de la classification selon (Shihavuddin et al., 2013)

Étape	Technique	Cas d'application
Prétraitement	CLAHS	Faible contraste et flou dans les images
Segmentation	Imagettes carrées centrées sur le point étiqueté	
Extraction des descripteurs	Matrices de cooccurrence des teintes de gris Réponses des filtres de Gabor CLBP Histogrammes de teinte et d'angle opposé	
Réduction de la dimension de l'espace de représentation	ACP	Moins de 5 000 échantillons
	Aucun	Plus de 12 000 échantillons
Classification	KNN	Moins de 5 000 échantillons
	SVM ou Réseau de neurones	Plus de 12 000 échantillons

Une autre initiative d'optimisation est menée par (Litimco et al., 2013). Ils ont proposé l'intégration au système d'une application Android, de manière à améliorer l'acquisition de données et l'accessibilité d'un algorithme simple, par segmentation simple et l'utilisation de MR8. Cette application sert de passerelle entre les experts en identification de coraux et les

fournisseurs d'images qui y téléchargent leurs clichés. Les experts, selon leur certitude, peuvent alors classer les images transmises dans les catégories acropora, porites, autres coraux ou non coraux. Les données sont subséquemment utilisées pour constituer le groupe d'entraînement de l'algorithme.

CHAPITRE 2

OPTIMISATION D'UN ALGORITHME DE CLASSIFICATION POUR LA BASE DE DONNÉES DE L'AIMS

2.1 Introduction

La base de données de l'AIMS a des caractéristiques propres. Nous devons sélectionner les techniques de traitement d'images de manière à maximiser la capacité du système à les classer. Chacune des étapes du processus doit donc être analysée. Le système optimisé servira ensuite à l'élaboration d'une stratégie de sélection de données d'entraînement pour la classification de nouvelles données pour l'AIMS (CHAPITRE 3).

2.2 Méthodologie d'optimisation

Afin d'optimiser le processus de classification, nous fixons tout d'abord des paramètres de base. Le tableau 2.1 résume ceux issus de l'optimisation de (Shihavuddin et al., 2013), dont la performance et la simplicité d'application justifient l'utilisation comme référence. Nous itérons ensuite sur le prétraitement, sur la segmentation, sur l'extraction de descripteurs, sur la réduction de dimension de l'espace des descripteurs et sur le classificateur. Il devient ensuite aisé de mesurer l'impact de chaque étape sur le taux de classification des données.

Le scénario de base consiste alors à faire des validations croisées à quatre ensembles, sur les données d'une période d'échantillonnage et d'un récif uniques, soit le récif Martin pour 2012-2013. Ce dernier démontre une performance moyenne par rapport aux autres cas et il contient 3000 points identifiés, répartis en 1422 algues, 197 coraux mous, 905 coraux durs, 413 abiotiques, 26 éponges, 1 millépore et 36 autres. Nous négligeons cependant les éponges, les millépores et les autres, vu leur faible population. La validation croisée contient, pour sa part, quatre ensembles plutôt que dix afin d'accélérer le traitement des données. Ce choix augmente l'incertitude sur la performance réelle du système, mais il permet tout de même

d’observer la tendance de son évolution. Pour compenser, nous répétons la validation croisée de manière à en estimer l’incertitude et à mesurer sa robustesse aux données d’entraînement.

Tableau 2.1 Paramètres de chaque étape de la classification pour le scénario de base

Étape	Technique	Paramètres
Prétraitement	Aucun	
Segmentation	Imagettes carrées centrées sur le point étiqueté	350 pixels de côté
Extraction des descripteurs	Matrices de cooccurrence des teintes de gris	22 statistiques (<i>Voir ANNEXE I, p.97</i>) pour 16 niveaux de gris et un rayon de 3 pixels
	Réponses aux filtres de Gabor	Moyenne, variance, asymétrie et kurtosis pour 4 tailles et 6 orientations de filtres
	CLBP	Fenêtres de 8, 16 et 24 pixels
	Histogrammes de teinte et d’angle opposé	36 plages chacun
	Intégration de la transformée de Fourier	50 plages
Réduction de la dimension de l’espace de représentation	Aucun	
Classification	SVM	«Un contre un», noyau RBF

2.3 Sélection des prétraitements

2.3.1 Caractéristiques des images

L'AIMS a collecté les images de la base de données à l'aide d'une caméra sous-marine selon un protocole normalisé. Les clichés ont conséquemment une haute résolution et une représentation uniforme. En effet, il faut que les images soient similaires pour permettre au système de les comparer. Les prétraitements doivent donc normaliser les caractéristiques des images, telles que le niveau de bruit, la couleur de la lumière d'illumination ainsi que la distribution de cette illumination dans l'image.

2.3.2 Méthode de comparaison des prétraitements

Nous utilisons, pour commencer, un cas sans prétraitement comme référence et nous intégrons ensuite divers prétraitements au système de classification. Nous comparons donc les taux de rappel résultants par un test de Student, suite à cinq répétitions de chaque cas. Nous reprenons ensuite la même démarche sur un total de dix récifs pour évaluer l'effet de la diversité des données.

Le premier prétraitement est un filtre médian et il vise à réduire le bruit dans les images. Les autres rectifient l'éclairage et optimisent la plage dynamique des couleurs. Nous les appliquons du plus simple au plus complexe : la balance des blancs avec un allongement du contraste, la CCN, une CLAHS et l'ICM-R. Un exemple typique de leurs effets est présenté à la figure 2.1.

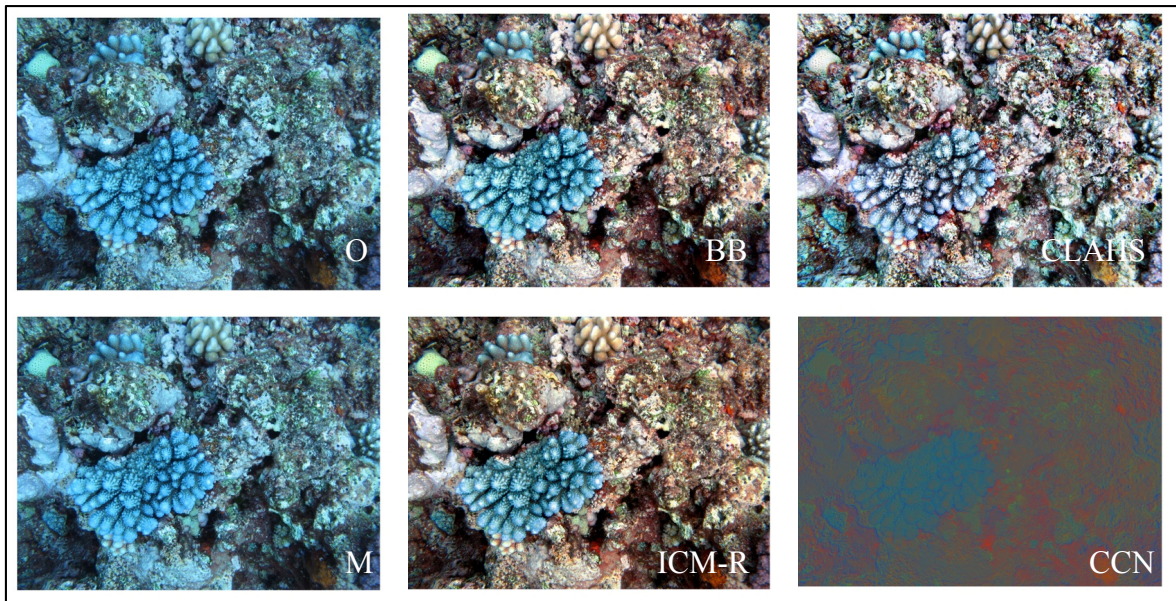


Figure 2.1 Application de la balance des blancs (BB), de CLAHS, d'un filtre médian (M), de ICM-R et de CCN à une image originale (O)
Photo reproduite et adaptée de l'AIMS (2013)

2.3.3 Résultats de la comparaison des prétraitements

La figure 2.2 illustre les taux de classification obtenus pour chacun des prétraitements évalués. Aucune des méthodes n'obtient de taux de classification significativement plus élevé que lorsqu'aucun prétraitement n'est mis en place. Cette conclusion est confirmée par des tests de Student, pour une précision $\alpha = 0,05$. La tendance est conservée sur l'ensemble des récifs testés.

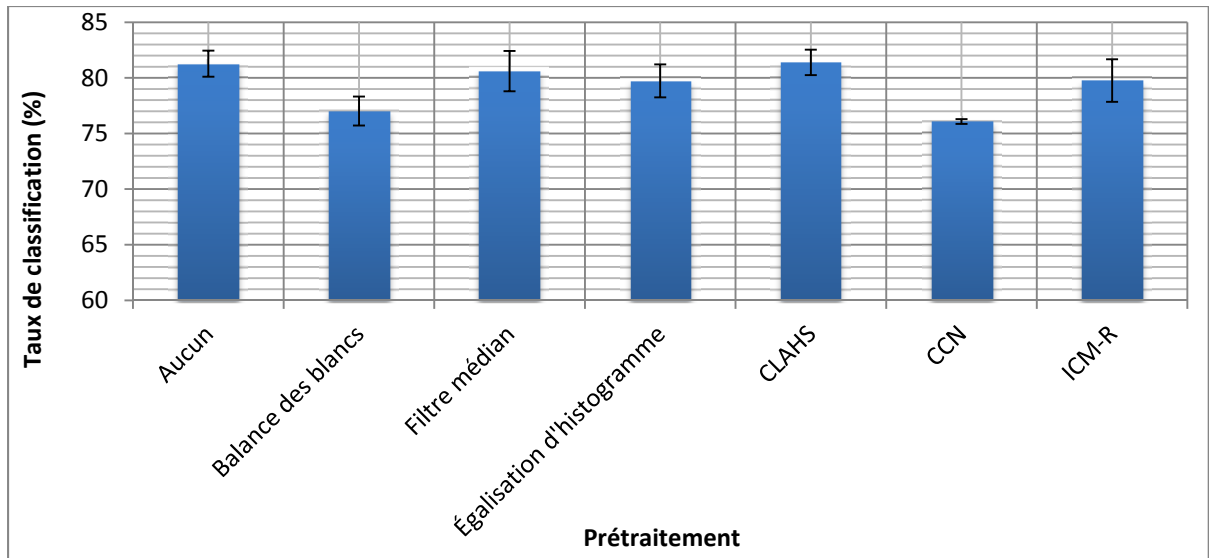


Figure 2.2 Taux de classification associés à divers prétraitements

2.3.4 Interprétation des conséquences du prétraitement

Les prétraitements conservent ou réduisent la performance sur la base de données de l'AIMS, mais ils ne l'améliorent pas. Conséquemment, ces images sont homogènes et les prétraitements n'éliminent pas suffisamment d'artéfacts pour justifier de les conserver, d'autant plus qu'ils sont coûteux en temps. Ces résultats sont cohérents avec les travaux de (Shihavuddin et al., 2013), qui a observé une amélioration de 1,2 % à 15,7 % de ses résultats lors de l'utilisation de l'égalisation d'histogramme ou d'une combinaison de techniques similaire à l'ICM-R. Ainsi, le comportement du système suivant un prétraitement ne varie pas nécessairement de manière significative selon les bases de données d'images sous-marines et colorées utilisées. La base de données de l'AIMS appartient à la catégorie qui ne nécessite pas cette étape, grâce à sa grande qualité.

2.4 Segmentation

2.4.1 Choix du type de segmentation

La segmentation efficace d'images de coraux est un problème complexe. En effet, les structures sont difficiles à délimiter. Conséquemment, nous découpons généralement de simples carrés autour du pixel étiqueté. La taille de ceux-ci doit être sélectionnée judicieusement puisque la performance de la reconnaissance de forme dépend des textures qui y sont incluses. Dans un cas idéal, seul l'objet d'intérêt y est. Néanmoins, cet objectif est irréalisable sans un processus de segmentation plus élaboré. Nous choisissons donc la taille de l'imagette de manière à ce que la majorité de l'espace soit occupée par l'objet d'intérêt et que suffisamment d'information y soit disponible. La mesure est ainsi caractéristique des clichés pris à partir du protocole normalisé de l'AIMS et elle offre une représentation statistique de l'ensemble de la base de données.

2.4.2 Méthode de sélection de la taille de l'imagette

Nous effectuons des itérations par pas de 25 pixels, sur une plage de 50 à 800 pixels, afin de déterminer la longueur optimale du côté des imagettes. Puis nous classifions les objets et nous comparons leurs taux de rappel pour chaque taille testée. Nous combinons ensuite les vecteurs de descripteurs de diverses dimensions et, ce faisant, nous évaluons s'ils contiennent plus d'information que ceux des imagettes uniques.

2.4.3 Résultats de la sélection de la taille de l'imagette

La figure 2.3 présente la courbe de performance du système selon la longueur du côté des imagettes. Les barres d'erreurs représentent la variabilité des résultats obtenus lors de la répétition des validations croisées. La figure 2.4 montre, pour sa part, les résultats de quelques combinaisons de tailles d'imagettes dont les vecteurs de descripteurs sont concaténés. La performance du système pour une imagette de 350 pixels de côté y est incluse à des fins de comparaison.

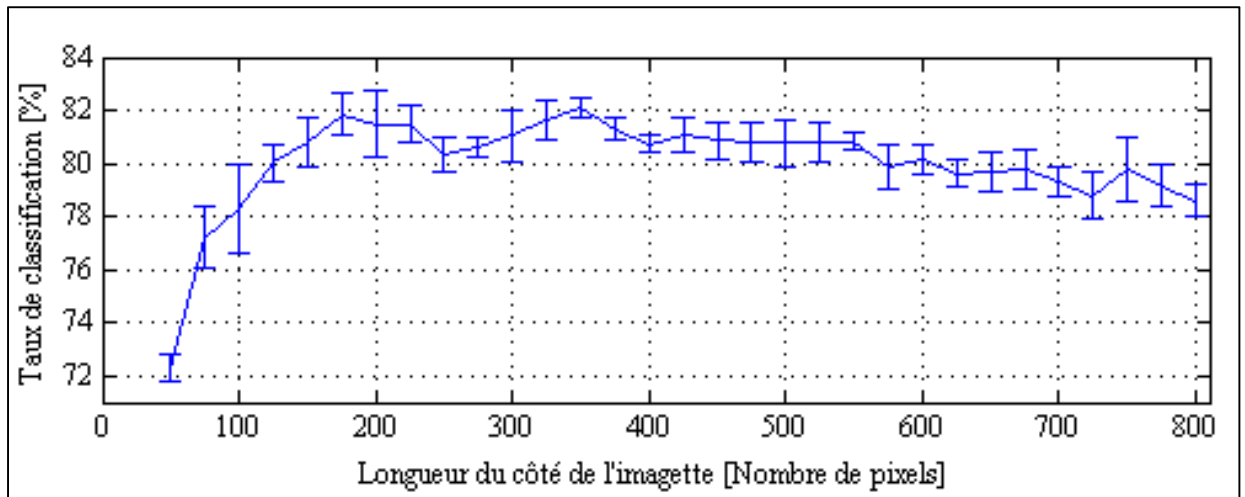


Figure 2.3 Taux de classification associé à la taille des imagerie testées

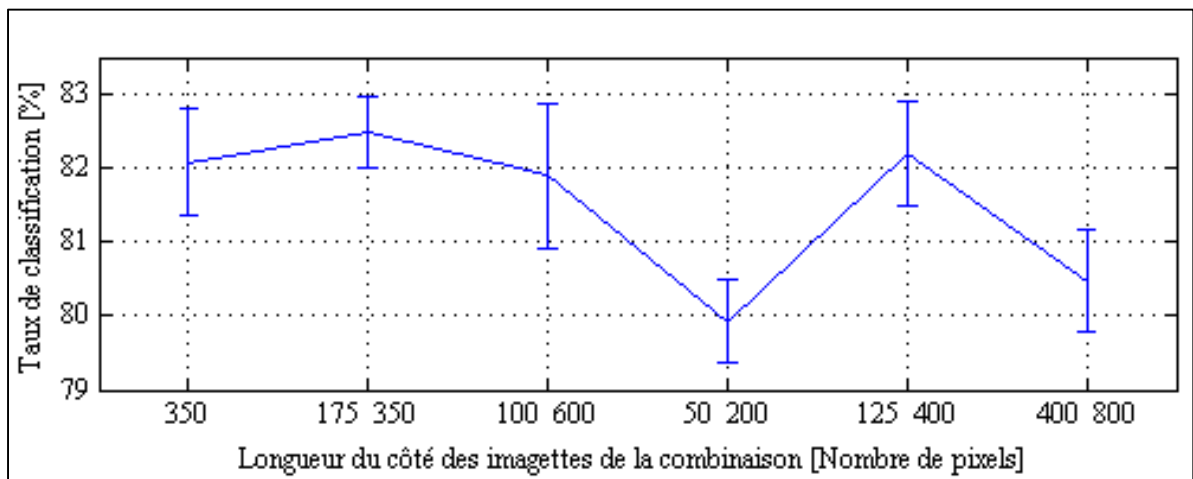


Figure 2.4 Taux de classification associé à chaque combinaison testée d'imagerie de tailles différentes

2.4.4 Impact du choix de taille d'imagerie

La quasi-totalité de la plage d'essai offre des résultats équivalents. Il n'est donc pas nécessaire d'inclure des combinaisons de tailles puisqu'elles ne sont pas significativement plus efficaces que les tailles uniques. Seules les longueurs de côté inférieures à 150 pixels sont à proscrire : elles n'incluent pas l'information nécessaire à la classification. Par contre,

l'inclusion d'une trop grande surface de l'image nuit également à la performance lorsque la longueur du côté est supérieure à 400 pixels. Il est donc finalement préférable d'utiliser une mesure de 350 pixels pour former les carrés puisqu'elle appartient à la plage du maximum de performance. Ce résultat est caractéristique à la base de données de l'AIMS et ne pourrait pas être extrapolé à une autre base de données. En effet, la taille des objets sur l'image pourrait varier dans une autre base.

2.5 Sélection des descripteurs

2.5.1 Définition de l'impact du choix des descripteurs

Le choix des descripteurs influence la performance de l'entièreté de la classification. Le système ne peut pas distinguer les classes si les descripteurs ne sont pas adaptés au type d'images traitées. Nous devons mettre en place différents mécanismes pour extraire l'information discriminante, puisque les coraux sont très colorés et que leurs textures varient selon l'échelle à laquelle ils sont observés. De manière à maximiser les résultats, nous implémentons des descripteurs reconnus dans la littérature et, par la suite, nous éliminons expérimentalement ceux qui sont redondants et inutiles pour ne conserver que ceux qui améliorent la classification par SVM en groupes benthiques.

2.5.2 Méthodologie de sélection des descripteurs

La liste contenue dans le tableau 2.2 présente l'ensemble des descripteurs considérés lors de l'extraction. Au point de départ, nous appliquons ces descripteurs de texture à une image en tons de gris et pour une seule échelle de l'image. Cependant, pour inclure des informations de couleur supplémentaires, nous extrayons également chacun des descripteurs selon les canaux de rouge, de vert, de bleu, de teinte et de saturation. De plus, nous transformons l'image par une pyramide gaussienne à trois couches, incluant l'image originale, et par une pyramide laplacienne à deux couches. En extrayant les descripteurs pour chacune des nouvelles échelles, nous ajoutons la notion de profondeur de la texture. Nous normalisons enfin l'ensemble des chiffres entre zéro et un, puis nous le concaténons dans un vecteur.

Tableau 2.2 Ensemble des descripteurs de texture et de couleur considérés

Technique	Paramètres	Dimensions évaluées
Matrices de cooccurrence des teintes de gris	22 statistiques (ANNEXE I, p.97) pour 16 niveaux de gris, 4 orientations et des rayons de 1, 3, 5 et 10 pixels	5 pyramides x 6 canaux de couleurs
	352 valeurs	
Réponses aux filtres de Gabor	Moyenne, variance, asymétrie et kurtosis pour 4 tailles et 6 orientations de filtres	5 pyramides x 6 canaux de couleurs
	96 valeurs	
CLBP	Fenêtres de 8, 16 et 24 pixels	5 pyramides x 6 canaux de couleurs
	108 valeurs	
Histogrammes de teinte et d'angle opposé	37 plages	5 pyramides
	74 valeurs	
Intégration de la transformée de Fourier	50 plages	5 pyramides x 6 canaux de couleurs
	50 valeurs	
Statistiques de l'histogramme de tons de gris	7 statistiques (<i>Voir</i> ANNEXE II, p. 101)	5 pyramides x 6 canaux de couleurs
	7 valeurs	

La totalité des descripteurs, pour une seule image, est donc constituée de 18 760 valeurs. Certaines d'entre elles sont toutefois redondantes ou inutiles. Nous devons donc sélectionner celles qui sont pertinentes pour la classification des groupes benthiques. Outre le besoin d'avoir un sous-ensemble bien constitué, le temps de traitement devient la contrainte principale dans le choix de la technique de sélection, vu la taille importante du vecteur. Nous analysons donc celui-ci grâce au logiciel Weka (version 3.7.11) (Machine Learning Group at the University of Waikato, 2014) étant donné sa disponibilité et la simplicité de son interface.

L'évaluation des descripteurs s'effectue par un filtrage basé sur la corrélation et un algorithme glouton sonde l'espace. Ces choix offrent un compromis raisonnable entre la qualité des résultats et le temps de calcul. Nous intégrons ensuite un peaufinage expérimental de l'ensemble de descripteurs sélectionnés. Pour ce faire, un second filtre, basé sur le gain d'information, indique la qualité de l'information fournie par chaque valeur et il les ordonne en fonction de leur capacité à distinguer les classes.

Martin 2012-2013, soit un récif unique, à une période d'échantillonnage précise, sert de modèle pour la sélection initiale des descripteurs. Puis, nous reprenons la sélection sur huit récifs différents pour donner un aperçu de la variabilité des résultats. Nous comparons la distribution des descripteurs par un coefficient de corrélation pour connaître leur degré de ressemblance. Nous classifions en outre les récifs à partir des descripteurs sélectionnés sur Martin 2012-2013. De cette manière, nous déduisons s'ils causent une dégradation de la performance par rapport aux ensembles des descripteurs sélectionnés sur les récifs eux-mêmes. Cette démarche nous indique si nous pouvons généraliser l'utilisation de cet ensemble au reste de la base de données.

Enfin, nous confrontons les taux de classification de l'ensemble réduit de descripteurs à ceux des descripteurs de (Shihavuddin et al., 2013), qui correspondent à la fine pointe de la technique, et à une réduction de la dimensionnalité par ACP, qui est la méthode de réduction la plus répandue. Ce faisant, nous mesurons la qualité du groupe par rapport aux autres méthodes.

2.5.3 Distribution des descripteurs sélectionnés à travers les récifs

Les figures 2.5, 2.6 et 2.7 illustrent la distribution des valeurs choisies par les algorithmes sur neuf récifs. Les barres d'erreur représentent la disparité des ensembles selon les récifs auxquels ils appartiennent. Le nombre d'instances sélectionnées par l'algorithme parmi les 18 760 initiales varie autour de 264 et leur analyse permet d'évaluer les combinaisons de descripteurs les plus fréquentes. Il apparaît alors que l'ensemble est hétérogène et que chaque

modalité est essentielle pour le constituer. Aucune échelle, couleur ou technique d'extraction n'est exclue, bien que certaines soient plus fréquemment employées que d'autres.

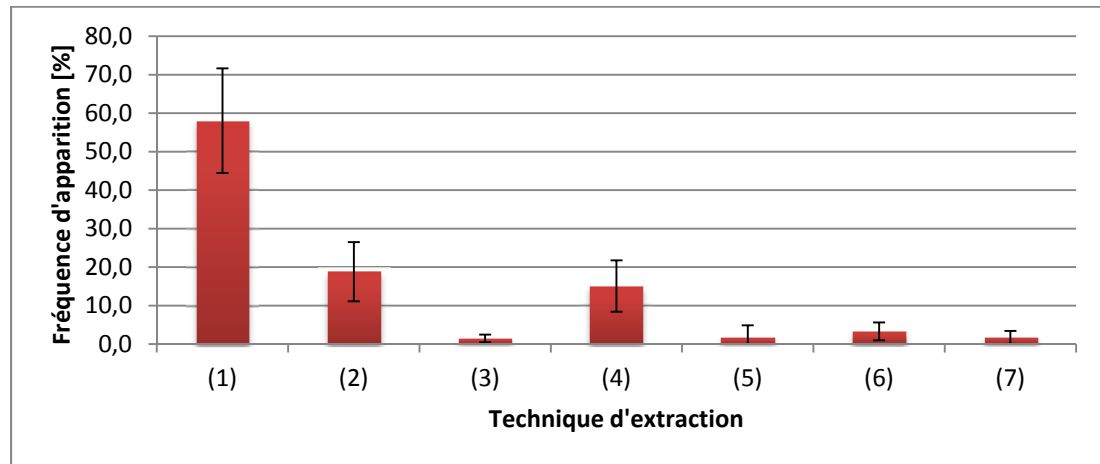


Figure 2.5 Fréquence relative d'occurrence de chaque technique d'extraction dans le groupe de descripteurs sélectionnés. Légende : (1) CLBP, (2) réponse aux filtres de Gabor, (3) statistiques sur l'histogramme, (4) matrice de cooccurrence, (5) intégration de la transformée de Fourier, (6) histogramme de teintes, (7) histogramme d'angle opposé

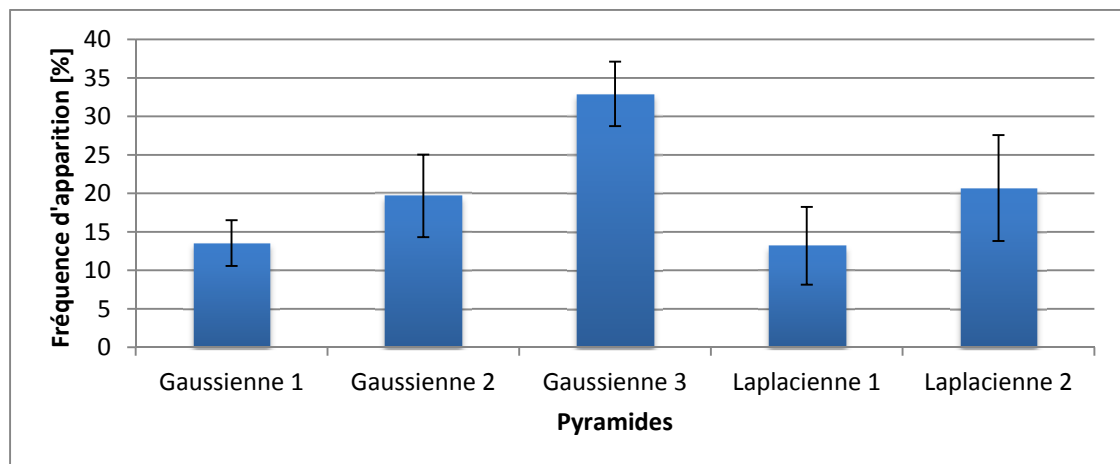


Figure 2.6 Fréquence relative d'occurrence de chaque échelle des pyramides gaussiennes et laplaciennes dans les descripteurs sélectionnés où la pyramide gaussienne 1 est l'image originale, à haute résolution

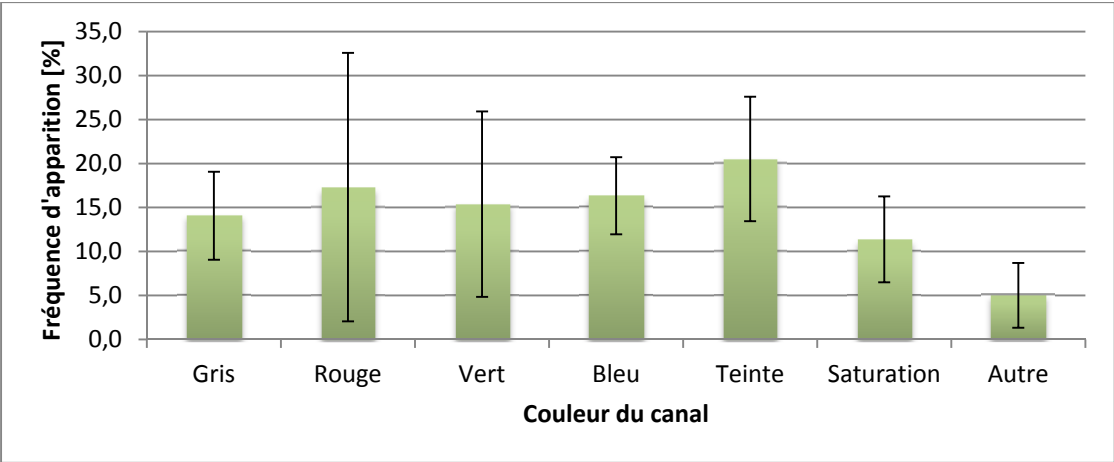


Figure 2.7 Fréquence relative d’occurrence de chaque dimension de couleur dans les attributs sélectionnés. Notez que «Autres» englobe les attributs qui ne sont pas associés à une couleur en particulier, comme les histogrammes de teinte et d’angle opposé

CLBP et les filtres de Gabor représentent 75 % de toutes les valeurs sélectionnées. L’échelle la plus utile est, pour sa part, la troisième couche, à basse résolution, de la pyramide gaussienne, alors que les couleurs apparaissent toutes approximativement à la même fréquence dans le vecteur. De plus, la composition des ensembles sur chacun des récifs est toujours corrélée au-delà d’un facteur 0,84 avec celle du récif Martin 2012-2013. Elles sont donc très similaires.

Tableau 2.3 Exemples de différence entre le taux de classification obtenu sur chaque récif selon que les descripteurs aient été sélectionnés sur eux-mêmes ou sur Martin 2012-2013

Récif	Écart du taux de classification [%]
Martin 2010-2011	0,2
No Name 2012-2013	1,1
No Name 2010-2011	-0,2

Le tableau 2.3 illustre la différence entre les taux de classification obtenus sur un récif à partir de son propre ensemble de descripteurs et à partir de celui extrait pour Martin 2012-2013. Dans la moitié des cas, le test de Student ($\alpha = 0,05$) indique que les résultats sont significa-

tivement différents, en faveur des ensembles formés sur le récif lui-même. Nous pouvons cependant conserver l'ensemble de Martin 2012-2013 puisque l'erreur associée à son utilisation est faible. En faisant ce choix, il ne faut pas perdre de vue qu'un biais est inévitable lors de la classification de récifs autres que Martin 2012-2013.

2.5.4 Apport de chaque descripteur à l'ensemble

Outre le nombre de valeurs sélectionnées associées à chaque technique, il est aussi intéressant de savoir quel est le poids de l'information qu'elles apportent. La figure 2.8 montre le gain d'information moyen fourni par chaque modalité d'acquisition. Puisque CLBP est la technique qui détient le plus grand nombre de valeurs sélectionnées et le plus haut pointage moyen au niveau du grain d'information, c'est la technique considérée la plus importante pour la classification. La troisième image de la pyramide gaussienne est également prometteuse en ce sens.

La figure 2.9 illustre donc la performance de sous-ensembles dérivés du groupe déjà réduit de descripteurs. Nous formons les sous-ensembles en se basant sur la fréquence d'apparition et sur le gain d'information lié à chaque modalité. Ainsi, ne conserver que les modalités qui offrent un haut gain d'information a un impact négatif sur les résultats. La performance des sous-groupes est toujours inférieure d'au moins 1,3 % au taux de classification des descripteurs sélectionnés par les algorithmes. Il n'est conséquemment pas souhaitable de réduire la taille de ce groupe.

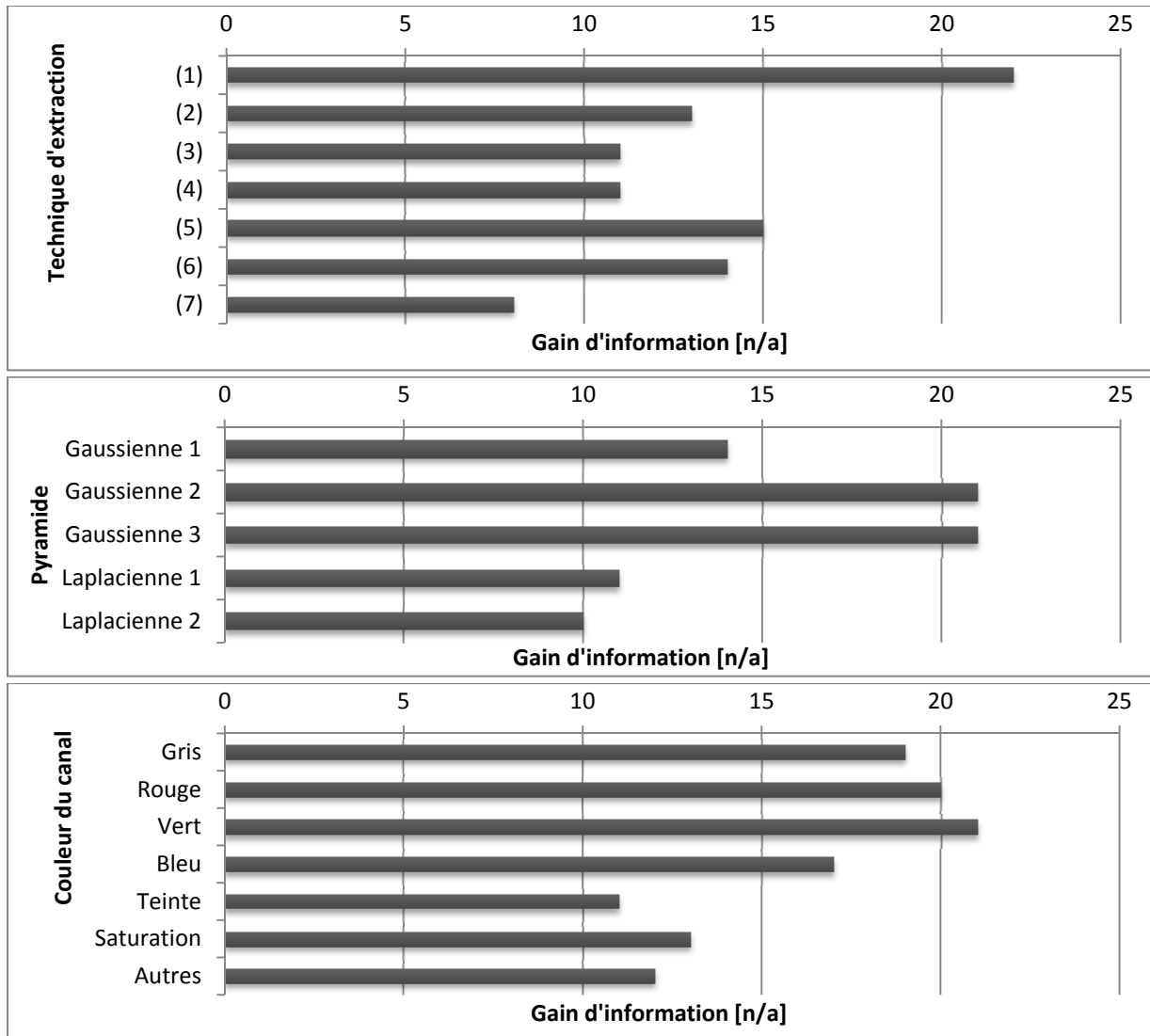


Figure 2.8 Gain moyen d'information associé à chaque modalité d'extraction d'attributs.
 (Légende : (1) CLBP, (2) réponse aux filtres de Gabor, (3) statistiques sur l'histogramme, (4) matrice de cooccurrence, (5) intégration de la transformée de Fourier, (6) histogramme de teintes, (7) histogramme d'angle opposé

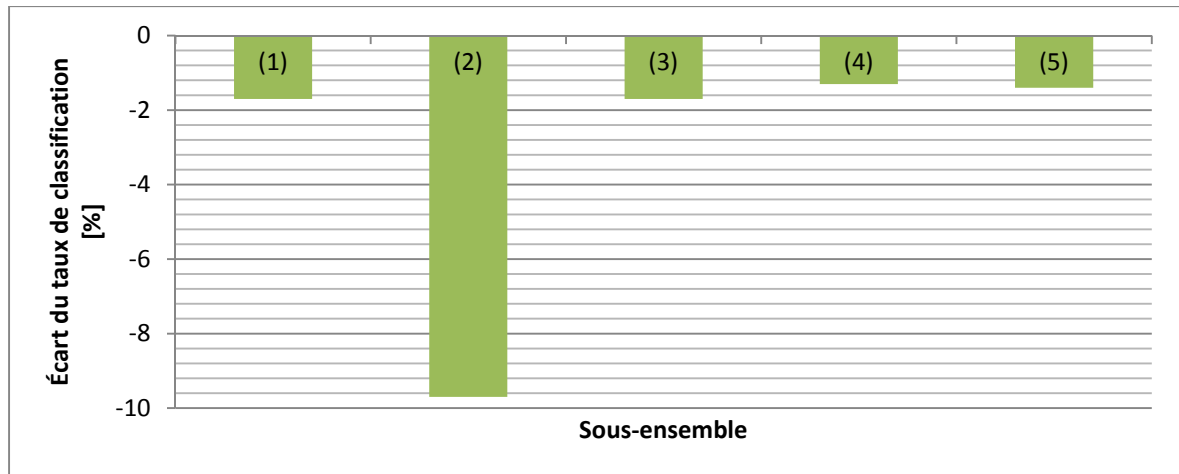


Figure 2.9 Différence de performance entre l'application des descripteurs sélectionnés et certains de leurs sous-ensembles. Légende : (1) CLBP, (2) Réponse aux filtres de Gabor, (3) CLBP + Réponse aux filtres de Gabor, (4) canaux gris, rouge, vert et bleu, (5) couches 2 et 3 de la pyramide gaussienne

2.5.5 Comparaison de la performance des descripteurs à la fine pointe de la technique

À des fins de comparaison, nous mesurons la performance du groupe de descripteurs sélectionnés à celle d'une ACP appliquée sur les 18 760 valeurs du vecteur initial et à la méthode de (Shihavuddin et al., 2013). Le tableau 2.4 présente les résultats de ces expérimentations. Il montre que la classification à l'aide des descripteurs sélectionnés permet de conserver un nombre de descripteurs significativement moindre que la plupart des autres méthodes, ce qui rend son temps d'exécution exponentiellement plus court.

L'analyse à l'aide du test de Student ($\alpha = 0,05$) permet également de déduire que l'ensemble de descripteurs sélectionnés est significativement plus performant que les autres. De plus, il surpasse la fine pointe de la technologie de près de 3 %. L'extraction d'informations supplémentaires, fournies par les pyramides et à travers les canaux de couleurs, est cependant un travail fastidieux et couteux en temps. Dans l'optique où le temps de traitement des images serait une contrainte, il serait préférable d'utiliser les descripteurs de (Shihavuddin et al., 2013). Ce n'est cependant pas le cas pour la classification de la base de données de l'AIMS. L'ensemble de descripteurs réduit peut donc être conservé.

Tableau 2.4 Taux de classification du système en fonction du nombre de valeurs utilisées et de leur méthode de sélection

Technique	Longueur du vecteur	Performance
Descripteurs sélectionnés	264	84 %
ACP	100	74 %
ACP	300	74 %
ACP	1000	63 %
ACP	2936	83 %
Descripteurs de (Shihavuddin et al., 2013)	250	81 %

2.6 Optimisation des paramètres des SVM

2.6.1 Utilisation d'un SVM

Le classificateur choisi est SVM parce que cette méthode est adaptée à une grande base de données, telle que celle de l'AIMS. Elle est aussi malléable et robuste face au débalancement des classes. Pour sa part, un noyau RBF est facile à optimiser et son efficacité a été prouvée dans le contexte de classes débalancées.

Une recherche en grilles est la technique qui sélectionne généralement les paramètres optimaux du noyau RBF d'un SVM. Ainsi, nous faisons des itérations sur γ et C de manière à explorer toutes leurs combinaisons sur la plage désignée et un SVM est créé à partir de chacune d'elles. Nous les testons ensuite à partir d'un couple entraînement-test constant et les taux de classification obtenus témoignent de leur performance. Puisque ce procédé est coûteux en temps, nous effectuons initialement une recherche grossière. Nous entreprenons ensuite des itérations plus fines autour du résultat grossier optimal et nous sélectionnons une combinaison finale. La performance du SVM dépend fortement de ces optimisations (Novakovic et Veljovic, 2011).

Ces valeurs sont, de plus, caractéristiques du groupe d’entraînement utilisé pour les produire. Puisque les ensembles de coraux contiennent des espèces similaires d’un récif à l’autre, il est possible que les paramètres optimisés y soient équivalents et réutilisables. La délimitation de plages de valeurs réduites pour l’optimisation des paramètres du RBF diminuerait significativement le temps d’entraînement du classificateur.

2.6.2 Méthodologie d’optimisation pour la base de données de l’AIMS

Nous faisons la classification de coraux à partir d’un SVM basé sur C-SVC et sur un noyau RBF, dont l’algorithme est compris dans la bibliothèque LibSVM (Chang et Lin, 2011). L’étape d’optimisation de γ et de C débute par une recherche en grille. Un premier balayage grossier sert à déterminer la combinaison (γ_0, C_0) la plus performante sur la grande plage de valeurs présentée au tableau 2.5 (Jiang, Missoum et Chen, 2014). La recherche fine ratisse ensuite une nouvelle plage de valeurs centrée sur (γ_0, C_0) , comme le montre le tableau 2.6. Le pas des itérations, lors de cette étape, est plus faible. Nous enregistrons alors les résultats des optimisations pour plusieurs ensembles d’entraînement de manière à juger s’ils sont compris dans une plage de valeurs réduite.

Tableau 2.5 Plage des valeurs des exposants pour $\gamma = 2^x$ et $C = 2^y$ lors de la recherche en grille grossière

Paramètre	Valeur minimale	Valeur maximale	Pas de l’itération
x	-5	15	2
y	-15	15	2

Par la suite, nous évaluons à l’aide de deux méthodes si les paramètres peuvent être réutilisés à travers les ensembles d’entraînements. Ces deux techniques visent à mettre en commun des paramètres optimisés pour des cas qui sont intuitivement similaires. Ainsi, la première démarche exploite la similitude des récifs à travers les années, alors que la seconde évalue des couples de récifs distincts à l’intérieur d’une même année.

Tableau 2.6 Plage de valeurs des exposants pour $\gamma = 2^x$ et $C = 2^y$ lors de la recherche en grille fine centrée sur $\gamma_0 = 2^{x_0}$ et $C_0 = 2^{y_0}$

Paramètre	Valeur minimale	Valeur maximale	Pas de l'itération
x	$x_0 - 2$	$x_0 + 2$	0,4
y	$y_0 - 2$	$y_0 + 2$	0,4

Le premier cas utilise donc un récif individuel, dont la composition, pour une certaine année, forme l'ensemble d'entraînement du SVM et, pour une autre année, compose le groupe test. Nous optimisons quatre couples de valeurs γ et C correspondant aux quatre périodes présentes dans la base de données. Nous formons ensuite un SVM pour chaque couple de paramètres et chaque SVM classe finalement le groupe test. Nous répétons l'ensemble des étapes pour trois récifs différents. Cette démarche permet ainsi d'évaluer si les paramètres optimisés du RBF peuvent être conservés à travers les années.

La seconde méthode utilise un couple de deux récifs pour l'entraînement et un troisième récif pour constituer le groupe test. Nous optimisons tout d'abord les paramètres du RBF sur la combinaison des deux récifs d'entraînement, puis sur chacun d'eux individuellement, pour un total de trois optimisations. Nous construisons ensuite des SVM à partir de chacun des couples de paramètres optimisés. Les trois SVM classifient finalement le groupe test. Cette démarche permet de savoir si l'optimisation sur les deux récifs est équivalente à l'optimisation sur un seul d'entre eux. Autrement dit, elle permet de juger si l'ajout de nouvelles données à l'ensemble d'entraînement fait varier ses caractéristiques de façon trop importante pour que l'optimisation initiale des paramètres soit toujours valable.

2.6.3 Résultats de l'optimisation des paramètres sur plusieurs récifs

Nous avons recueilli les valeurs optimisées des paramètres du noyau RBF, γ et C , tant à l'étape de la recherche grossière que de la recherche fine pour 59 optimisations sur 13 ensembles d'entraînement distincts. Pour la recherche grossière, γ ne prend que les valeurs

$2^1, 2^3, 2^4$ et $C, 2^{-5}, 2^{-3}$. La recherche fine donne, pour sa part, des résultats divers autour de ces valeurs. Nous pouvons donc réduire la plage de recherche pour les valeurs grossières.

Lors de la conservation des paramètres à travers les périodes d'échantillonnage, les optimisations faites sur le récif et l'année qui constitue le groupe d'entraînement servent de référence. Ainsi, nous comparons leurs taux de classification à ceux des scénarios où l'optimisation est effectuée sur le même récif, mais pour d'autres années (*Voir* ANNEXE IV, p.107). La figure 2.10 montre la variation des résultats dans ces cas. Il est à noter qu'une prise de valeur atypique des paramètres d'optimisation cause les écarts de 10 %, 14 % et 16 %. Ils sont donc considérés aberrants. Globalement, il n'y a donc pas de variation significative du taux de classification pour les 18 optimisations, puisque l'ensemble des valeurs se concentre autour de 0 % d'écart.

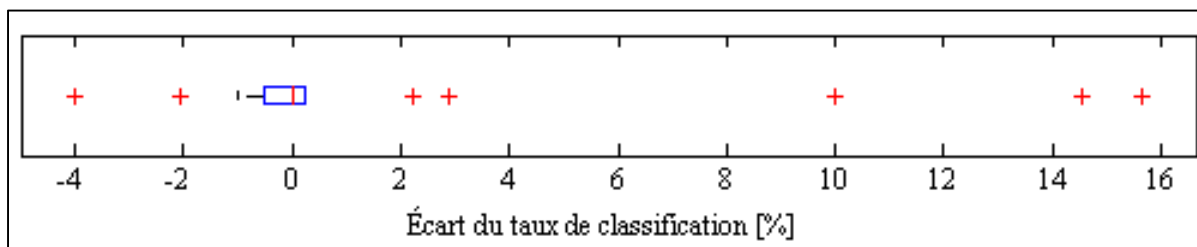


Figure 2.10 Écarts entre les taux de classifications lors de l'optimisation des paramètres γ et C sur le groupe d'entraînement et lorsque nous conservons ces paramètres à travers les groupes d'entraînement

Les résultats sont similaires dans le cas de la conservation des paramètres optimisés lors de l'ajout de données supplémentaires à l'ensemble d'entraînement (*Voir* ANNEXE IV, p.107), comme le montre la figure 2.11. Une amélioration des taux de classification est d'ailleurs plus fréquente qu'une dépréciation pour les six combinaisons testées. N'utiliser qu'un seul des récifs d'entraînement plutôt que deux pour faire une optimisation des paramètres du RBF ne nuit donc pas à la classification.

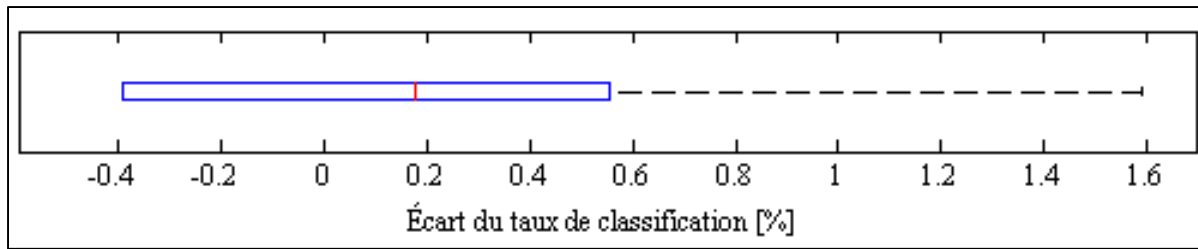


Figure 2.11 Écarts entre les taux de classifications des cas d’optimisation des paramètres sur les deux récifs d’entraînement et sur un seul des récifs d’entraînement

2.6.4 Interprétation des valeurs prises par les paramètres

À la suite des observations, pour ce qui est de la recherche grossière, la zone de recherche de paramètres optimaux peut être réduite aux valeurs recensées pour γ et C . Pour sa part, l’optimisation des paramètres du SVM sur diverses années pour de mêmes récifs n’introduit pas une variation statistiquement significative. De même, les résultats ne montrent pas de variation significative des taux de classification selon que les paramètres aient été optimisés sur les deux récifs d’entraînement ou seulement sur un seul. Selon les cas, il peut y avoir détérioration ou amélioration de la performance du système.

Ce phénomène est explicable par le fait qu’un groupe d’entraînement construit le SVM alors qu’un groupe distinct le teste. La variabilité entre ces groupes introduit donc des variations plus grandes au niveau du SVM que les paramètres optimisés en tant que tels. Nous n’avons donc pas à généraliser l’optimisation des paramètres du SVM à chaque expérimentation. Des paramètres précis peuvent être choisis et conservés pour toutes les classifications de groupes benthiques à l’intérieur de la base de données de l’AIMS. En l’occurrence, $\gamma = 2^{-1.4}$ et $C = 2^{+1.8}$ correspondent à une optimisation sur 75 % de la base de données de l’AIMS et seraient un choix judicieux.

2.6.5 Discussion sur le processus d’optimisation

La bibliothèque LibSVM (Chang et Lin, 2011) utilise le critère d’exactitude pour sélectionner ses paramètres. Ainsi, l’algorithme valorise ceux qui permettent d’obtenir le plus grand

nombre d'échantillons bien classifiés. Cependant, étant donné le débalancement des classes de la base de données de l'AIMS, les classes faiblement peuplées sont négligées par ce critère. Dans l'éventualité où on voudrait favoriser une classe particulière, le critère *F – measure* serait à privilégier. Autrement, pour que la performance de classification soit uniforme à travers les classes, *G – mean* deviendrait une alternative avantageuse. Toutefois, la classification actuelle vise à donner un portrait global des groupes benthiques présents dans les récifs et non à examiner particulièrement les cas des classes faiblement peuplées. Conséquemment, utiliser l'exactitude comme critère d'optimisation est acceptable.

2.7 Configuration finale

À la suite de l'analyse de chaque étape du système de classification, une configuration optimale pour la base de données de l'AIMS est fixée. Ses paramètres permettront en définitive d'identifier le contenu des images de manière efficace. Le tout est présenté au tableau 2.7.

Tableau 2.7 Configuration finale du système de classification

Étape	Technique	Paramètres
Prétraitement	Aucun	
Segmentation	Imagettes carrées centrées sur le point étiqueté	350 pixels de côté
Extraction des descripteurs	Ensemble de descripteurs sélectionné par les algorithmes	
Réduction de la dimension de l'espace de représentation	Aucun	
Classification	SVM	«Un contre un» Noyau RBF : $\gamma = 2^{-1,4}$, $C = 2^{+1,8}$

CHAPITRE 3

GÉNÉRALISATION SUR LA BASE DE DONNÉES DE L'AIMS

3.1 Introduction

Nous devons classifier les images de la base de données de l'AIMS afin de fournir des statistiques sur les populations de groupes benthiques des récifs. À la suite de l'optimisation de l'algorithme pour les caractéristiques propres des images (CHAPITRE 2), il faut évaluer sa performance maximale sur toutes les images en faisant une généralisation de son application. De plus, nous évaluons la compatibilité de divers groupes d'entraînement et de test pour augmenter la probabilité d'obtenir un haut taux de classification.

3.2 Seuils de performance du système

Le système de classification développé a des limites liées à son utilisation avec la base de données de l'AIMS. Des combinaisons similaires de descripteurs et de SVM ont obtenu des taux de classification de plus de 99 % sur des bases de données de textures telles que CUR-TeT (Shihavuddin et al., 2013). Les limites ne sont donc pas liées aux distinctions de textures, mais plutôt aux proportions des populations de groupes benthiques dans chaque récif, aux faibles différences entre les classes et à la grande diversité de chacune d'elles.

Afin de déterminer l'ampleur des difficultés, nous évaluons des cas optimaux de performance. Ils correspondent à des classifications de données à partir d'un modèle bâti sur d'autres données appartenant à un même groupe uniforme. Nous effectuons le tout par l'application de validations croisées à dix ensembles. La performance de la validation croisée peut, par la suite, servir d'étalon pour les autres simulations : lorsqu'un couple de groupes entraînement-test est performant, son taux de classification tend vers celui de la validation croisée du groupe test. Par ailleurs, il existe aussi une incertitude quant aux résultats des validations croisées. En effet, ils sont dépendants du choix des groupes d'entraînement. Leur variation reflète donc la robustesse du système en ce qui regarde les données.

3.2.1 Méthodologie d'évaluation des limites du processus de classification

La capacité du système à classifier la base de données de l'AIMS est initialement évaluée par des validations croisées pour chaque récif et pour chaque année individuellement. Ainsi, nous évaluons les récifs Martin, No Name, Carter, Linnet, Lizard, Macgillivray, North direction et Yonge de façon distincte pour les périodes 2012-2013, 2010-2011, 2008-2009 et 2006-2007. Ce faisant, nous analysons l'ensemble des données disponibles. Pour mesurer la robustesse du système, nous répétons ensuite ces validations croisées cinq fois chacune et nous entreprenons des validations croisées pour chaque période d'échantillonnage, pour quelques récifs. Subséquemment, nous appliquons la validation croisée aux huit récifs d'une période, puis aux quatre périodes des récifs Martin, Yonge et Carter.

3.2.2 Résultats des validations croisées

Les résultats de toutes validations sont consignés à l'ANNEXE V. Les validations croisées sont d'abord menées sur chaque récif pris individuellement au cours d'une seule période d'échantillonnage. La figure 3.1 montre la matrice de confusion moyenne obtenue sur l'ensemble des données. Elle reflète donc la capacité du système à classifier la base de données de l'AIMS. Nous remarquons qu'elle est dominée par la haute confusion de toutes les classes avec le groupe benthique des algues.

La figure 3.2A expose la distribution des taux de classification à travers les périodes d'échantillonnage. La figure 3.2B illustre, quant à elle, leur variation lorsqu'on répète chacune des validations croisées cinq fois. Ainsi, les résultats sont équivalents à travers le temps, mais la période 2006-2007 est beaucoup plus variable que les autres lors des répétitions. Elle est donc moins robuste et moins fiable.

		Classe prédite							N _M
% CR		AL	CM	CD	AB	E	M	AU	
Classe réelle (CR)	AL	88	0	10	1	0	0	0	39,7
	CM	18	52	29	0	0	0	0	5,0
	CD	21	1	74	3	0	0	0	25,2
	AB	23	0	3	74	0	0	0	4,5
	E	73	2	23	0	1	0	0	0,7
	M	34	7	29	0	0	30	0	0,5
	AU	70	1	23	0	0	0	6	1,5

Figure 3.1 Matrice de confusion moyenne pour les validations croisées menées sur chaque récif et chaque période
(Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres, N_M – nombre d'échantillons en milliers dans la classe réelle)

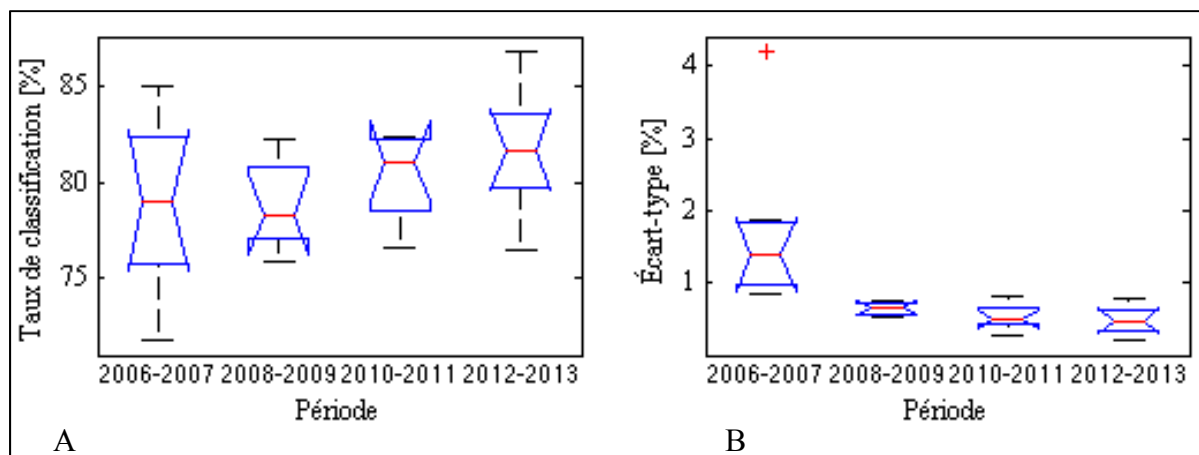


Figure 3.2 (A) Distribution des taux de classification à travers les années pour les validations croisées sur chaque couple de récif et de période et (B) distribution des écart-types correspondants liés à la robustesse du système

La figure 3.3 montre que, tout comme les périodes d'échantillonnage, tous les récifs ne s'équivalent pas. Ainsi, Yonge et Macgillivray ont des performances fortes, alors que les taux de classification de Carter, Lizard et No Name sont sous la barre des 80 %. La figure 3.4 présente la matrice de confusion obtenue pour Linnet 2012-2013, un exemple de haute performance du système. De plus, dans la figure 3.5, nous avons éliminé les périodes et les récifs difficiles à classifier de la matrice de confusion moyenne. Elle ne tient donc plus compte de la période 2006-2007, ainsi que des récifs Carter, Lizard et No Name. Son analyse permet d'extraire les statistiques du tableau 3.1 qui reflètent la performance maximale du système.

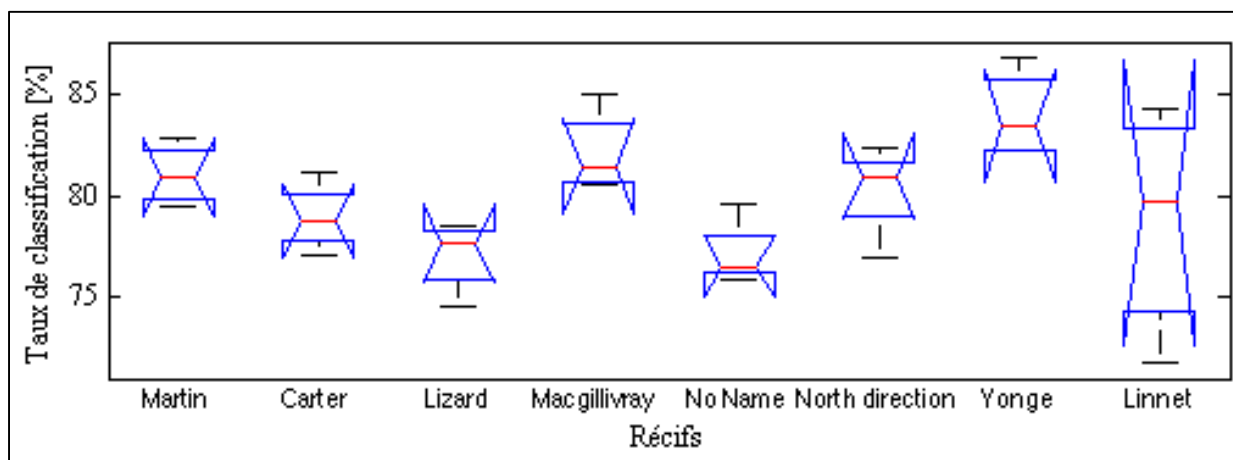


Figure 3.3 Taux de classification selon les récifs pour les validations croisées sur chaque couple de récif et de période

Nous regroupons ensuite les données par période, puis par récif pour faire les validations croisées. Dans les figures 3.6 et 3.7 respectivement, nous comparons les résultats des validations croisées à une moyenne des validations croisées correspondant à leurs groupes de données. L'intégration en ensembles atteint de meilleures performances que les moyennes de validations croisées individuelles.

		Classe prédite							N
% CR		AL	CM	CD	AB	E	M	AU	
Classe réelle (CR)	AL	90	0	9	0	0	0	0	1579
	CM	17	52	32	0	0	0	0	163
	CD	15	1	84	0	0	0	0	1172
	AB	34	0	23	43	0	0	0	35
	E	50	0	25	0	25	0	0	4
	M	22	11	56	0	0	11	0	9
	AU	45	0	42	0	0	0	13	38

Figure 3.4 Matrice de confusion obtenue pour la validation croisée touchant le récif Linnet pour la période 2012-2013
(Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres, N – nombre d'échantillons dans la classe réelle)

Tableau 3.1 Taux de rappel, précision et population des groupes benthiques, moyennés sur toutes les validations croisées pour chaque couple de récif et de période, à l'exception de ceux éliminés (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe benthique	AL	CM	CD	AB	E	M	AU	Tous
Population [%]	51,0	4,2	33,7	8,4	0,8	0,4	1,5	100
Taux de rappel [%]	86,3	38,6	77,1	75,7	0,6	26,4	9,8	81,8
Précision [%]	81,9	86,9	74,0	72,5	65,8	79,0	96,7	79,0

		Classe prédite							N _M
% CR		AL	CM	CD	AB	E	M	AU	
Classe réelle (CR)	AL	86	0	12	2	0	0	0	22,9
	CM	20	39	40	1	0	0	0	1,9
	CD	18	0	77	5	0	0	0	15,2
	AB	21	0	3	76	0	0	0	3,8
	E	73	0	25	1	1	0	1	0,4
	M	28	6	40	0	0	26	0	0,2
	AU	61	1	28	0	0	0	10	0,7

Figure 3.5 Matrice de confusion moyenne pour les validations croisées menées sur chaque récif et chaque période, à l'exception des récifs éliminés (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres, N_M – nombre d'échantillons en milliers dans la classe réelle)

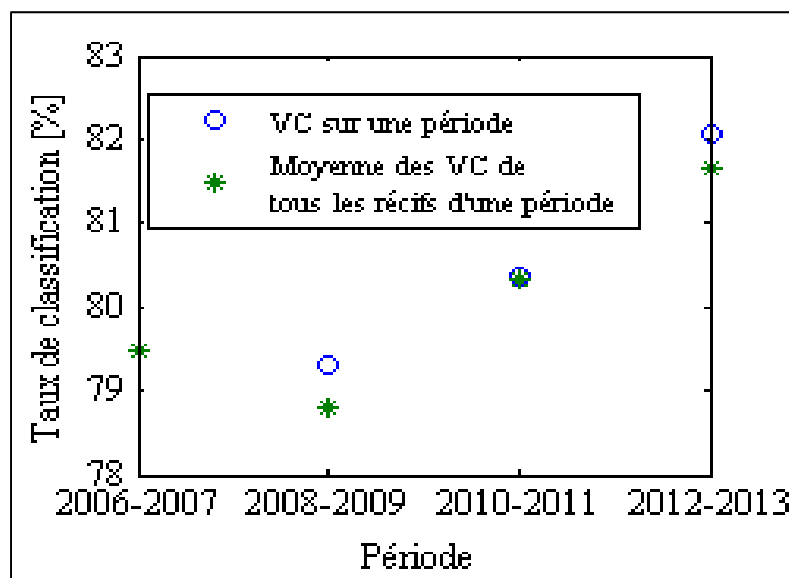


Figure 3.6 Comparaison des taux de classification de validations croisées (VC) effectuées sur tous les récifs d'une période aux moyennes des résultats des validations croisées des récifs de cette période pris isolément

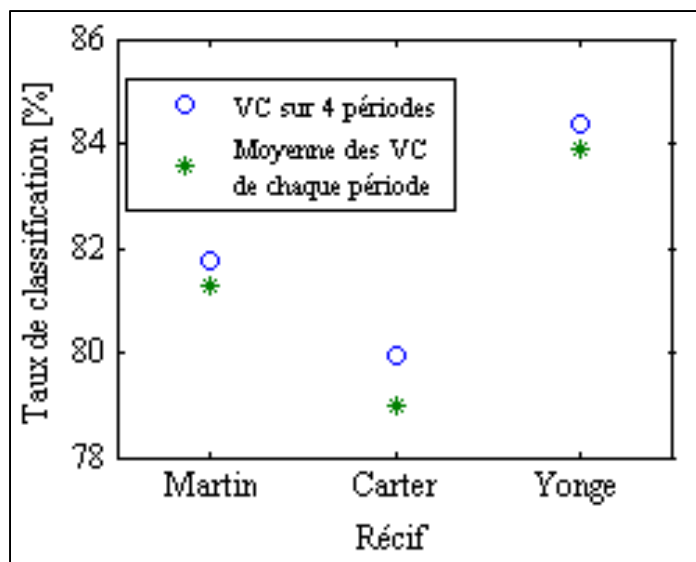


Figure 3.7 Comparaison des taux de classification de validations croisées (VC) effectuées sur toutes les périodes d'un récif aux moyennes des résultats des validations croisées de chaque période de ce récif prise isolément

3.2.2.1 Interprétation des limites de performance du système de classification

La performance des validations croisées est variable selon les récifs, les années et les groupes benthiques. La précision touchant les classes est supérieure à 80 % : le taux d'erreur pour chaque classe est faible relativement à la population. Globalement, le taux de rappel atteint 81,8 % d'échantillons convenablement classifiés. Selon les récifs, il s'échelonne entre 76,5 % et 83,8 % comme le montre la figure 3.3. Certains récifs contiennent donc des images plus compatibles avec le système que d'autres.

Les taux de classification varient peu à travers les années, ce qui signifie que la qualité des images de la base de données est constante. Toutefois, en comparant les écart-types de 2012-2013 et de 2006-2007 (*Voir figure 3.2B*) à leurs populations correspondantes (*Voir figure 1.6*), nous déduisons que la robustesse de l'algorithme diminue lorsque le nombre d'échantillons d'un récif d'entraînement est réduit. Conséquemment, avoir un grand groupe d'entraînement très diversifié fait augmenter la fiabilité du SVM.

Dans le cas de l'intégration des données à travers les périodes d'échantillonnage ou les récifs, nous améliorons en règle générale les résultats. Ainsi, les récifs supplémentaires du groupe d'entraînement ajoutent de l'information qui n'est pas disponible dans le cas d'un récif unique. Ce n'est que dans de rares occasions qu'ils induisent une plus grande erreur.

Les validations croisées révèlent également que certains groupes benthiques sont significativement plus faciles à différencier que d'autres. Par exemple, les classes « autres » et « éponges » ont un taux de rappel très faible et la quasi-totalité de leur confusion est avec les algues. Du point de vue du système, il y a donc une forte ressemblance des groupes benthiques « algues », « éponges » et « autres ». D'ailleurs, cette performance ne peut être imputée uniquement au débalancement des classes, puisque les millépores, dont la population est plus petite que celle des « éponges » ou des « autres », a un taux de rappel de plus de 15 % supérieur au leur. La forte confusion de toutes les classes avec les algues peut cependant être expliquée par la présence résiduelle d'algues dans la plupart des images à la suite de la segmentation.

Un cas similaire est celui des coraux mous. Ils ont une quantité d'échantillons équivalente à celle des abiotiques, alors qu'ils ont un taux de rappel inférieur de 37,1 %. En effet, près de 40,3 % de tous les coraux mous sont classés dans les coraux durs. Les classes « coraux mous » et « coraux durs » peuvent donc être fusionnées de manière à réduire les conséquences de ces similitudes sur le système. Sa performance subséquente est présentée au tableau 3.2 et à la figure 3.8, où le nouveau taux de rappel atteint alors 83,7 %. La distinction des sous-classes pourrait être entreprise par la suite, à un niveau plus fin. Cette étude permettrait de mettre en lumière les caractéristiques précises les distinguant.

Tableau 3.2 Taux de classification, précision et population des groupes benthiques fusionnés

Groupe benthique	Algues	Coraux	Abiotiques	Autres	Tous
Population [%]	51,0	37,9	8,4	2,7	100
Taux de rappel [%]	86,3	77,7	75,7	9,8	83,7
Précision [%]	81,9	79,6	72,5	92,3	81,3

		Classe prédite				N _M
% CR		AL	CO	AB	AU	
Classe réelle (CR)	AL	86	12	2	0	39,7
	CO	18	78	4	0	30,1
	AB	21	3	76	0	4,5
	AU	60	30	0	10	2,7

Figure 3.8 Matrice de confusion moyenne pour les groupes benthiques fusionnés (Légende : AL – algues, CO – coraux, AB – abiotiques, AU – autres, N_M – nombre d'échantillons en milliers dans la classe réelle)

3.3 Scénarios de classification des images de l'AIMS

Dans un contexte réel, il est nécessaire d'utiliser des connaissances préalables pour classer de nouveaux échantillons. Pour ce faire, nous pouvons adopter diverses stratégies. La classification peut tout d'abord reposer sur les images des récifs prises aux années précédentes ou postérieures. Le groupe d'entraînement ne contient alors que les représentations, pour un récif précis, de périodes d'échantillonnage autres que celle qui est testée. Le modèle est donc spécifique au récif. Le groupe d'entraînement peut également contenir l'ensemble des récifs au cours des années précédentes ou suivantes, pour avoir une représentation plus globale des groupes benthiques. Autrement, nous pouvons étiqueter certains récifs de la période

d'échantillonnage testée, puis généraliser ces identifications à l'ensemble des clichés. L'ensemble de ces scénarios doit être évalué afin de déterminer lequel est optimal pour la classification de coraux.

De plus, il est possible que certains couples entraînement-test soient plus performants que d'autres à cause de la ressemblance de leurs échantillons. Pour cette raison, nous calculons le coefficient de corrélation de Pearson qui lie leurs compositions à un niveau taxonomique fin, soit la famille, et nous nommons ce coefficient C_{comp} . Cela permet de déterminer si leurs groupes benthiques contiennent des familles similaires et si cette affinité est reflétée lors de leur comparaison par le système. Ensuite, l'évaluation du coefficient de corrélation de Pearson qui lie C_{comp} et la performance de la classification indique si la ressemblance des récifs est un facteur déterminant pour garantir une bonne classification. La figure 1.7 donne une comparaison de haut niveau de la composition des récifs.

3.3.1 Méthodologie d'évaluation des couples entraînement-test

La première étape consiste à déterminer quelles combinaisons de récifs peuvent entraîner un SVM qui classifie un tiers récif de manière performante. Nous formons donc les groupes d'entraînement et de test au cours d'une unique période d'échantillonnage. Les SVM sont entraînés à partir d'un, de deux, de trois ou de sept récifs. Nous reprenons le même type de démarche en fixant le récif, mais en faisant varier les années. Nous constituons ainsi le groupe d'entraînement d'une, de deux, puis de trois périodes d'un même récif et nous en classifions la quatrième. Finalement, nous mettons en place une combinaison des classifications par de multiples récifs et par de multiples périodes d'échantillonnage. Dans ce cas, l'ensemble des récifs d'une ou de plusieurs années différentes de l'année à classifier composent le groupe d'entraînement du SVM. Pour toutes ces simulations, nous comparons enfin les résultats entre eux, ainsi qu'aux validations croisées pertinentes. Tous les résultats sont consignés à l'ANNEXE VI.

3.3.2 Résultats d'entraînements et tests sur deux récifs distincts, au cours d'une période d'échantillonnage

Nous avons formé différents couples de récifs d'entraînement et de test au hasard à l'intérieur de périodes d'échantillonnage prises isolément. L'analyse d'une dizaine de combinaisons forme la figure 3.9, où nous déduisons les taux de rappel des récifs de test de ceux de leurs validations croisées qui sont les étalons. Ces écarts s'échelonnent entre -2,0 % et -26,2 %. La figure 3.10 montre que le taux de rappel de chaque groupe benthique est également modifié.

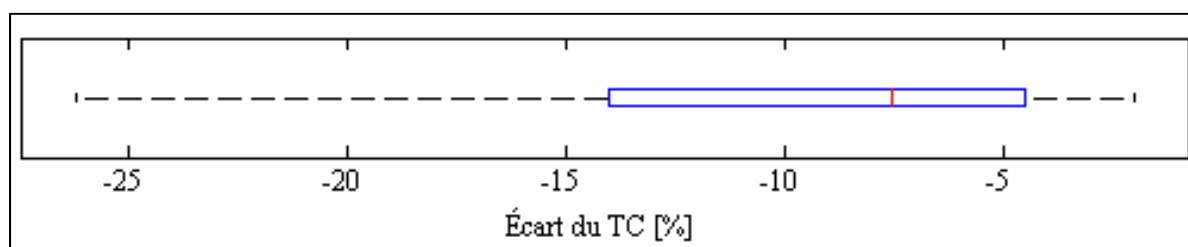


Figure 3.9 Écart du taux de classification par rapport à l'étalon lors de l'entraînement sur un seul autre récif de la même année pour dix cas

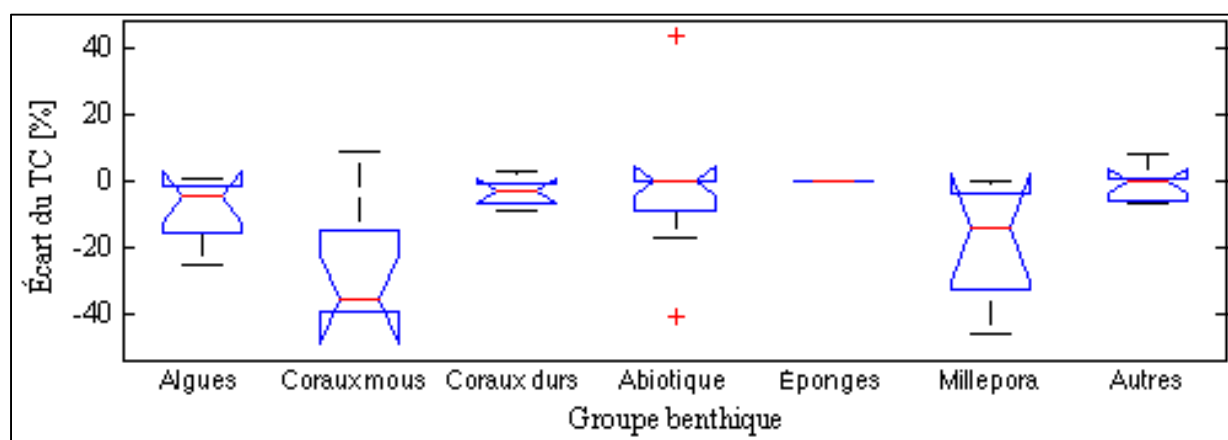


Figure 3.10 Écart entre le taux de classification de chaque classe benthique et l'étalon lors d'un entraînement sur un seul autre récif

Nous inversons, par la suite, certains des couples entraînement-test de manière à ce que le récif constituant le groupe d'entraînement devienne le récif de test et vice-versa. Les résultats de ces expériences sont présentés au tableau 3.3, dans lequel les valeurs expriment la varia-

tion de l'écart à l'étalon avant et après l'inversion. Ainsi, pour le cas particulier du couple Martin 2012-2013 et No Name 2012-2013, la différence est grande. Ce couple ne peut donc pas être interverti sans conséquence. La figure 3.11 présente, pour sa part, une comparaison entre la performance de la classification et C_{comp} , la corrélation des familles dans les récifs d'entraînement et de test.

Tableau 3.3 Comparaison des écarts, avec l'étalon, des taux de classification pour des couples entraînement-test intervertis

Couples de récifs	Différence absolue des écarts [%]
Carter 2010-2011 / No Name 2010-2011	4,0
Carter 2012-2013 / Yonge 2012-2013	1,7
Martin 2012-2013 / No Name 2012-2013	9,6

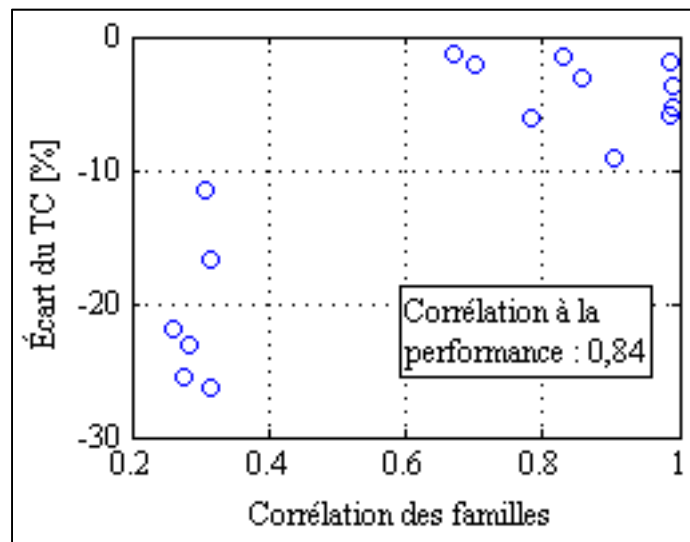


Figure 3.11 Corrélation entre C_{comp} et l'écart à l'étalon lors d'un entraînement sur un autre récif de la même année que le récif test

3.3.3 Résultats d'entraînements sur plusieurs récifs et de test sur un récif distinct, au cours d'une période

Nous ajoutons des récifs supplémentaires à chaque groupe d'entraînement. Sans être exhaustive du point de vue du nombre de combinaisons prises en compte, la figure 3.12 illustre les résultats. La courbe tend vers une réduction de l'écart à l'étalon au fur et à mesure des ajouts. Puisque huit récifs sont disponibles, l'entraînement sur sept récifs représente le maximum atteignable. La figure 3.13 donne l'ensemble des résultats de ce cas particulier.

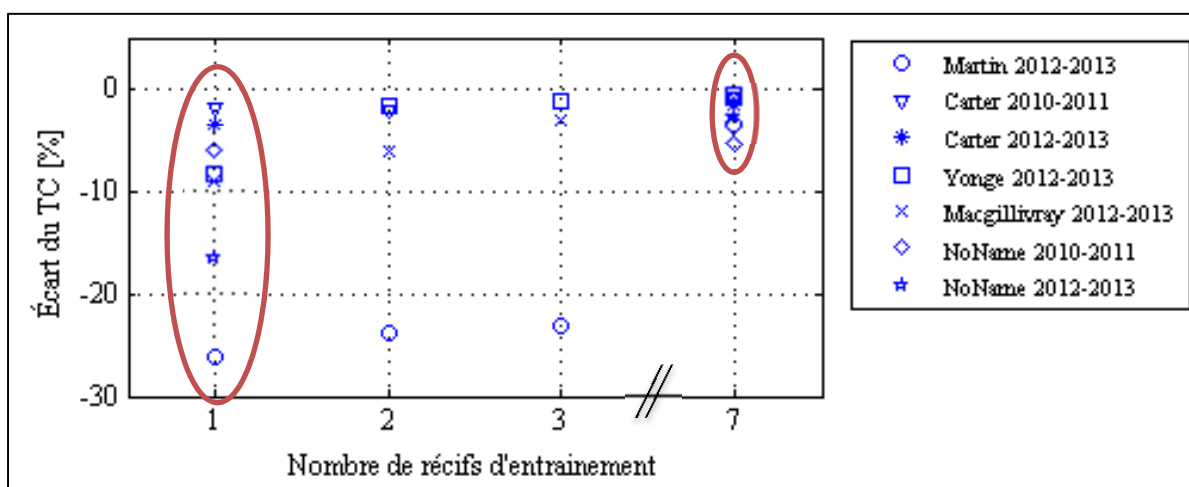


Figure 3.12 Écart entre l'étalon et le taux de classification d'un entraînement sur un, deux, trois ou sept récifs d'une même année que le récif test

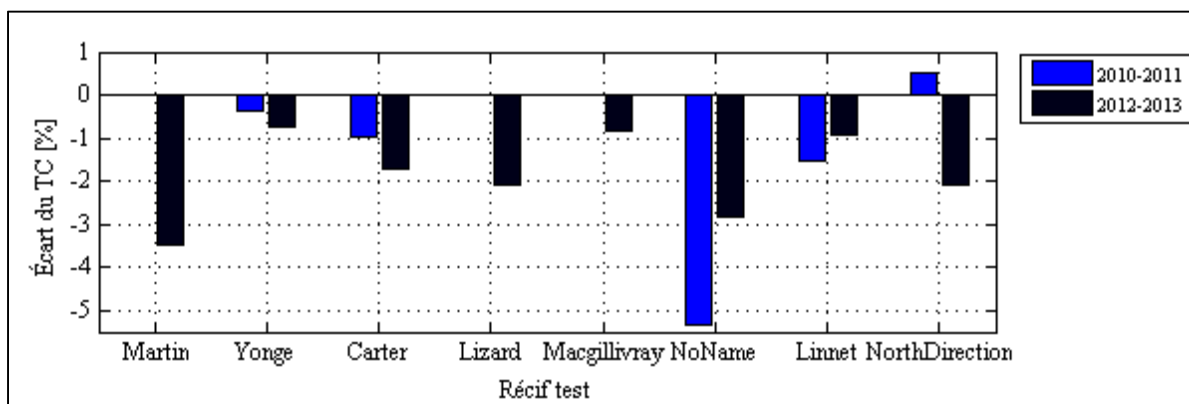


Figure 3.13 Écart entre le taux de classification et l'étalon lors de l'entraînement sur les sept autres récifs de la même année

3.3.4 Résultats d'entraînements sur plusieurs périodes d'échantillonnage pour un seul récif

Nous entraînons un SVM à partir de données appartenant au même récif que le groupe test, mais à une période différente. Les plages d'écarts par rapport à l'étalon sont alors larges, comme le montre la figure 3.14. Il y a également une variation du taux de rappel des groupes benthiques selon la figure 3.15 et, comme en témoigne la figure 3.16, un haut C_{comp} n'assure pas une haute performance. Donc, la ressemblance entre les images d'un récif à travers les années n'a pas d'impact tangible sur la classification. Dans le même ordre d'idées, la figure 3.17 illustre que l'éloignement temporel entre les périodes d'échantillonnage du récif d'entraînement et de test ne joue pas de rôle significatif sur la performance de la classification. Toutefois, l'ajout de périodes d'échantillonnage supplémentaires au groupe d'entraînement offre, en général, une amélioration marquée des résultats selon la figure 3.18.

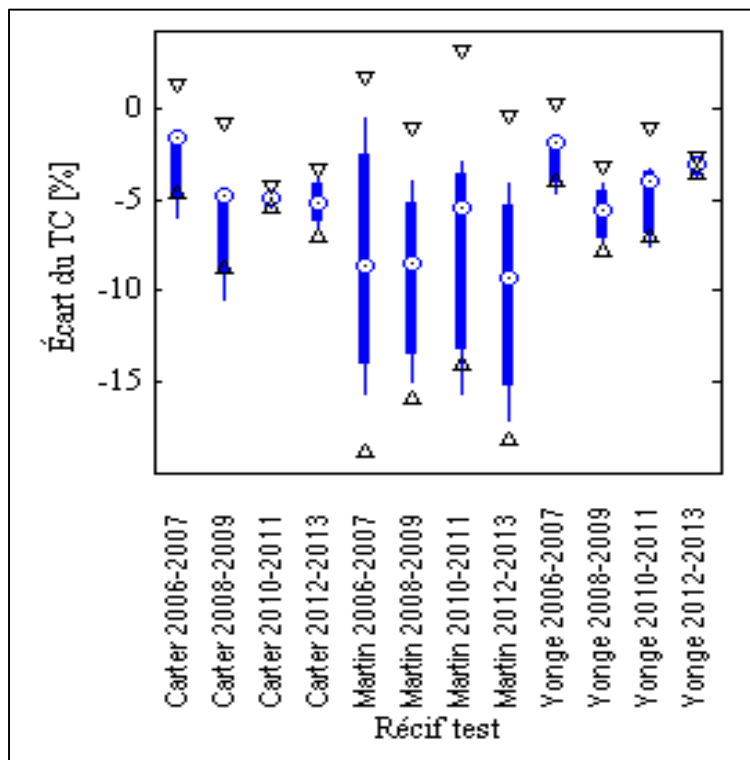


Figure 3.14 Écart entre le taux de classification et l'étalon lors d'un entraînement par une autre année du même récif

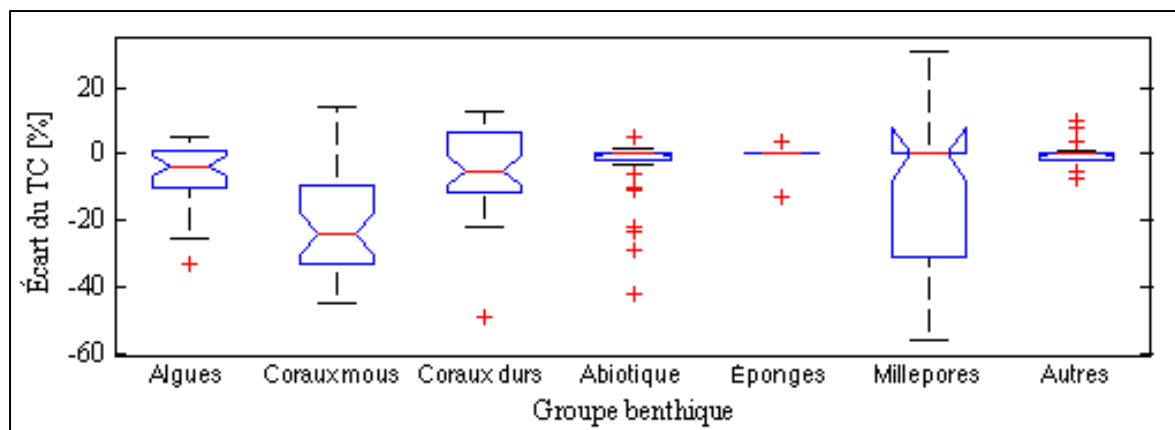


Figure 3.15 Variation par rapport à l'étalon du taux de rappel de chaque groupe benthique lors de l'utilisation d'une autre année d'un même récif comme groupe d'entraînement sur 36 cas

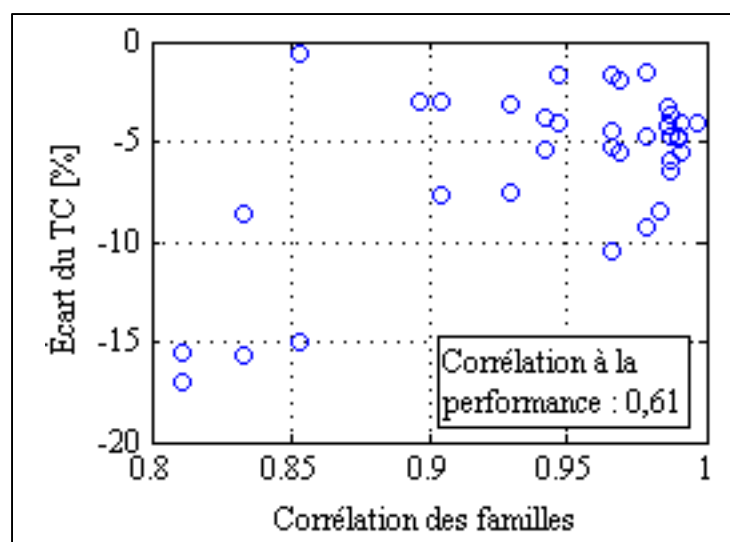


Figure 3.16 Corrélation entre C_{comp} et l'écart entre le taux de classification d'un récif test par un entraînement sur le même récif à une autre année

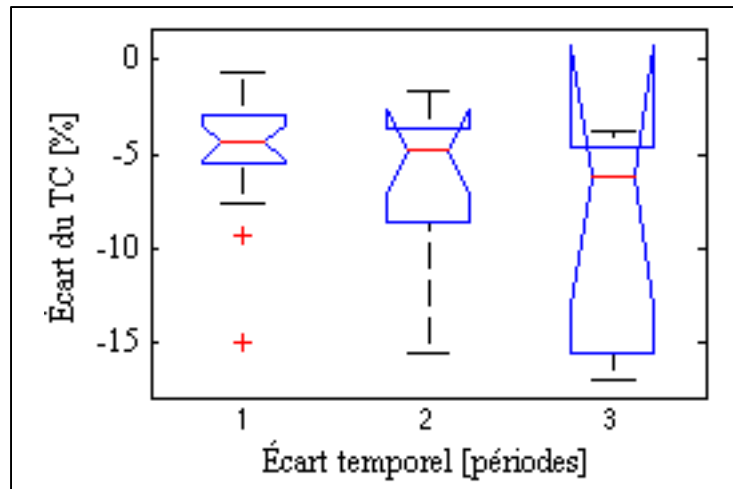


Figure 3.17 Écarts entre les taux de classification et l'étalon selon que l'entraînement ait été fait avec le même récif dont l'échantillonnage s'est fait sur une, deux ou trois périodes dans le temps sur 36 cas

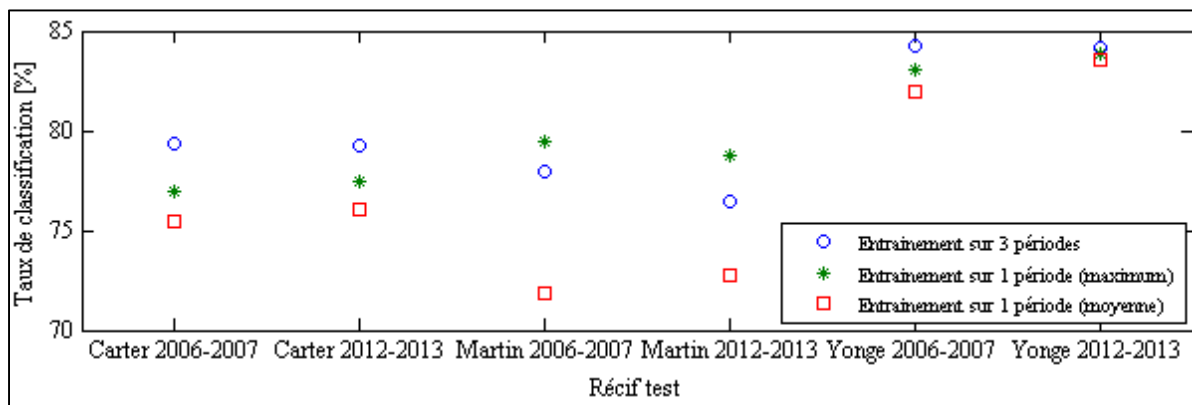


Figure 3.18 Évolution du taux de classification d'un récif selon le nombre de périodes utilisées lors de l'entraînement

3.3.5 Résultats d'entraînements sur d'autres années de tous les récifs

La dernière expérience entreprise est la mise en commun de tous les récifs de périodes d'échantillonnage pour créer les groupes d'entraînement et de test des SVM. De manière à conserver un scénario réaliste, nous utilisons seulement les récifs d'autres années que celle testée pour former les groupes d'entraînement. La figure 3.19 présente les écarts entre les résultats et les validations croisées effectuées sur l'ensemble des récifs de la période de test.

Le tableau 3.4 contient les taux de classification sur sept, puis quatre classes correspondant à l'utilisation de tous les périodes autres lors de l'entraînement. Les figures 3.20 et 3.21 illustrent quant à elles la matrice de confusion dans le meilleur cas pour ce type d'expérimentation, soit l'entraînement sur tous les récifs des années 2006-2007, 2008-2009 et 2010-2011. Le test est mené sur tous les récifs de l'année 2012-2013. Toutefois, puisque les récifs No Name, Carter et Lizard sont peu performants, nous les avons retiré de l'analyse, ce qui offre une amélioration de 1,2 % par rapport au tableau 3.4. Le tableau 3.5 en donne subséquemment les statistiques.

Tableau 3.4 Taux de classification (TC) de tous les récifs d'une période lors de l'entraînement sur tous les récifs de toutes les autres périodes

Période de test	TC sur sept classes [%]	TC sur quatre classes [%]
2006-2007	79,9	81,0
2008-2009	77,7	80,2
2010-2011	78,2	80,7
2012-2013	80,3	82,7

Tableau 3.5 Taux de classification, précision et population des groupes benthiques fusionnés pour un entraînement sur les périodes 2006-2007, 2008-2009 et 2010-2011 et un test sur 2012-2013, sans les récifs No Name, Lizard et Carter

Groupe benthique	Algues	Coraux	Abiotiques	Autres	Tous
Population [%]	54,5	35,8	6,9	2,8	100
Taux de rappel [%]	92,0	80,1	71,0	6,0	83,9
Précision [%]	82,5	85,5	93,2	100,0	90,1

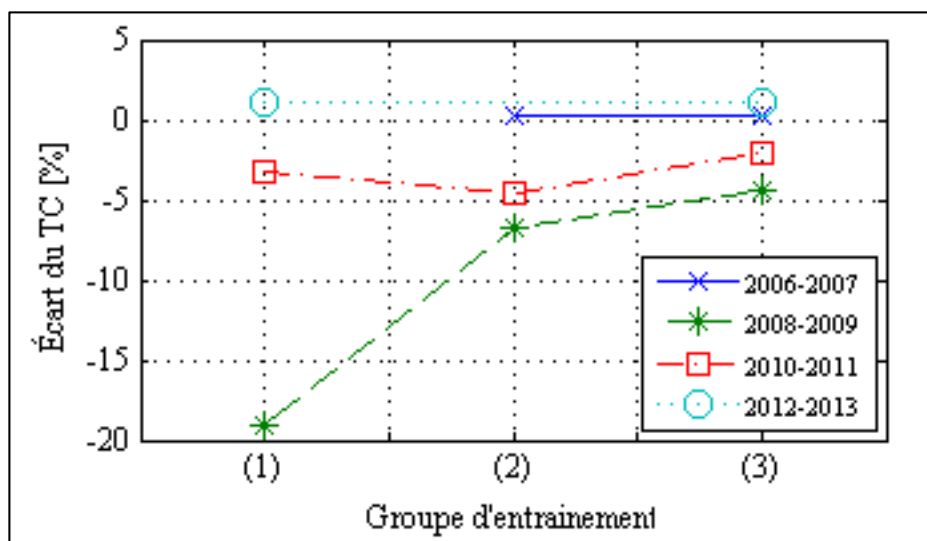


Figure 3.19 Écart à l'étalon des taux de classification de périodes de test lors d'entraînements sur tous les récifs (1) de toutes les périodes précédentes, (2) de toutes les périodes suivantes et (3) de toutes les périodes autres

		Classe prédite							
% CR		AL	CM	CD	AB	E	M	AU	N _M
Classe réelle (CR)	AL	92	1	6	1	0	0	0	8,2
	CM	18	57	24	0	0	0	0	0,7
	CD	20	4	76	0	0	0	0	4,7
	AB	24	0	5	71	0	0	0	1,0
	E	81	2	16	0	1	0	0	0,1
	M	46	8	34	0	0	12	0	0,1
	AU	71	1	22	0	0	0	7	0,2

Figure 3.20 Matrice de confusion résultant de la classification de la période 2012-2013 à partir d'un SVM entraîné sur les périodes 2006-2007, 2008-2009 et 2010-2011, sans les récifs No Name, Lizard et Carter (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres, N_M – nombre d'échantillons en milliers dans la classe réelle)

		Classe prédite				N _M
% CR		AL	CO	AB	AU	
Classe réelle (CR)	AL	92	7	1	0	8,2
	CO	20	80	0	0	5,4
	AB	24	5	71	0	1,0
	AU	70	24	0	6	0,4

Figure 3.21 Matrice de confusion consolidée résultante de la classification de la période 2012-2013 à partir d'un SVM entraîné sur les périodes 2006-2007, 2008-2009 et 2010-2011, sans les récifs No Name, Lizard et Carter
(Légende : AL – algues, CO – coraux, AB – abiotiques, AU – autres, N_M – nombre d'échantillons dans la classe réelle)

3.3.6 Interprétation des résultats des différents scénarios

La comparaison des taux de classification avec l'étalon démontre que l'utilisation d'un seul récif pour constituer le groupe d'entraînement réduit, règle générale, de manière significative la performance du système. Certaines combinaisons de récif d'entraînement et de test sont tout de même prometteuses. Cependant, un récif qui entraîne un SVM qui classe correctement un second récif n'est pas nécessairement bien classifié par un SVM entraîné à partir de ce même récif. Il n'y a pas de lien de réciprocité. Donc, il ne peut pas y avoir généralisation de l'utilisation de récifs uniques quelconques comme modèles de SVM. Certains sont compatibles, alors que d'autres ne le sont pas du tout, comme les récifs Martin et No Name 2012-2013.

Au niveau des groupes benthiques, la performance varie d'un couple entraînement-test à l'autre. Parfois, la réduction du taux de classification est uniforme sur toutes les classes. Dans d'autres cas, le taux de classification de certaines classes varie énormément, alors que celui

des autres reste constant. Ainsi, l'utilisation de récifs aléatoires à l'entraînement peut permettre d'introduire de l'information qui n'était pas disponible dans le récif initial. En contrepartie, elle peut également être manquante dans le récif d'entraînement choisi, ce qui explique les baisses importantes de performance. La corrélation C_{comp} entre les familles de coraux présentes dans deux récifs est conséquemment un bon indicateur de la capacité du couple à performer. Par l'analyse des données de la figure 3.13, nous concluons qu'un C_{comp} très faible, soit sous 0,4, a pour résultat de grands déficits par rapport à l'étalon (>11 %). Au-delà de cette valeur, la relation n'est cependant pas linéaire. Une corrélation plus forte n'est donc pas garante d'un déficit faible.

L'utilisation des images d'un même récif à travers les années assure un C_{comp} supérieur à 0,8 comme le montre la figure 3.17. Par contre, la corrélation entre C_{comp} et le déficit du taux de classification par rapport à l'étalon est faible, soit 0,6. Encore une fois, cela montre que, bien qu'une corrélation minimale soit nécessaire entre les familles des récifs, une corrélation forte ne mène pas nécessairement à un haut taux de classification. Ce type de couple atteint d'ailleurs une performance similaire à l'utilisation d'un récif aléatoire. Les groupes benthiques sont influencés de la même façon, sans égard au nombre d'années qui séparent les groupes d'entraînement et de test. Les écarts par rapport à l'étalon sont également équivalents. Bref, ce type de représentation assure une cohérence minimale des images, sans plus.

La combinaison de trois périodes d'un même récif pour faire l'entraînement ne donne pas nécessairement de meilleurs résultats que l'utilisation d'une seule de ces périodes. Par contre, il y a amélioration en moyenne : il est peu probable qu'une période individuelle soit plus performante qu'un regroupement. Lors de la combinaison de deux, trois ou quatre récifs quelconques pour former le groupe d'entraînement, chaque ajout fait augmenter le taux de classification d'environ 3 % lorsque le récif atteint un C_{comp} minimum de 0,4 avec le récif test.

Le cas où sept récifs sont utilisés pour l'entraînement du SVM exploite tous les récifs disponibles dans la base de données pour une même période d'échantillonnage. Cette situation améliore le taux de classification par rapport à tous les autres cas. Comme cela est montré à

la figure 2.22, le cas de No Name 2010-2011 ne s'approche cependant jamais de moins de 5 % de son étalon. Il est considéré difficile à classer, tout comme No Name 2012-2013 et Martin 2012-2013, qui sont tous au-dessus de la barre de 2 % d'écart avec l'étalon. Globalement, outre ces cas problématiques, l'approche de l'entraînement par plusieurs récifs quelconques tend vers l'étalon.

Conséquemment, l'utilisation de grands groupes d'entraînement augmente la performance de la classification et les récifs de périodes différentes de celle du groupe test forment des modèles de SVM aussi valables que ceux d'une même année. Puisqu'annoter de nouvelles images est coûteux en temps, nous favorisons l'alternative d'utiliser tous les récifs déjà annotés pour faire la classification. Ils correspondent donc aux images des années précédant le groupe test. Les résultats sont alors représentatifs d'un cas réel de classification de coraux.

Comme le montre la figure 3.21, les résultats s'améliorent avec la quantité de données utilisées pour faire l'entraînement et ils tendent vers l'étalon. Dans les cas les plus complets, soit l'entraînement sur tous les récifs de 2006-2007, 2008-2009 et 2010-2011 pour le test de tous les récifs de 2012-2013, puis l'entraînement sur 2008-2009, 2010-2011 et 2012-2013 pour le test sur 2006-2007, le taux de classification surpasse celui de la validation croisée effectuée sur l'ensemble des récifs de 2012-2013. Le taux de rappel et la précision sont également au-dessus de l'objectif de 80 %. Conséquemment, cette stratégie de sélection de données est prometteuse pour une application subséquente aux autres images de la base de données de l'AIMS.

CONCLUSION

L'objectif principal de ce mémoire était de développer un outil effectif de classification d'images de récifs coralliens spécifiquement pour la base de données de l'AIMS. Nous avons donc fait la recherche des paramètres optimaux pour chacune des étapes d'un système de reconnaissance de forme, nous avons évalué sa performance maximale dans un cas idéal et nous l'avons reproduite dans un contexte d'application réel.

L'étude a débuté, dans le chapitre 1, par la définition des caractéristiques des coraux, de la base de données de l'AIMS et des algorithmes de reconnaissance de forme. Nous avons ainsi remarqué que les coraux sont des structures colorées et riches en textures, mais que leur apparence diffère beaucoup à l'intérieur d'une même famille, alors que des individus de familles différentes peuvent être très similaires. Les images sous-marines sont de plus affectées par plusieurs artéfacts tels que la déficience en rouge et les flous. Toutefois, dans le cas particulier de la base de données de l'AIMS, les images sont acquises par un protocole strict et elles sont donc composées uniformément. Les algorithmes de reconnaissance de forme et les cinq étapes qui les composent doivent ainsi tenir compte des contraintes de ces dernières.

Le chapitre 1 présente donc divers prétraitements éprouvés pour compenser la faible qualité des images et il énumère des méthodes de segmentation permettant de distinguer les objets entre eux. Il inventorie également plusieurs techniques d'extraction de descripteurs correspondant à la fine pointe dans le domaine de la classification de coraux et nous y comparons divers classificateurs. SVM apparaît d'ailleurs comme étant l'un des mieux adaptés à la problématique et l'un des plus malléables. Les renseignements compris dans le chapitre 1 ont ensuite servi de guide pour construire l'algorithme de classification de données exploité aux chapitres 2 et 3.

Le second chapitre a servi à optimiser chaque étape du système pour les images de récifs coralliens de l'AIMS. Pour ce faire, nous avons utilisé des techniques et des paramètres connus dans la littérature comme référence pour la recherche. Nous avons donc commencé

par tester les divers prétraitements qui auraient pu améliorer la qualité des images. Il s'est cependant avéré que ces derniers ne menaient à aucune augmentation significative de la performance de l'algorithme. Ensuite, nous avons utilisé une méthode de segmentation simple et facile à mettre en place : le découpage de carrés autour du pixel identifié. Nous avons ainsi déterminé par itérations une plage idéale de longueurs des côtés des carrés pour la base de données de l'AIMS et nous avons choisi à l'intérieur de cette plage une valeur de 350 pixels puisqu'elle correspond approximativement à un maximum de performance.

Nous avons également effectué une pré-optimisation des paramètres du noyau RBF du SVM de manière à réduire le temps de son entraînement. En effet, la différence de composition entre le groupe d'entraînement et le groupe de test engendre plus de variations de performance qu'une optimisation précise des paramètres. Conséquemment, certaines valeurs prédéterminées de ces derniers offrent une performance équivalente à celle issue de leur optimisation au groupe traité.

Par la suite, nous avons extrait une série de descripteurs de texture et de couleur correspondant à la fine pointe de la technique. Nous avons répété le procédé à plusieurs échelles et à travers plusieurs canaux de couleur pour tenir compte de la grande richesse d'information contenue dans les images de récifs coralliens. Puis, nous avons choisi un sous-ensemble de valeurs utiles et non-redondantes parmi toutes celles générées grâce à un filtre basé sur la corrélation et à un algorithme glouton. Le temps d'extraction des descripteurs est alors significativement plus long que la fine pointe de la technique, mais le vecteur réduit est plus performant de 3 % sur un échantillon de la base de données de l'AIMS. Le système construit au cours de cette optimisation permet donc de classer ces images de manière aussi efficiente que possible selon les normes actuelles.

Le chapitre 3 avait quant à lui pour but d'évaluer les limites du système de classification optimisé, puis de faire une généralisation efficace de son application à la base de données. Nous avons circonscrit les bornes en étudiant un cas idéal : la validation croisée à dix ensembles sur des sous-ensembles uniformes de données. Ainsi, toute l'information pertinente à la clas-

sification était disponible pour entraîner les SVM. En répétant le processus de manière à classer toutes les données, nous avons en moyenne atteint un taux de rappel de 83,7 % sur quatre classes. Nous avons donc obtenu une performance supérieure à la cible de 80,0 % de l'AIMS au niveau des groupes benthiques. Le système optimisé permet ainsi une qualité de classification surpassant les attentes initiales dans un modèle idéal.

L'application du système de classification à la base de données de l'AIMS a ensuite passé par l'élaboration d'une stratégie de sélection de données pour faire l'entraînement du SVM : il était nécessaire de reproduire un cas réel de classification des images. Nous avons donc testé diverses méthodes de constitution des groupes d'entraînement et de test telles que l'utilisation de données appartenant à la même période temporelle ou, au contraire, l'utilisation de données appartenant au même emplacement spatial, mais séparées temporellement. L'ensemble des simulations a montré que l'entraînement du SVM à l'aide d'un ensemble large et diversifié permet une meilleure classification que l'utilisation de certaines combinaisons de récifs qui devraient intuitivement être compatibles.

D'un point de vue global, l'approche la plus performante pour l'ensemble des données est d'utiliser un large éventail de valeurs étiquetées au cours des années antérieures ou postérieures pour constituer le groupe d'entraînement. Par cette méthode, nous réussissons à atteindre les valeurs se rapprochant le plus de celles des validations croisées, ou même des valeurs qui leur sont supérieures. Ainsi, dans les cas réalistes, nous obtenons des taux de classification sur sept classes de 80,3 % pour tous les récifs de 2012-2013, de 78,2 % pour tous ceux de 2010-2011, de 77,7 % pour 2008-2009 et de 79,9 % pour 2006-2007. En éliminant une grande part de la confusion par la fusion de certaines classes clé, cela équivaut respectivement à 82,7 %, 80,7 %, 80,2 % et 81,0 %. Nous obtenons également un gain d'environ 1,2 % lorsque les récifs peu performants sont retirés. Encore une fois, l'objectif d'avoir un taux de rappel d'au moins 80 % est surpassé.

Tous les objectifs de ce mémoire ont ainsi été abordés : nous avons fixé les paramètres optimaux d'un outil de classification des images de récifs coralliens et nous avons développé une

stratégie d'utilisation de cet outil à partir de la base de données de l'AIMS. Ces expériences ont permis d'explorer les particularités de la reconnaissance de forme appliquée à des images naturelles, qui diffèrent de celles des applications typiques, telles que la reconnaissance de visage ou la reconnaissance dans le domaine biomédical. En effet, les images de coraux ont pour particularité d'être denses en information et d'être complexes. Elles exigent ainsi une approche adaptée par rapport au débalancement des classes, au choix des techniques de pré-traitement, de segmentation, d'extraction des descripteurs et de réduction de la dimensionnalité, puis au choix des paramètres du classificateur.

En utilisant ce système de classification, les biologistes marins de l'AIMS pourront plus facilement étiqueter d'anciennes et de nouvelles images de récifs de la Grande Barrière de Corail et ils pourront effectuer un suivi plus précis de leur santé. Cet outil permettra d'obtenir des statistiques à l'intérieur d'un très court intervalle de temps après l'acquisition des données, ce qui accélérera l'élaboration des stratégies de sauvegarde de l'environnement. Il libérera des ressources et les rendra disponible à l'analyse du cœur des problèmes.

RECOMMANDATIONS

Les expérimentations sur le système de reconnaissance de forme ont permis de distinguer plusieurs facteurs ayant un effet significatif tant sur la vitesse d'exécution de l'algorithme que sur la qualité des résultats obtenus. Ces observations permettent donc de formuler les recommandations suivantes :

- L'acquisition des données par un protocole strict permet d'obtenir des images uniformément constituées. Les éléments des images sont alors plus faciles à associer lors de la classification et cela rend l'utilisation de prétraitement superflue. Le temps de traitement des images est donc diminué et une meilleure performance est garantie.
- La segmentation par le découpage de carrés autour du pixel étiqueté doit inclure suffisamment d'information pour bien caractériser l'image, mais elle doit également éviter de noyer l'information pertinente. Bien que la plage acceptable de longueurs du côté des carrés soit large, il est important de ne pas y déroger puisque la baisse de performance pourrait alors être significative.
- Les descripteurs de (Shihavuddin et al., 2013) offrent une bonne performance et ils permettent une extraction rapide des caractéristiques des images. Toutefois, l'inclusion de dimensions supplémentaires d'échelles et de canaux de couleur améliore le taux de classification. La sélection de sous-ensembles de descripteurs permet donc d'obtenir un vecteur équivalent en longueur, mais plus performant que la fine pointe de la technologie.
- Pour une distribution à peu près fixe des classes, les valeurs des paramètres du noyau RBF du classificateur SVM peuvent être fixées. La réduction de la plage d'optimisation réduit le temps d'entraînement des divers SVM et le rend négligeable.

- Le groupe d'entraînement du classificateur doit être représentatif de l'ensemble des échantillons de la base de données. Conséquemment, le meilleur groupe est grand et diversifié. Ce faisant, le taux d'erreur est minimisé et les résultats sont répétables.

Grâce à l'application de ces règles, le système de classification de coraux répond aux critères de performance posés, mais nous pourrions encore améliorer plusieurs de ses aspects. Un peaufinage de certains paramètres pourrait ainsi le rendre plus malléable et lui permettre d'offrir des résultats plus précis.

Tout d'abord, il serait utile que le système distingue les images à des niveaux taxonomiques fins. Les informations extraites seraient alors plus riches et plus utiles aux biologistes marins. De plus, en exploitant des étiquettes plus précises, nous augmenterions le nombre de classes et nous pourrions alors réduire la variabilité intraclasse. Cependant, il est possible que la sélection de descripteurs doive être ajustée en fonction des classes qu'il faut distinguer. En effet, elles peuvent faire varier l'information discriminante.

Ensuite, il pourrait être pertinent de changer le critère d'optimisation du SVM. En utilisant l'exactitude, les classes les plus grandes sont généralement favorisées par rapport aux autres. L'intégration de *F-measure* permettrait d'obtenir les meilleurs résultats possibles pour une classe en particulier. Ce serait une alternative appropriée si une classe avait plus d'importance que les autres pour les biologistes. Les statistiques par rapport à cette classe seraient alors plus fiables. Le processus pourrait également être répété pour différentes classes et il mènerait à un outil modulable selon les besoins des experts.

Finalement, il serait pertinent de développer un processus de segmentation adapté aux particularités des coraux. En effet, la principale difficulté n'est pas dans la distinction des objets entre eux, mais plutôt dans l'association de l'étiquette à la structure précise qu'elle identifie. Conséquemment, nous optons généralement pour des techniques simples, comme le découpage de carrés autour de l'étiquette, qui conservent sciemment des corps étrangers dans

l'image. Elles induisent ainsi automatiquement un biais dans la représentation des objets, mais celui-ci est moindre que si nous sélectionnions la mauvaise structure.

Pour compenser cette contrainte, il serait possible de développer une technique de segmentation applicable lors de l'étiquetage manuel des images par les biologistes. Notre équipe œuvre à adapter une méthode de segmentation de type lasso qui pourra être utilisée à cet effet. La détermination de la structure ne serait alors plus une contrainte et les modèles de chaque classe seraient significativement plus uniformes. Par cette démarche, une importante part de la confusion interclasse serait éliminée.

ANNEXE I

STATISTIQUES APPLIQUÉES À LA MATRICE DE COOCCURRENCE

Le tableau-A I-1 contient les statistiques appliquées sur les matrices de cooccurrences normalisées $C(i, j)$ où i est le numéro de la ligne et, j , celui de la colonne. i et j prennent des valeurs de 1 à N , où N représente le nombre de niveaux de gris pris en compte lors de la génération des matrices. Les opérations de bases décrites par les équations-A I-1 à I-14 sont exploitées par plusieurs statistiques du tableau. Elles sont basées sur des variantes des matrices de cooccurrence normalisées $C(i, j)$, de la moyenne μ , de l'écart-type S et de l'entropie H .

$$C_x(i) = \sum_j C(i, j) \quad (\text{A I-1})$$

$$C_y(j) = \sum_i C(i, j) \quad (\text{A I-2})$$

$$C_{x+y}(k) = \sum_{i=1}^k C(i, k - i + 1) \mid k \leq N \quad (\text{A I-3})$$

$$C_{x-y}(k) = \begin{cases} \sum_{i=1}^{N-k+1} C(i, i + k - 1) + C(i + k - 1, i) \mid 1 < k < N \\ \sum_{i=1}^N C(i, j) \mid k = 1 \end{cases} \quad (\text{A I-4})$$

$$\mu = \frac{1}{N} \sum_{i,j} C(i, j) \quad (\text{A I-5})$$

$$\mu_x = \sum_{i,j} i C(i, j) \quad (\text{A I-6})$$

$$\mu_y = \sum_{i,j} j C(i,j) \quad (\text{A I-7})$$

$$S_x = - \sum_{i,j} (i - \mu_x)^2 C(i,j) \quad (\text{A I-8})$$

$$S_y = - \sum_{i,j} (j - \mu_y)^2 C(i,j) \quad (\text{A I-9})$$

$$H_{xy} = \sum_{i,j} C(i,j) \log C(i,j) \quad (\text{A I-10})$$

$$H_{xy1} = - \sum_{i,j} C(i,j) \log (C_x(i) C_y(j)) \quad (\text{A I-11})$$

$$H_{xy2} = - \sum_{i,j} C_x(i) C_y(j) \log (C_x(i) C_y(j)) \quad (\text{A I-12})$$

$$H_x = - \sum_i C_x(i) \log C_x(i) \quad (\text{A I-13})$$

$$H_y = - \sum_j C_y(j) \log C_y(j) \quad (\text{A I-14})$$

Tableau-A I-1 Statistiques associées à la matrice de cooccurrence

Adapté de (Shihavuddin et al., 2013)

	Statistique	Formules
1	Autocorrélation	$AC = \sum_{i,j} i \cdot j \cdot C(i,j)$
2	Contraste	$\sum_{(i,j)} C(i,j) i - j ^2$
3	Corrélation 1	$\sum_{i,j} \frac{(i - \mu_x)(j - \mu_y)C(i,j)}{S_x S_y}$
4	Corrélation 2	$\frac{AC - \mu_x \mu_y}{S_x S_y}$
5	Différence inverse	$\sum_{i,j} \frac{C(i,j)}{1 + i - j }$
6	Différence inverse normalisée	$\sum_{i,j} \frac{C(i,j)}{1 + \frac{ i - j }{N}}$
7	Dissimilarité	$\sum_{i,j} C(i,j) i - j $
8	Énergie	$\sum_{i,j} C(i,j)^2$
9	Entropie	$\sum_{i,j} C(i,j) \log C(i,j)$
10	Entropie de la différence	$-\sum_{k=1}^N C_{x-y}(k) \log C_{x-y}(k)$
11	Entropie de la somme	$S_e = -\sum_{k=1}^{2N-1} C_{x+y}(k) \log C_{x+y}(k)$
12	Mesure de la corrélation 1	$\frac{H_{xy} - H_{xy1}}{\max(H_x, H_y)}$
13	Mesure de la corrélation 2	$(1 - e^{-2(H_{xy2} - H_{xy})})^{0,5}$

	Statistique	Formules
14	Moment de la différence inverse	$\sum_{i,j} \frac{C(i,j)}{1 + i - j ^2}$
15	Moment de la différence inverse normalisée	$\sum_{i,j} \frac{C(i,j)}{1 + \frac{(i-j)^2}{N^2}}$
16	Moyenne de la somme	$\sum_{k=1}^{2N-1} (k+1)C_{x+y}(k)$
17	Ombre de l'amas schématique	$\sum_{i,j} (i+j - \mu_x - \mu_y)^3 C(i,j)$
18	Probabilité maximale	$\max\{C(i,j) \forall (i,j)\}$
19	Protubérance de l'amas schématique	$\sum_{i,j} (i+j - \mu_x - \mu_y)^4 C(i,j)$
20	Variance de la différence	$\sum_{k=0}^{N-1} k^2 C_{x-y}(k+1)$
21	Variance de la somme	$\sum_{k=1}^{2N-1} (k+1 - S_e)^2 C_{x+y}(k)$
22	Variance de la somme des carrés	$\sum_{i,j} (i - \mu)^2 C(i,j)$

ANNEXE II

STATISTIQUES APPLIQUÉES AUX HISTOGRAMMES DE TONS DE GRIS

Les statistiques de la moyenne, de l'écart-type, de l'asymétrie et du kurtosis, qui sont présentés dans les équations 1.6 à 1.9, sont les premières appliquées aux histogrammes de tons de gris $H(k)$, où k est le niveau d'intensité des pixels. Les statistiques supplémentaires sont présentées dans les équations-A II-1 à II-3.

$$Probabilité\ maximale = \max\{H(k) \forall k\} \quad (\text{A II-1})$$

$$Énergie = \sum_k H(k)^2 \quad (\text{A II-2})$$

$$Entropie = - \sum_k H(k) \log H(k) \quad (\text{A II-3})$$

ANNEXE III

BASE DE DONNÉES DE L'AIMS

Le tableau-A III-1 donne la répartition des échantillons de la base de données de l'AIMS.

Tableau-A III-1 Quantité d'échantillons pour chaque récif et chaque période de la base de données de l'AIMS

Récif	Année	Quantité totale	Algues	Coraux mous	Coraux durs	Abiotiques	Éponges	Millépores	Autres
Tous	2012-2013	23999	13017	1717	7287	1153	222	167	436
Tous	2010-2011	23994	11496	1690	8718	1292	213	142	443
Tous	2008-2009	24006	12291	1382	7775	1648	198	165	547
Tous	2006-2007	4955	2858	198	1378	381	32	34	74
Martin	2012-2013	3000	1422	197	905	413	26	1	36
Martin	2010-2011	3000	1353	133	958	513	25	4	14
Martin	2008-2009	3000	1311	126	956	562	16	5	24
Martin	2006-2007	200	60	7	59	69	2	0	3
Carter	2012-2013	3000	1814	109	908	0	17	69	83

Récif	Année	Quantité totale	Algues	Coraux mous	Coraux durs	Abiotiques	Éponges	Millépores	Autres
Carter	2010-2011	3000	1446	150	1264	1	16	42	81
Carter	2008-2009	3010	1593	115	1086	4	22	40	150
Carter	2006-2007	955	555	40	323	0	7	10	20
Lizard	2012-2013	3000	1432	593	758	102	66	8	40
Lizard	2010-2011	3000	1504	547	752	75	61	17	43
Lizard	2008-2009	2999	1609	503	665	110	44	24	43
Lizard	2006-2007	400	245	31	65	47	5	1	6
Macgillivray	2012-2013	3000	1462	192	839	415	44	5	43
Macgillivray	2010-2011	3000	1386	205	809	482	44	9	64
Macgillivray	2008-2009	3000	1502	109	640	647	39	3	59
Macgillivray	2006-2007	200	79	0	54	62	5	0	0
No Name	2012-2013	3000	1593	311	961	9	14	25	87
No Name	2010-2011	3000	1358	346	1148	6	12	28	102
No Name	2008-2009	3000	1503	243	1067	5	19	49	114
No Name	2006-2007	600	396	43	133	5	3	7	13

Récif	Année	Quantité totale	Algues	Coraux mous	Coraux durs	Abiotiques	Éponges	Millépores	Autres
North direction	2012-2013	3000	1570	73	1087	177	39	2	52
North direction	2010-2011	3000	1649	67	1029	187	21	2	43
North direction	2008-2009	3000	1734	43	870	276	39	2	35
North direction	2006-2007	1200	641	37	333	165	9	2	13
Yonge	2012-2013	3000	2145	79	657	2	12	48	57
Yonge	2010-2011	3000	1545	100	1238	0	24	25	68
Yonge	2008-2009	3000	1754	105	1005	1	11	27	97
Yonge	2006-2007	1200	829	33	307	1	0	13	17
Linnet	2012-2013	3000	1579	163	1172	35	4	9	38
Linnet	2010-2011	3000	1255	142	1520	28	10	15	28
Linnet	2008-2009	3000	1285	138	1486	43	8	15	25
Linnet	2006-2007	200	53	7	104	32	1	1	2

ANNEXE IV

RÉSULTATS DE L'OPTIMISATION DES PARAMÈTRES DU NOYAU RBF DU SVM

Le tableau-A IV-1 et IV-2 illustrent les résultats obtenus lors des optimisations des paramètres γ et C du noyau RBF des SVM selon diverses méthodes.

Tableau-A IV-1 Résultats de l'optimisation des paramètres γ et C du noyau RBF des SVM en gardant le récif constant et en faisant varier les années utilisées pour l'optimisation. Le groupe testé est toujours l'année 2012-2013 du récif d'entraînement

Récif traité	Année d'entraînement	Année d'optimisation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégories	Différence de TC avec l'optimisation sur l'année d'entraînement [%]	Moyenne des différences [%]	Écart-type des différences [%]
Martin	2006-2007	2006-2007	61,02	54,77	3	-	13,40	3,00
	2006-2007	2008-2009	71,02	65,58	3	10,00		
	2006-2007	2010-2011	76,68	70,53	3	15,66		
	2006-2007	2012-2013	75,55	70,09	3	14,53		
	2008-2009	2008-2009	77,73	66,24	4	-	1,36	2,08

Récif traité	Année d'en- trainement	Année d'opti- misation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégo- ries	Différence de TC avec l'optimisation sur l'année d'entraînement [%]	Moyenne des diffé- rences [%]	Écart-type des diffé- rences [%]
	2008-2009	2006-2007	76,71	62,70	4	-1,02		1,70
	2008-2009	2010-2011	80,60	68,04	4	2,87		
	2008-2009	2012-2013	79,95	68,04	4	2,22		
	2010-2011	2010-2011	75,04	65,22	4	-	-2,24	
	2010-2011	2006-2007	71,02	58,91	4	-4,02		
	2010-2011	2008-2009	72,97	62,46	4	-2,07		
	2010-2011	2012-2013	74,40	64,53	4	-0,64		
Yonge	2006-2007	2006-2007	88,83	78,50	2	-	0,12	0,11
	2006-2007	2008-2009	88,90	79,76	2	0,07		
	2006-2007	2010-2011	88,87	79,42	2	0,04		
	2006-2007	2012-2013	89,08	79,99	2	0,25		
	2008-2009	2008-2009	87,26	58,41	3	-	-0,22	0,23
	2008-2009	2006-2007	86,78	55,40	3	-0,48		
	2008-2009	2010-2011	87,23	57,92	3	-0,03		
	2008-2009	2012-2013	87,12	59,36	3	-0,14		

Récif traité	Année d'entraînement	Année d'optimisation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégories	Différence de TC avec l'optimisation sur l'année d'entraînement [%]	Moyenne des différences [%]	Écart-type des différences [%]
	2010-2011	2010-2011	87,19	58,17	3	-	0,00	0,18
	2010-2011	2006-2007	87,09	55,40	3	-0,10		
	2010-2011	2008-2009	87,40	58,51	3	0,21		
	2010-2011	2012-2013	87,09	58,46	3	-0,10		
Carter	2006-2007	2006-2007	82,48	77,20	2	-	-0,03	0,05
	2006-2007	2008-2009	82,48	77,58	2	0,00		
	2006-2007	2010-2011	82,40	77,23	2	-0,08		
	2006-2007	2012-2013	82,48	77,58	2	0,00		
	2008-2009	2008-2009	78,52	49,54	4	-	-0,63	0,55
	2008-2009	2006-2007	77,63	49,28	4	-0,89		
	2008-2009	2010-2011	77,52	48,11	4	-1,00		
	2008-2009	2012-2013	78,52	49,54	4	0,00		
	2010-2011	2010-2011	81,46	64,77	3	-	-0,01	0,27
	2010-2011	2006-2007	81,14	64,53	3	-0,32		
	2010-2011	2008-2009	81,60	65,24	3	0,14		

Récif traité	Année d'entraînement	Année d'optimisation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégories	Différence de TC avec l'optimisation sur l'année d'entraînement [%]	Moyenne des différences [%]	Écart-type des différences [%]
	2010-2011	2012-2013	81,60	65,24	3	0,14		
Résultats globaux							-0,21	0,65

Tableau-A IV-2 Résultats de l'optimisation des paramètres γ et C du noyau RBF des SVM en utilisant un ou deux récifs du groupe d'entraînement. Les groupes d'entraînement et de test sont toujours l'année 2012-2013 des récifs

Récifs d'entraînement	Récif de test	Récif d'optimisation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégories	Différence de TC avec l'optimisation sur l'année d'entraînement [%]	Moyenne des différences [%]	Écart-type des différences [%]
Carter No Name	Martin	Carter No Name	67,35	60,05	3	-	0,60	1,40
Carter	Martin	Carter	68,94	58,08	3	1,59		

Récifs d'entraînement	Récif de test	Récif d'optimisation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégories	Différence de TC avec l'optimisation sur l'année d'entraînement [%]	Moyenne des différences [%]	Écart-type des différences [%]
No Name								
Carter No Name	Martin	No Name	66,96	58,98	3	-0,39		
Carter Yonge	Martin	Carter Yonge	81,74	82,91	2	-	-0,19	0,76
Carter Yonge	Martin	Carter	82,08	83,27	2	0,34		
Carter Yonge	Martin	Yonge	81,01	82,31	2	-0,73		
Martin No Name	Carter	Martin No Name	80,38	70,63	3	-	0,20	0,02
Martin No Name	Carter	Martin	80,56	70,42	3	0,18		
Martin No Name	Carter	No Name	80,59	73,46	3	0,21		

Récifs d'entrainement	Récif de test	Récif d'optimisation des paramètres du noyau RBF du SVM	TC [%]	TC moyen par classe [%]	Quantité de catégories	Différence de TC avec l'optimisation sur l'année d'entrainement [%]	Moyenne des différences [%]	Écart-type des différences [%]
Résultats globaux							0,20	0,73

ANNEXE V

RÉSULTATS DES VALIDATIONS CROISÉES

Les tableaux-A V-1, V-2 et V-3 contiennent les résultats de validations croisées effectuées sur différents groupes de la base de données de l'AIMS.

Tableau-A V-1 Validations croisées effectuées sur un récif au cours d'une seule période (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Récif	Période	TC [%]	Incertitude sur le TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (AL, CM, CD, AB, E, M, AU)
Martin	2012-2013	82,83	0,66	45,50	7	91, 43, 86, 82, 0, 0, 0
Martin	2010-2011	81,65	0,43	41,23	7	87, 29, 84, 82, 0, 0, 0
Martin	2008-2009	79,43	0,74	40,34	7	87, 27, 82, 76, 13, 0, 0
Martin	2006-2007	80,10	4,21	42,59	7	83, 0, 88, 88, 0, -, 0
Carter	2012-2013	81,20	0,21	47,32	7	92, 50, 71, -, 0, 72, 0
Carter	2010-2011	78,94	0,74	37,24	7	88, 50, 78, 0, 0, 50, 1
Carter	2008-2009	77,04	-	34,71	7	91, 38, 74, 0,0, 35, 5

Récif	Période	TC [%]	Incertitude sur le TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (AI, CM, CD, AB, E, M, AU)
Carter	2006-2007	78,53	-	36,14	7	92, 57, 68, -, 0, 0, 0
Lizard	2012-2013	78,49	0,63	44,69	7	91, 76, 70, 72, 0, 0, 8
Lizard	2010-2011	78,12	0,81	42,95	7	93, 71, 68, 63, 0, 0, 2
Lizard	2008-2009	77,19	0,55	40,99	7	91, 70, 63, 52, 5, 8, 12
Lizard	2006-2007	74,51	1,89	32,11	7	93, 23, 51, 51, 0, 0, 0
Macgillivray	2012-2013	82,15	0,80	43,99	7	92, 53, 82, 80, 0, 0, 0
Macgillivray	2010-2011	80,53	0,45	48,42	7	92, 45, 80, 79, 0, 0, 39
Macgillivray	2008-2009	80,73	0,70	43,57	7	92, 28, 76, 79, 0, 0, 25
Macgillivray	2006-2007	85,00	1,41	65,43	7	89, -, 93, 84, 0, -, -
No Name	2012-2013	76,40	0,47	32,53	7	92, 49, 69, 0, 0, 20, 0
No Name	2010-2011	76,54	0,48	34,30	7	89,58, 76, 0, 0, 14, 7
No Name	2008-2009	75,88	0,65	33,91	7	90, 47, 73, 0, 0, 24, 2
No Name	2006-2007	79,53	1,28	28,86	7	97, 49, 59, 0, 0, 0, 0
North direction	2012-2013	80,93	0,31	37,93	7	89, 22, 81, 62, 0, 0, 8
North direction	2010-2011	82,33	0,53	35,12	7	90, 16, 84, 56, 0, 0, 0

Récif	Période	TC [%]	Incertitude sur le TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (AI, CM, CD, AB, E, M, AU)
North direction	2008-2009	80,93	0,57	34,56	7	91, 12, 77, 65, 0, 0, 3
North direction	2006-2007	76,98	0,88	34,38	7	88, 11, 78, 62, 0, 0, 0
Yonge	2012-2013	86,86	0,38	38,82	7	96, 47, 72, 0, 0, 46, 7
Yonge	2010-2011	82,22	0,59	41,34	7	92, 43, 81, -, 0, 32, 1
Yonge	2008-2009	82,27	0,58	39,26	7	93, 50, 77, 0, 0, 56, 5
Yonge	2006-2007	84,68	0,84	32,22	7	96, 27, 69, 0, -, 0, 0
Linnet	2012-2013	84,26	0,50	42,20	7	90, 49, 86, 40, 0, 11, 13
Linnet	2010-2011	82,38	0,28	41,11	7	86, 43, 86, 75, 0, 0, 0
Linnet	2008-2009	76,94	0,77	34,48	7	82, 29, 81, 48, 0, 0, 0
Linnet	2006-2007	71,70	1,67	29,51	7	53, 0, 88, 50, 0, 0, 0

Tableau-A V-2 Validations croisées effectuées sur tous les récifs au cours d'une seule période (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Récifs	Année	TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)
Tous	2006-2007	79,72	40,09	7	91, 38, 74, 63, 0, 12, 3
Tous	2008-2009	82,08	50,84	7	93, 64, 67, 75, 0, 24, 20
Tous	2010-2011	80,38	48,45	7	93, 68, 75, 78, 0, 0, 25
Tous	2012-2013	79,31	46,13	7	93, 60, 78, 74, 1, 44, 7

Tableau-A V-3 Validations croisées effectuées sur un récif au cours de toutes les périodes (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Récif	Année	TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)
Martin	Toutes	81,78	44,14	7	88, 42, 84, 80, 6, 0, 9
Carter	Toutes	79,94	39,51	7	92, 52, 75, 0, 0, 54, 5
Yonge	Toutes	84,36	38,08	7	94, 50, 78, 0, 0, 42, 4

ANNEXE VI

RÉSULTATS DE LA GÉNÉRALISATION DE L'APPLICATION DE L'ALGORITHME

Les tableaux-A VI-1 à VI-8 contiennent les résultats de la généralisation de l'application de l'algorithme selon diverses contraintes.

Tableau-A VI-1 Résultats de classifications utilisant deux récifs à l'intérieur d'une période respectivement comme groupe d'entraînement et de test (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – mil-lépores, AU – autres)

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de caté- gories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Carter	2010-2011	No Name	2010-2011	70,57	26,89	7	90, 19, 72, 0, 0, 7, 0	0,99
Carter	2012-2013	Yonge	2012-2013	81,53	30,51	7	90, 27, 75, 0, 0, 21, 2	0,99
Martin	2012-2013	No Name	2012-2013	59,77	23,37	7	71, 11, 64, 44, 0, 0, 8	0,32
No Name	2010-2011	Carter	2010-2011	76,97	33,26	7	87, 59, 76, 0, 0, 10, 1	0,99
No Name	2012-2013	Martin	2012-2013	56,60	25,19	7	66, 34, 77, 65, 0, 0, 0	0,32

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Yonge	2012-2013	Carter	2012-2013	77,57	40,29	7	89, 17, 71, 0, 64, 0, 1	0,99
Linnet	2012-2013	Macgil-livray	2012-2013	73,03	32,01	7	90, 14, 81, 39, 0, 0, 0	0,91
Macgilli-vray	2012-2013	Yonge	2012-2013	75,40	22,40	7	86, 8, 63, 0, 0, 0, 0	0,31

Tableau-A VI-2 Résultats de classifications utilisant trois, quatre ou cinq récifs à l'intérieur d'une période. Un seul d'entre eux est mis dans le groupe test; les autres appartiennent au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Carter, No Name	2012-2013	Martin	2012-2013	57,33	26,34	7	72, 42, 68, 0, 0, 0, 3	0,27
Carter,	2012-2013	Martin	2012-2013	60,93	25,75	7	77, 28, 75, 0, 0, 0, 0	0,26

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Yonge								
Carter, No Name, Yonge	2012-2013	Martin	2012-2013	59,77	27,37	7	72, 37, 77, 0, 0, 0, 6	0,28
Martin, No Name	2010-2011	Carter	2010-2011	76,83	33,00	7	86, 64, 76, 0, 0, 2, 2	0,70
Macgillivray, Carter	2012-2013	Yonge	2012-2013	85,30	29,10	7	97, 24, 68, 0, 0, 15, 0	0,83
Macgillivray, Carter, Linnet	2012-2013	Yonge	2012-2013	85,63	29,81	7	97, 29, 71, 0, 0, 10, 2	0,67
Linnet, No Name	2012-2013	Macgil- livray	2012-2013	75,97	35,11	7	92, 17, 79, 59, 0, 0, 0	0,78
Linnet, No Name, Lizard	2012-2013	Macgil- livray	2012-2013	79,03	40,18	7	92, 40, 78, 72, 0, 0, 0	0,86

Tableau-A VI-3 Résultats de classifications utilisant les huit récifs à l'intérieur d'une période. Un seul d'entre eux est mis dans le groupe test; les autres appartiennent au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)
Tous les récifs sauf Martin	2012-2013	Martin	2012-2013	79,37	44,26	7	87, 63, 85, 57, 0, 0, 17
Tous les récifs sauf Yonge	2012-2013	Yonge	2012-2013	86,10	33,77	7	96, 51, 74, 0, 0, 15, 57
Tous les récifs sauf Carter	2012-2013	Carter	2012-2013	79,50	30,48	7	92, 49, 72, 0, 0, 0, -
Tous les récifs sauf Lizard	2012-2013	Lizard	2012-2013	76,39	44,89	7	93, 58, 73, 58, 0, 25, 8
Tous les récifs sauf Macgillivray	2012-2013	Macgillivray	2012-2013	81,30	42,82	7	91, 44, 81, 84, 0, 0, 0
Tous les récifs sauf No Name	2012-2013	No Name	2012-2013	73,57	33,71	7	93, 27, 65, 22, 0, 24, 5
Tous les récifs sauf Linnet	2012-2013	Linnet	2012-2013	83,33	42,87	7	92, 64, 78, 40, 0, 22, 3
Tous les récifs sauf North Direction	2012-2013	North direction	2012-2013	78,83	36,16	7	91, 38, 76, 48, 0, 0, 0
Tous les récifs sauf No	2010-2011	No	2010-2011	71,20	26,42	7	89, 22, 74, 0, 0, 0, 0

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)
Name		Name					
Tous les récifs sauf Carter	2010-2011	Carter	2010-2011	77,97	33,40	7	88, 54, 78, 0, 0, 14, 0
Tous les récifs sauf Yonge	2010-2011	Yonge	2010-2011	82,00	38,28	7	90, 58, 82, 0, 0, 0
Tous les récifs sauf Linnet	2010-2011	Linnet	2010-2011	80,72	40,17	7	81, 49, 86, 61, 0, 0, 4
Tous les récifs sauf North direction	2011	North direction	2010-2011	82,82	41,39	7	94, 48, 77, 57, 5, 0, 9

Tableau-A VI-4 Résultats de classifications utilisant un seul récif à l'intérieur de deux périodes. Une période est mise dans le groupe test; l'autre appartient au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Carter	2008-2009	Carter	2006-2007	76,96	39,86	7	89, 43, 67, -, 0, 30, 10	0,98

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Carter	2010-2011	Carter	2006-2007	76,86	31,86	7	87, 30, 74, -, 0, 0, 0	0,97
Carter	2012-2013	Carter	2006-2007	72,57	32,08	7	76, 15, 81, -, 0, 20, 0	0,99
Carter	2006-2007	Carter	2008-2009	72,26	25,36	7	91, 14, 65, 0, 0, 8, 0	0,98
Carter	2010-2011	Carter	2008-2009	72,36	24,86	7	80,12, 82, 0, 0, 0, 0	0,99
Carter	2012-2013	Carter	2008-2009	66,58	24,34	7	70, 4, 81, 0, 0, 15, 0	0,97
Carter	2006-2007	Carter	2010-2011	74,43	30,92	7	88, 39, 70, 0, 0, 19, 0	0,97
Carter	2008-2009	Carter	2010-2011	74,03	33,62	7	91, 42, 66, 0, 0, 36, 1	0,99
Carter	2012-2013	Carter	2010-2011	73,53	29,40	7	89, 24, 69, 0, 0, 24, 0	0,94
Carter	2006-2007	Carter	2012-2013	74,73	36,13	7	93, 28, 55, -, 0, 41, 0	0,99
Carter	2008-2009	Carter	2012-2013	76,00	40,06	7	97, 46, 49, -, 0, 48, 1	0,97
Carter	2010-2011	Carter	2012-2013	77,43	39,26	7	91, 38, 66, 0, 41, 0	0,94
Martin	2008-2009	Martin	2006-2007	79,50	44,27	7	73, 14, 88, 90, 0, -, 0	0,85
Martin	2010-2011	Martin	2006-2007	71,50	37,63	7	50, 0, 83, 93, 0, -, 0	0,83
Martin	2012-2013	Martin	2006-2007	64,50	33,80	7	78, 0, 39, 86, 0, -, 0	0,81
Martin	2006-2007	Martin	2008-2009	64,47	29,17	7	64, 1, 87, 53, 0, 0, 0	0,85
Martin	2010-2011	Martin	2008-2009	75,37	34,11	7	79, 6, 89, 65, 0, 0, 0	1,00
Martin	2012-2013	Martin	2008-2009	70,93	32,79	7	88, 15, 69, 54, 0, 0, 4	0,98

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Martin	2006-2007	Martin	2010-2011	66,03	29,42	7	62, 0, 91, 53, 0, 0, 0	0,83
Martin	2008-2009	Martin	2010-2011	78,70	37,80	7	91, 22, 73, 79, 0, 0, 0	0,90
Martin	2012-2013	Martin	2010-2011	76,17	37,94	7	91, 32, 67, 72, 4, 0, 0	0,99
Martin	2006-2007	Martin	2012-2013	65,80	28,63	7	69, 0 91, 40, 0, 0, 0	0,81
Martin	2008-2009	Martin	2012-2013	78,73	40,02	7	83, 19, 90, 80, 0, 0, 8	0,99
Martin	2010-2011	Martin	2012-2013	73,57	37,32	7	70, 21, 95, 76, 0, 0, 0	0,98
Yonge	2008-2009	Yonge	2006-2007	82,75	35,92	7	88, 30, 82, 0, 15, 0	0,97
Yonge	2010-2011	Yonge	2006-2007	83,07	29,11	7	90, 3, 82, 0, -, 0, 0	0,95
Yonge	2012-2013	Yonge	2006-2007	80,00	35,95	7	86, 21, 78, 0, 31, 0	0,99
Yonge	2006-2007	Yonge	2008-2009	76,73	24,53	7	94, 10, 63, 0, 0, 4, 0	0,97
Yonge	2010-2011	Yonge	2008-2009	78,10	24,91	7	89, 9, 77, 0, 0, 0, 0	0,99
Yonge	2012-2013	Yonge	2008-2009	74,77	27,15	7	83, 5, 77, 0, 0, 22, 3	0,93
Yonge	2006-2007	Yonge	2010-2011	78,17	29,42	7	95, 10, 70, -, 0, 0, 1	0,95
Yonge	2008-2009	Yonge	2010-2011	78,93	31,40	7	95, 4, 71, -, 0, 8, 0	0,99
Yonge	2012-2013	Yonge	2010-2011	74,60	36,59	7	91, 24, 65, -, 0, 40, 0	0,90
Yonge	2006-2007	Yonge	2012-2013	83,13	25,24	7	98, 22, 57, 0, 0, 0, 0	0,99
Yonge	2008-2009	Yonge	2012-2013	83,73	24,88	7	97, 11, 63, 0, 0, 2, 0	0,93

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)	C_{comp}
Yonge	2010-2011	Yonge	2012-2013	83,87	25,15	7	98, 14, 62, 0, 0, 2, 0	0,90

Tableau-A VI-5 Résultats de classifications utilisant un seul récif à l'intérieur de quatre périodes. Une seule d'entre elles est mise dans le groupe test; les autres appartiennent au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)
Martin	2006-2007, 2008-2009, 2010-2011	Martin	2012-2013	76,50	38,32	7	77, 22, 93, 76,0, 0, 0
Martin	200807, 2010-2011, 2012-2013	Martin	2006-2007	78,00	45,54	7	67, 29, 88, 90, 0, -, 0
Carter	2006-2007,	Carter	2012-2013	79,20	44,46	7	94, 47, 64, -, 0, 61, 1

Groupe d'entraine- ment		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégo- ries	Diagonale de la matrice de confusion (AI, CM, CD, AB, E, M, AU)
	2008-2009, 2010-2011						
Carter	200807, 2010-2011, 2012-2013	Carter	2006-2007	79,37	41,29	7	89, 55, 73, -, 0, 30, 0
Yonge	2006-2007, 2008-2009, 2010-2011	Yonge	2012-2013	84,13	27,32	7	98, 25, 62, 0, 0, 6, 0
Yonge	200807, 2010-2011, 2012-2013	Yonge	2006-2007	84,25	43,16	7	90, 42, 81, 0, 46, 0

Tableau-A VI-6 Résultats de classifications utilisant les huit récifs à l'intérieur de deux périodes. Une période est mise dans le groupe test; l'autre appartient au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entraînement		Groupe de test		TC [%]	TC moyen par classe [%]	Quan- tité de catégo- ries	Diagonale de la ma- trice de confusion (AL, CM, CD, AB, E, M, AU)	C_{comp}
Tous les récifs	2006-2007	Tous les récifs	2010-2011	64,02	31,21	7	94, 37, 38, 50, 0, 0, 0	0,98
Tous les récifs	2008-2009	Tous les récifs	2010-2011	69,84	37,92	7	87, 48, 57, 72, 0, 0, 1	0,91
Tous les récifs	2012-2013	Tous les récifs	2010-2011	75,77	40,19	7	92, 39, 70, 65, 1, 11, 3	0,98
Tous les récifs	2006-2007	Tous les récifs	2012-2013	64,61	29,89	7	91, 41, 33, 41, 2, 0, 0	0,84
Tous les récifs	2008-2009	Tous les récifs	2012-2013	71,51	37,91	7	83, 48, 65, 67, 0, 0, 1	0,88
Tous les récifs	2010-2011	Tous les récifs	2012-2013	79,01	42,04	7	91, 54, 73, 67, 0, 7, 2	0,98

Tableau-A VI-7 Résultats de classifications utilisant les huit récifs à l'intérieur de plusieurs périodes. Une seule période est mise dans le groupe test; les autres appartiennent au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entrainement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (AL, CM, CD, AB, E, M, AU)
Tous les récifs	2006-2007, 2008-2009, 2010-2011	Tous les récifs	2012-2013	80,34	45,87	7	92, 60, 73, 71, 1, 17, 6
Tous les récifs	2006-2007, 2008-2009	Tous les récifs	2010-2011	77,08	42,19	7	92, 51, 70, 74, 0, 6, 2
Tous les récifs	2006-2007, 2008-2009, 2012-2013	Tous les récifs	2010-2011	78,23	44,87	7	92, 50, 72, 74, 0, 22, 4
Tous les récifs	2006-2007	Tous les récifs	2008-2009	62,99	28,09	7	94, 31, 33, 37, 2, 0, 0
Tous les récifs	2010-2011, 2012-2013	Tous les récifs	2008-2009	75,39	38,94	7	82, 34, 84, 57, 1, 6, 8
Tous les récifs	2006-2007, 2010-2011, 2012-2013	Tous les récifs	2008-2009	77,74	40,71	7	91, 41, 76, 59, 0, 14, 4

Groupe d'entrainement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégo- ries	Diagonale de la ma- trice de confusion (Al, CM, CD, AB, E, M, AU)
Tous les récifs	2008-2009, 2010-2011, 2012-2013	Tous les récifs	2006-2007	79,94	41,23	7	89, 41, 79, 62, 0, 15, 3

Tableau-A VI-8 Résultats de classifications utilisant les récifs Martin, Yonge, Macgillivray, North Direction et Linnet à l'intérieur de plusieurs périodes. Une seule période est mise dans le groupe test; les autres appartiennent au groupe d'entraînement (Légende : AL – algues, CM – coraux mous, CD – coraux durs, AB – abiotiques, E – éponges, M – millépores, AU – autres)

Groupe d'entrainement		Groupe de test		TC [%]	TC moyen par classe [%]	Quantité de catégories	Diagonale de la matrice de confusion (Al, CM, CD, AB, E, M, AU)
Tous les récifs	2006-2007, 2008-2009, 2010-2011	Tous les récifs sauf No Name, Carter et Lizard	2012-2013	81,56	45,10	7	92, 57, 76, 71, 01, 12, 7
Tous les récifs	2008-2009, 2010-2011, 2012-2013	Tous les récifs sauf No Name, Carter et Lizard	2006-2007	79,93	40,72	7	88, 32, 82, 65, 0, 19, 0

LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES

- Abdul Ghani, Ahmad Shahrizan, et Nor Ashidi Mat Isa. 2015. « Underwater image quality enhancement through integrated color model with Rayleigh distribution ». *Applied Soft Computing*, vol. 27, p. 219-230.
- Agence spatiale canadienne. 2012. « Des images prises depuis la station spatiale aident les nations insulaires à gérer les ressources des récifs coralliens ». < http://www.asc-csa.gc.ca/fra/iss/avantages_20_images.asp >. Consulté le 5 septembre 2014.
- Anand, Ashish, Ganesan Pugalenth, Gary B Fogel et P. N. Suganthan. 2010. « An approach for classification of highly imbalanced data using weighting and undersampling ». *Amino Acids*, vol. 39, n° 5, p. 1385-1391.
- Bankman, Isaac N. 2009. *Handbook of medical image processing and analysis*, 2nd ed. Amsterdam: Elsevier/Academic Press, 984 p.
- Beijbom, Oscar, Peter J Edmunds, David I Kline, B Greg Mitchell et David Kriegman. 2012. « Automated annotation of coral reef survey images ». In *IEEE Conference on Computer Vision and Pattern Recognition*. p. 1170-1177. IEEE.
- Bouchard, Jonathan. 2011. « Methodes de vision et d'intelligence artificielles pour la reconnaissance de specimens coralliens ». M.Ing. Montréal, Canada, École de Technologie Supérieure (Canada), 190 p. In ProQuest Dissertations & Theses Full Text. < <http://search.proquest.com/docview/929293255?accountid=27231> >.
- Bryson, M., M. Johnson-Roberson, O. Pizarro et S. Williams. 2013. « Automated Registration for Multi-year Robotic Surveys of Marine Benthic Habitats ». In *2013 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. (3-7 novembre 2013), p. 3344-3349.
- Chang, Chih-Chung, et Chih-Jen Lin. 2011. « LIBSVM : a library for support vector machines ». *ACM Transactions on Intelligent Systems and Technology*. Vol. 2, n° 3, p. 27:1-27:27. < <http://www.csie.ntu.edu.tw/~cjlin/libsvm> >.
- Chaudhury, Nandini Ray, et Ajai. 2014. « Impact of Climate Change on Coral Reefs ». In *Geospatial Technologies and Climate Change*, sous la dir. de Sundaresan, Janardhanan, K. M. Santosh, Andrea Déri, Rob Roggema et Ramesh Singh. Vol. 10,

- p. 37-52. Coll. « Geotechnologies and the Environment »: Springer International Publishing. < http://dx.doi.org/10.1007/978-3-319-01689-4_3 >.
- Duin, R.P.W., et D.M.J. Tax. 2005. « Statistical pattern recognition ». In *Handbook of pattern recognition and computer vision*, sous la dir. de Chen, C. H., et P.S.P. Wang, 3rd ed., p. 3-23. Singapore / Hackensack, NJ: World Scientific.
- Ebner, Marc. 2009. « Color constancy based on local space average color ». *Machine Vision and Applications*, vol. 20, n° 5, p. 283-301.
- Finlayson, Graham D., Bernt Schiele et James L. Crowley. 1998. « Comprehensive colour image normalization ». In *Computer Vision — ECCV'98*, sous la dir. de Burkhardt, Hans, et Bernd Neumann. Vol. 1406, p. 475-490. Coll. « Lecture Notes in Computer Science »: Springer Berlin Heidelberg. < http://download.springer.com/static/pdf/577/chp%3A10.1007%2Fbfb0055685.pdf?auth66=1409833227_1a54348a0fea296cf4a91adb8db3ad88&ext=.pdf >.
- Guo, Zhenhua, et David Zhang. 2010. « A completed modeling of local binary pattern operator for texture classification ». *IEEE Transactions on Image Processing*, vol. 19, n° 6, p. 1657-1663.
- Hall, Mark A. 1999. « Correlation-based feature selection for machine learning ». Hamilton, New Zealand, The University of Waikato, 178 p. < <http://www.cs.waikato.ac.nz/~mhall/thesis.pdf> >.
- Haralick, Robert M, Karthikeyan Shanmugam et Its' Hak Dinstein. 1973. « Textural features for image classification ». *IEEE Transactions on Systems, Man and Cybernetics*, vol. 3, n° 6, p. 610-621.
- Jähne, Bernd. 2004. « Scale and Texture ». In *Practical handbook on image processing for scientific and technical applications*, 2nd ed.. p. 443-472. Boca Raton [Fla.]: CRC Press. < <http://marc.crcnetbase.com/isbn/9780849390302> >.
- Jähne, Bernd. 2005. « Feature Extraction : Texture ». In *Digital image processing*, 6th rev. and ext. ed., p. 435-446. Berlin / New York: Springer. < <http://www.books24x7.com/marc.asp?bookid=16224> >.

- Jiang, Peng, Samy Missoum et Zhao Chen. 2014. « Optimal SVM parameter selection for non-separable and unbalanced datasets ». *Structural and Multidisciplinary Optimization*, vol. 50, n° 4, p. 523-535.
- Jonker, M., K. Johns et K. Osborne. 2008. *Long term Monitoring GBR Standard Operational Procedure*. Coll. « Surveys of benthic reef communities using underwater digital photography and counts of juvenile corals », Number 10. Townsville: Australian Institute of Marine Science, 85 p.
< <http://www.aims.gov.au/docs/research/monitoring/reef/technical-reports.html> >.
- Kudo, Mineichi, et Jack Sklansky. 2000. « Comparison of algorithms that select features for pattern classifiers ». *Pattern Recognition*, vol. 33, n° 1, p. 25-41.
- Litimco, C. E. O., M. G. A. Villanueva, N. G. Yecla, M. N. Soriano et P. C. Naval. 2013. « Coral Identification Information System ». In *2013 IEEE International Underwater Technology Symposium*. (New York). IEEE.
- Liu, Guangcan, Zhouchen Lin et Yong Yu. 2009. « Radon representation-based feature descriptor for texture classification ». *IEEE Transactions on image processing*, vol. 18, n° 5, p. 921-928.
- Machine Learning Group at the University of Waikato. 2014. « WEKA ».
- Manjunath, B. S., et W. Y. Ma. 1996. « Texture features for browsing and retrieval of image data ». *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 18, n° 8, p. 837-842.
- Marcos, Ma Shiela Angeli, Laura David, Eileen Penaflor, Victor Ticzon et Maricor Soriano. 2008. « Automated benthic counting of living and non-living components in Ngedarrak Reef, Palau via subsurface underwater video ». *Environmental Monitoring and Assessment*, vol. 145, n° 1-3, p. 177-184.
- Naval, P. C., M. Soriano, B. C. Esmero et Z. M. Abad. 2014. « A Coral Mapping and Health Assessment System Based on Texture Analysis ». In *Intelligent Information and Database Systems, Pt 1*, sous la dir. de Nguyen, N. T., B. Attachoo, B. Trawinski et K. Somboonviwat. Vol. 8397, p. 601-609. Coll. « Lecture Notes in Computer Science ». Berlin: Springer-Verlag Berlin.

- Ninio, R, John Steven Craig Delean, K Osborne et Hugh Sweatman. 2003. « Estimating cover of benthic organisms from underwater video images: variability associated with multiple observers ». *Marine Ecology-Progress Series*, vol. 265, p. 107-116.
- Novakovic, Jasmina. 2009. « Using information gain attribute evaluation to classify sonar targets ». In *17th Telecommunications forum TELFOR*. (Belgrade, Serbie, 24- 25 Novembre 2009), p. 24-26.
- Novakovic, Jasmina, et A. Veljovic. 2011. « C-Support Vector Classification: Selection of kernel and parameters in medical diagnosis ». In *Intelligent Systems and Informatics (SISY), 2011 IEEE 9th International Symposium on*. (8-10 Sept. 2011), p. 465-470. < <http://ieeexplore.ieee.org/ielx5/6027501/6034292/06034373.pdf?tp=&arnumber=6034373&isnumber=6034292> >.
- Pal, Sankar K ., et Pabitra Mitra. 2004. « Unsupervised Feature Selection ». In *Pattern Recognition Algorithms for Data Mining*. p. 59-82. Coll. « Chapman & Hall/CRC Computer Science & Data Analysis »: Chapman and Hall/CRC. < <http://dx.doi.org/10.1201/9780203998076.ch3> >. Consulté le 2014/12/08.
- Petrou, Maria, et Pedro Garcia Sevilla. 2006. *Image processing: dealing with texture*, 10. John Wiley and Sons, 634 p.
- Pizarro, O., P. Rigby, M. Johnson-Roberson, S. B. Williams et J. Colquhoun. 2008. « Towards Image-based Marine Habitat Classification ». In *OCEANS 2008*. (Ville de Québec, Québec, 15-18 septembre 2008), p. 1-7. IEEE.
- Shih, Frank Y. 2010. *Image processing and pattern recognition : fundamentals and techniques* (2010). Hoboken, N.J.: IEEE/Wiley, 537 p.
- Shihavuddin, A. S. M., N. Gracias, R. Garcia, A. C. R. Gleason et B. Gintert. 2013. « Image-Based Coral Reef Classification and Thematic Mapping ». *Remote Sensing*, vol. 5, n° 4, p. 1809-1841.
- Stokes, M. D., et G. B. Deane. 2009. « Automated processing of coral reef benthic images ». *Limnology and Oceanography-Methods*, vol. 7, p. 157-168.
- Sun, Yanmin, Mohamed S Kamel, Andrew KC Wong et Yang Wang. 2007. « Cost-sensitive boosting for classification of imbalanced data ». *Pattern Recognition*, vol. 40, n° 12, p. 3358-3378.

- Van De Weijer, Joost, et Cordelia Schmid. 2006. « Coloring local feature extraction ». In *Computer Vision–ECCV 2006*. p. 334-348. Springer.
- Varma, Manik, et Andrew Zisserman. 2005. « A statistical approach to texture classification from single images ». *International Journal of Computer Vision*, vol. 62, n° 1-2, p. 61-81.
- Whelan, Paul F, et Ovidiu Ghita. 2008. « Colour texture analysis ». In *Handbook of texture analysis*, sous la dir. de Mirmehdi, Majid, Xianghua Xie et Jasjit Suri. London: Imperial College Press.
- Yogamangalam, R., et B. Karthikeyan. 2013. « Segmentation Techniques Comparison in Image Processing ». *International journal of engineering and technology*, vol. 5, n° 1, p. 307-313.
- Yu, Hualong, Chaoxu Mu, Changyin Sun, Wankou Yang, Xibei Yang et Xin Zuo. 2015. « Support vector machine-based optimized decision threshold adjustment strategy for classifying imbalanced data ». *Knowledge-Based Systems*, vol. 76, p. 67-78.
- Zuiderveld, Karel. 1994. « Contrast limited adaptive histogram equalization ». In *Graphics gems IV*, sous la dir. de Paul, S. Heckbert. p. 474-485. Academic Press Professional, Inc.