



Détection automatique des opacités en tomosynthèse numérique du sein

Giovanni Palma

► To cite this version:

Giovanni Palma. Détection automatique des opacités en tomosynthèse numérique du sein. Traitement du signal et de l'image [eess.SP]. Télécom ParisTech, 2010. Français. NNT : . pastel-00005948

HAL Id: pastel-00005948

<https://pastel.hal.science/pastel-00005948>

Submitted on 8 Apr 2010

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

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

THESE
présentée par
Giovanni Palma
pour l'obtention du
GRADE DE DOCTEUR

Spécialité : Signal et Images

Détection automatique des opacités en
tomosynthèse numérique du sein

Soutenue le : 23 février 2010
devant le jury composé de :

Didier Dubois

Rapporteur

Lionel Moisan

Rapporteur

Isabelle Bloch

Directeur de thèse

Serge Muller

Codirecteur de thèse

Bruno Boyer

Examineur

Corinne Vachier

Examineur

Djamel Abdelkader Zighed

Examineur

à mon grand père, ma grand mère et ma mère
à Aurélie

Abstract

Digital breast tomosynthesis is a new imaging technique that may potentially overcome some limitations of standard mammography like tissue superimposition. Unfortunately, the amount of data to review by the radiologist is also increased. In this context, it makes sense to design a tool to detect suspicious radiological findings in order to help him to have an acceptable reading time and to keep a high sensitivity. In this imaging modality, several patterns may indicate the presence of cancer: calcification clusters, masses and architectural distortions. In our work, we focus on the detection of masses and architectural distortions, which is challenging because of great variability and sometimes low contrast of these signs. Masses can be seen as over-densities in the image while architectural distortions are more subtle convergence patterns without dense kernel.

In order to detect masses, we propose a sound theoretical work aiming at extending connected filters into the fuzzy sets framework. The proposed work describes a generic and comprehensive processing chain composed of image imprecision representation, filters for various common tasks and efficient algorithms allowing to actually implement them on real problems. Although this formalism is illustrated in the context of mass detection in DBT, the theoretical study we conducted allows not to be limited to that specific purpose.

Structure segmentation, which is a key step for characterization, is a difficult part of an automatic mass detection chain. A wrong segmentation can invalidate all the processing steps that come after and thus lead to a wrong decision. Beyond the limitation of any given segmentation method, a wrong segmentation can be caused by imprecision and uncertainty coming from the image. Actually, it is a quite difficult task to define precisely a single contour for a given mass in 2D/3D mammography images. For this reason, we propose to use the fuzzy contour formalism which is suitable to deal with this limitation. Thus, we conduct a study on various fuzzy segmentation methods applied to DBT data. We also discuss the link between fuzzy segmentation and the former fuzzy connected filters formalism.

As an additional contribution, we also propose a completely different kind of tool in order to detect architectural distortions. This one is based on an contrario modeling of the suspicious convergence detection problem. In addition to the method, another contribution, which mainly consists of a fast algorithm, is also proposed. This one enables us to actually process DBT datasets.

These tools are finally combined together in order to build a multi-channels Computer-Aided Detection system. Both of the channels that have been implemented have been validated on clinical data corresponding to the two kinds of radiological findings previously presented. The proposed approach is also compared to the state of the art techniques that can be found in the literature.

Résumé

La tomosynthèse numérique du sein est une nouvelle technique d'imagerie 3D qui peut potentiellement pallier certaines limitations de la mammographie standard comme la superposition de tissus. Ces améliorations se font au prix d'une plus grande quantité de données à examiner pour le radiologue. Dans ce contexte, l'élaboration d'un outil de détection automatique de signes radiologiques suspects, pour permettre au praticien de conserver un temps de lecture acceptable avec une haute sensibilité, prend tout son sens. Dans ce type d'imagerie, les cancers peuvent se traduire par plusieurs types de signes radiologiques : les amas de microcalcifications, les masses et les distorsions architecturales. Nos travaux se sont focalisés sur la détection des masses et des distorsions architecturales, tâche particulièrement difficile compte tenu de la grande variabilité morphologique de ces signes radiologiques et de leur contraste qui peut être particulièrement faible. Les masses se traduisent par une sur-densité dans l'image alors que les distorsions architecturales ne présentent quant à elles que des motifs subtils de convergence sans noyau dense.

Pour détecter les masses, nous proposons l'établissement d'un cadre théorique fort visant à étendre les filtres connexes dans le formalisme des ensembles flous. Les travaux proposés couvrent de manière complète une chaîne générique de traitement, en commençant par l'introduction d'imprécision dans l'image, puis en formalisant différents filtres possibles dédiés à différentes tâches courantes pour finir par le développement d'algorithmes de filtrage efficaces rendant l'approche dans son ensemble réaliste. Bien que le formalisme proposé soit illustré dans le contexte de la détection des masses en tomosynthèse, son étude théorique permet de ne pas s'y limiter.

La segmentation, point essentiel pour leur caractérisation, est une étape assez délicate dans une chaîne de détection automatique des masses à noyaux denses. Une mauvaise segmentation peut invalider tous les traitements qui lui sont postérieurs et amener à une mauvaise prise de décision. Au delà des limitations intrinsèques à la méthode de segmentation choisie, un tel cas peut se produire principalement à cause de l'imprécision et de l'incertitude émanant de l'image. En effet, il est souvent difficile de définir avec précision un contour unique pour une masse donnée. Pour cette raison nous avons proposé d'utiliser le formalisme des contours flous qui est un outil permettant de prendre en compte de tels facteurs limitant de détection. Nous proposons ainsi une étude de différentes méthodes de segmentation floue et illustrons leur utilisation en tomosynthèse. Nous montrons aussi le lien avec le précédent formalisme des filtres connexes flous.

De manière parallèle, les distorsions architecturales sont traitées grâce à l'adaptation du modèle a contrario pour détecter des signes de convergences qui sont hautement suspects en tomosynthèse. Une reformulation du problème est aussi introduite permettant la mise en œuvre rapide de l'approche proposée autorisant ainsi son utilisation pour traiter des volumes de tomosynthèse.

Ces différents outils sont finalement combinés pour l'élaboration d'une chaîne de détection de signes radiologiques à plusieurs canaux. Les deux canaux mis en œuvre et validés sur des données cliniques traitent respectivement les deux types de signes radiologiques précédemment évoquées. Une comparaison avec l'état de l'art des performances obtenues est présentée.

Remerciements

Cette thèse est le résultat d'un partenariat entre GE Healthcare (Buc, France) et Télécom ParisTech (Paris, France) dans le cadre d'une convention CIFRE (20061165) grâce au concours de l'Association Nationale Recherche Technique.

Je voudrai tout d'abord remercier ma directrice de thèse et mon co-directeur de thèse, à savoir Isabelle Bloch et Serge Muller pour leur encadrement d'une très grande qualité pendant ces trois années de thèse. Je souhaite tout particulièrement souligner leur confiance qui m'a permis d'avoir une grande liberté de mouvement me permettant ainsi d'explorer des pistes scientifiques toujours plus intéressantes.

Je tiens aussi à remercier Didier Dubois et Lionel Moisan, qui ont accepté d'évaluer de manière constructive mes travaux, sans oublier les autres membres du jury, à savoir Bruno Boyer, Corinne Vachier, et Djamel Abdelkader Zighed.

Mes remerciements vont aussi à Răzvan pour son implication dans ces travaux de thèse ainsi que Gero qui m'a permis de prendre un bon départ sans oublier toutes les personnes qui ont constitué un environnement de travail stimulant, aussi bien humainement que techniquement. Parmi ces dernières je voudrai citer Sylvie, Xavier, Laurence et il y a bien longtemps Fanny. De manière plus éphémère, le passage de nombreux stagiaires a participé aussi à cette ambiance de travail. Parmi ces derniers je souhaite citer Louis qui a apporté sa contribution dans les travaux présentés dans ce manuscrit. De même je souhaite remercier les personnes de l'équipe TSI à Télécom Paritech, et notamment Olivier, David, Nicolas, Jérémie, Geoffroy et tous les autres.

Un grand merci à mes amis, tout particulièrement Guillaume mais aussi Polo, Mélanie, Gence, et Julie, qui m'ont encouragé pendant cette longue période. J'aimerais aussi exprimer ma reconnaissance à ma chère et tendre Aurélie qui a été à mes côtés tout le long de cette aventure. Merci aussi à mamie et Maryline pour leur soutien. Je souhaiterais aussi avoir une pensée pour Yvette. Enfin je voudrai avoir un dernier mot pour ma grand-mère et ma mère qui m'ont soutenu bien au delà de ces trois dernières années.

Table des matières

Abstract	i
Résumé	iii
Remerciements	v
Introduction	1
1 Le sein, la mammographie, et la détection automatique de cancers	5
1.1 Le cancer du sein	5
1.1.1 La mammographie	5
1.1.2 Caractérisation des différents types de cancer	6
1.2 La tomosynthèse du sein	7
1.2.1 Motivation	8
1.2.2 Principe général	9
1.2.3 Méthodes de reconstruction	9
1.2.4 Limitations et bénéfices	10
1.3 La détection automatique de cancers en mammographie conventionnelle	11
1.3.1 Motivations	11
1.3.2 Différents composants d'un système de détection automatique	12
1.3.3 Pré-traitement des images	13
1.3.4 Détection des microcalcifications	13
1.3.5 Détection des opacités	14
1.3.6 Prise de décision	16
1.3.7 Autres approches	18
1.4 Aide à la détection pour des volumes de tomosynthèse	19
1.4.1 Motivations	19
1.4.2 Différents designs de système de détection	19
1.4.3 Microcalcifications	19
1.4.4 Densités	20
1.4.5 Approche proposée	21
1.5 Conclusion	22
2 Filtres connexes flous	25
2.1 Les filtres connexes	25
2.1.1 Opérateurs connexes binaires	26
2.1.2 Extension aux images à niveaux de gris et filtres de nivellement	27
2.1.3 Ouvertures d'attribut et filtres d'amincissement (thinning)	27
2.2 Représentation de l'imprécision dans les images	27
2.2.1 Ensembles flous	28
2.2.2 Les image floues	28
2.2.3 Images d'ombres floues	28

2.2.4	Images d'intervalles flous et images de nombres flous	29
2.2.5	Autres utilisations du flou en traitement d'image	30
2.3	Les filtres connexes flous	31
2.3.1	Connexité des ensembles flous	31
2.3.2	Régions homogènes et zones plates floues	33
2.3.3	Familles d'opérateurs connexes flous	34
2.3.4	Persistance d'un ensemble flou	36
2.3.5	Opérateurs connexes flous pour les images à niveaux de gris flous	36
2.3.6	Lien avec les opérateurs connexes classiques	39
2.3.7	Définition de filtres dans un contexte de reconnaissance	45
2.4	Utilisation pratique des filtres connexes	46
2.4.1	Outils	46
2.4.2	Filtre de détection de lésions circonscrites	48
2.4.3	Extraction à partir d'une <i>IOF</i>	49
2.4.4	Extraction à partir d'une <i>IIF</i>	50
2.5	Conclusion	52
3	Mise en œuvre des filtres connexes flous	53
3.1	Notations et définitions	54
3.2	Croissance d'un arbre représentant un ensemble flou	56
3.3	Décroissance d'un arbre associé à un ensemble flou	66
3.4	Application au filtrage d'image de quantités floues	68
3.4.1	Images floues et représentation des niveaux de gris	68
3.4.2	Croissance d'un ensemble flou	68
3.4.3	Réduction d'un ensemble flou	69
3.4.4	Filtrage et discussion	70
3.5	Conclusion	74
4	Construction d'images floues	77
4.1	Méthodes générales de construction	77
4.1.1	Construction d'images d'ombres floues nettes	77
4.1.2	Construction en utilisant un patron	78
4.2	Construction d'image et débruitage	80
4.2.1	Débruitage d'images et construction d'image floues	80
4.2.2	Construction d'image floues à partir de filtres de rang	80
4.2.3	Application du principe d'extension pour construire des images floues	81
4.3	Débruitage flous par ondelettes	82
4.3.1	Décomposition en ondelettes	82
4.3.2	Seuillage des coefficients d'ondelettes	84
4.3.3	Coefficients flous	85
4.3.4	Génération d'une image floue	86
4.3.5	Quantité de flou	87
4.3.6	Expérimentations sur images synthétiques	89
4.3.7	Expérimentation sur une image réelle	92
4.4	Conclusion	92
5	Segmentation floue	95
5.1	Intérêt des contours flous	95
5.1.1	Interprétation d'une image : imprécision contre incertitude	95
5.1.2	Formalisme des contours flous	96
5.2	Segmentation par seuillages multiples	96
5.2.1	Principe	96
5.2.2	Lien avec les images d'ombres floues	97
5.2.3	Illustration et limitations	98
5.3	Formulation du problème dans le cadre des ensembles de niveaux	98

5.3.1	Ensemble de niveaux	98
5.3.2	Méthode originale	99
5.3.3	Segmentation par ensembles de niveaux flous	99
5.3.4	Limitations	100
5.4	Approche par programmation dynamique	100
5.4.1	Travaux existants	100
5.4.2	Obtentions de contours flous	104
5.4.3	Résultats	106
5.5	Conclusions et perspectives	112
6	Utilisation des contours flous pour la caractérisation de lésions	115
6.1	Schéma global	115
6.1.1	Marqueurs	116
6.1.2	Hypothèses multiples	116
6.2	Segmentation et conditionnement de l'information	116
6.2.1	Extraction de contours flous	116
6.2.2	Conditionnement de l'information	116
6.3	Classification à l'aide d'arbres de décision flous	117
6.3.1	Description d'un arbre flou	118
6.3.2	Utilisation d'un arbre flou	118
6.3.3	Construction d'un arbre	119
6.3.4	Prise de décision	121
6.4	Résultats	121
6.4.1	Base de données	121
6.4.2	Extraction de contours	121
6.4.3	Extraction de caractéristiques	122
6.4.4	Classification	122
6.5	Conclusion	123
7	Détection de motifs de convergence	125
7.1	Motifs de convergence et cancers	125
7.2	Approches courantes	126
7.2.1	Détection de densité et analyse de la périphérie	126
7.2.2	Détection de convergences	126
7.3	Discussion sur la définition de convergence d'un point vers un autre	128
7.3.1	Définitions possibles	128
7.3.2	Comparaison des définitions	130
7.4	Modélisation a contrario pour la détection de zones de convergence	132
7.4.1	Variables aléatoires et modèle naïf	133
7.4.2	Convergences ϵ -significatives	134
7.4.3	Détection de convergence ϵ -significatives	135
7.5	Mise en œuvre rapide	136
7.5.1	Décomposition du problème	136
7.5.2	Complexité	138
7.6	Réduction de faux positifs	140
7.6.1	Agrégation des événements ϵ -significatifs	140
7.6.2	Différenciation lésion/croisement de fibres	140
7.6.3	Illustration	141
7.7	Conclusion	142

8	Détection automatique de lésions malignes en mammographie 3D	145
8.1	Description globale de la méthode	145
8.1.1	Approche	145
8.1.2	Base de données	146
8.2	Détection des lésions à noyaux denses	147
8.2.1	Détection des densités	147
8.2.2	Extraction de marqueurs à partir de la carte de détection floue	148
8.2.3	Segmentation	150
8.2.4	Caractéristiques	151
8.2.5	Séparateur à Vaste Marge	151
8.2.6	Résultats	154
8.3	Détection des lésions sans noyau dense	155
8.3.1	Détection de convergence	156
8.3.2	Réduction de faux positifs	157
8.3.3	Performance	157
8.4	Performances globales de la chaîne	158
8.4.1	Fusion des deux canaux	158
8.5	Conclusion	160
9	Conclusions et perspectives	163
A	Quelques notions sur les filtres en traitement d'images	167
A.1	Image d'ombres	167
A.2	Connexité	167
A.3	Différence filtre/opérateur	168
A.4	Rappels sur les filtres connexes	169
A.4.1	Ouverture d'attribut, amincissements	169
A.4.2	Opérateur d'extinction	169
B	Quelques notions sur les ensembles flous	171
B.1	Opérateurs d'agrégation	171
B.2	Quantités floues	171
C	Preuves sur les mises à jour d'arbres	173
C.1	Différents cas de l'opérateur $\overline{M}$	173
C.2	Preuve du théorème 3.2.1	173
C.3	Preuve du théorème 3.2.4	176
C.4	Preuve du théorème 3.2.5	177
C.5	Preuve du théorème 3.3.2	178
D	Nombre d'éléments dans B_γ et $\overline{B_\gamma}$	181
D.1	Cardinal de B_γ	181
D.2	Cardinal de $\overline{B_\gamma}$	182
	Liste des publications	185
	Bibliographie	201
	Table des figures	205
	Liste des algorithmes	207
	Liste des tableaux	209
	Index	211

Introduction

La mammographie est un domaine en forte évolution depuis plusieurs décennies. Ainsi, les systèmes analogiques ont bénéficié d'avancées technologies concernant les tubes à rayon X et les films, améliorant significativement la détection de lésions dans le sein. Plusieurs essais sur le dépistage de masse ont montré la capacité de cette modalité à réduire le taux de mortalité par cancer du sein d'environ 30% pour les femmes de plus de 50 ans, et de 18% pour les femmes entre 40 et 50 ans (Kerlikowske *et al.*, 1995; National Cancer Institute Consensus Development Panel, 1998). Plus récemment, les systèmes numériques ont été introduits permettant d'améliorer les performances cliniques (Pisano *et al.*, 2005), de réduire la durée de la procédure médicale, et ainsi d'améliorer la précision du geste et le confort de la patiente. Cependant certaines limitations de ce type d'imagerie sont toujours d'actualité, comme par exemple la superposition de tissus pouvant potentiellement cacher des lésions ou en suggérer au travers d'images construites. En effet, certaines études montrent que 76% des lésions qui sont manquées par les radiologues se trouvent dans des seins denses. Ainsi la superposition des tissus serait la principale cause de la réduction de visibilité des signes radiologiques (Holland *et al.*, 1982).

Une solution face à ce genre de limitations est de considérer l'information tridimensionnelle du sein. De nos jours, la tomosynthèse numérique du sein (Niklason *et al.*, 1997; Claus et Eberhard, 2002; Claus *et al.*, 2002) est une voie qui semble prometteuse aussi bien pour le dépistage que pour le diagnostic du cancer du sein. En ce qui concerne le dépistage, une telle technique favorise l'exploration volumique des seins denses, pouvant potentiellement limiter le nombre de femmes rappelées pour des examens complémentaires suite à la détection d'un signe radiologique provenant de la superposition de tissus. Dans le cas du diagnostic, le bénéfice attendu se situe surtout dans l'aide à la caractérisation des signes radiologiques permettant ainsi une meilleure différenciation entre signes malins et bénins.

L'apport de l'information 3D rendue disponible par la tomosynthèse se fait au prix d'une plus grande quantité de données à examiner pour le radiologue. Dans ce contexte, l'élaboration d'un outil de détection automatique de signes radiologiques suspects, pour permettre au praticien de conserver un temps de lecture acceptable avec une haute sensibilité, prend tout son sens. Dans ce type d'imagerie, les cancers peuvent se traduire par plusieurs types de signes radiologiques : les amas de microcalcifications, les masses et les distorsions architecturales. Nos travaux se sont focalisés sur la détection des masses et des distorsions architecturales, tâche particulièrement difficile compte tenu de la grande variabilité morphologique de ces signes radiologiques et de leur contraste qui peut être particulièrement faible. Les opacités se traduisent par une sur-densité dans l'image alors que les distorsions architecturales ne présentent quant à elles que des motifs subtils de convergence sans noyau dense.

Le problème de détection pour ce type d'imagerie est en fait assez peu abordé dans la littérature. Cela s'explique principalement par le caractère récent des données concernées. En effet, la tomosynthèse du sein est en cours d'introduction et la quantité des données disponibles pour le développement de nouvelles applications est extrêmement faible en comparaison d'autres types d'imagerie comme la mammographie standard. Parmi les travaux existant on citera ceux de Chan *et al.* (2008a) qui présentent les résultats les plus aboutis pour la détection d'opacités. Ils reposent sur l'utilisation de l'analyse de flot de gradient dans le volume pour la détection de zones suspectes couplée à une étape de classification pour réduire le nombre faux positifs. L'élaboration d'un outil de détection reposant sur la théorie de l'information a aussi été proposée par Singh (2008). Ce dernier type d'approche semble cependant moins performant que le premier bien que n'ayant pas été évalués sur la même base de données. De manière générale, ces travaux, bien que prometteurs, n'atteignent pas encore les performances des systèmes de détection utilisés en mammographie standard. Ainsi l'étude d'outils de détection reste un problème d'actualité.

Pour accomplir les tâches de détection des opacités et des distorsions architecturales, nous avons développé un ensemble d'outils originaux. Certains d'entre eux, destinés à la détection des opacités, reposent sur la prise en compte de l'imprécision et l'incertitude contenues dans les images pour effectuer des tâches de détection et de caractérisation. D'autres proposent une modélisation a contrario couplée à une mesure de convergence pour détecter les motifs stellaires suspects.

Dans le chapitre 1 nous exposerons plus en détail la problématique de la détection de signes radiologiques en mammographie et en tomosynthèse. Nous rappellerons les notions fondamentales de la formation des images à traiter, puis nous étudierons les techniques existantes dans la littérature pour accomplir la détection des différents signes radiologiques suspects aussi bien en mammographie conventionnelle qu'en tomosynthèse. L'étude de ces deux types d'imagerie est proposée pour tirer partie du lien qui existe entre les caractéristiques des signes suspects pouvant provenir des deux modalités. Nous présenterons aussi de manière générale notre approche pour résoudre ce même problème de détection. Pour chaque signe radiologique considéré, à savoir les opacités et les distorsions architecturales, nous proposons un canal de détection reposant sur une première étape de filtrage permettant de localiser de manière grossière les zones suspectes suivie d'une étape de caractérisation. Dans les travaux présentés dans ce document, nous introduirons et étudierons, entre autres, des filtres pour effectuer la première étape de détection.

Dans le chapitre 2, nous introduirons un formalisme permettant d'exprimer un nouveau type de filtres. Ces derniers sont une extension des filtres connexes dans le formalisme des ensembles flous. Les filtres proposés, filtres connexes flous, utilisent des versions assouplies de la notion de composante connexe ainsi qu'une représentation de l'imprécision des niveaux de gris par l'intermédiaire d'images floues. En introduisant ce dernier concept, nous proposerons différentes stratégies pour l'extraction de telles composantes, laissant la possibilité d'exprimer des extensions pour les filtres de nivellement, d'amincissement ou encore de traitement de zones plates. Nous exprimerons aussi dans ce formalisme de nouveaux filtres destinés à la détection de structures dans des images et illustrons leur utilisation comme étape de marquage pour des opacités en tomosynthèse du sein.

La représentation des données sous forme d'images floues augmente considérablement la quantité d'information à traiter. Dans le chapitre 3, nous proposerons des techniques permettant de rendre réalisable en pratique l'utilisation de tels filtres. Nous utiliserons le formalisme de représentation arborescente couramment utilisé dans le domaine des filtres connexes. Nous proposerons des algorithmes de mise à jour d'arbres permettant la prise en compte de la redondance entre les différentes composantes connexes floues extraites à partir des images.

Dans le chapitre 4, nous étudierons différentes stratégies pour introduire de manière utile de l'imprécision dans les images pour obtenir des images floues. Nous étudierons l'impact du bruit sur la construction de ces images. Nous proposerons des techniques pour convertir ce type de défaut, qui est d'ordre statistique et par conséquent non représentable par le formalisme d'images floues, en imprécision de débruitage. Nous étudierons de manière approfondie une méthode de construction reposant sur une décomposition en ondelettes ainsi que son impact sur la détectabilité des structures contenues dans une image.

Après ces trois chapitres sur les filtres connexes flous nous nous intéresserons à l'utilisation d'outils de segmentation par contours flous pour pouvoir modéliser l'incertitude et l'imprécision provenant des images de tomosynthèse. Nous détaillerons trois méthodes permettant d'accomplir cette tâche en mettant en avant leurs qualités et défauts. Nous validerons aussi sur des données réelles de tomosynthèse certaines des méthodes de segmentation non floues qui sont à l'origine des approches proposées.

Dans le chapitre 6 nous proposerons un exemple pratique d'utilisation du formalisme de contours flous précédemment évoqué. Dans cette étude, nous nous intéresserons à la différenciation entre lésion spiculée et lésion circonscrite en utilisant un schéma à deux hypothèses. La validation sur des cas réels permettra de juger de l'apport du formalisme employé.

Les différents outils introduits dans les chapitres précédemment cités sont applicables à la détection de lésions ayant un noyau dense. Pour pouvoir détecter des lésions qui ne présentent pas de noyau dense, c'est-à-dire les distorsions architecturales, un outil spécifique et radicalement différent de l'approche jusque là détaillée sera proposé au chapitre 7. Ce dernier propose la définition d'un critère de convergence dans une modélisation a contrario pour détecter les formes stellaires suspectes dans les coupes de tomosynthèse du sein.

Enfin, nous proposerons au chapitre 8 un exemple pratique d'une chaîne de détection automatique de

signes radiologiques suspects dans des volumes de tomosynthèse du sein utilisant les différents outils introduits tout au long de ce manuscrit. Une validation ainsi qu'une comparaison avec les méthodes proposées dans la littérature seront présentées.

Chapitre 1

Le sein, la mammographie, et la détection automatique de cancers

La cancer du sein est la première cause de mortalité chez la femme toutes causes confondues entre 35 et 55 ans. Avec 42000 nouveaux cas dépistés par an en France, c'est le cancer chez la femme le plus fréquent (Remontet *et al.*, 2003). Dans le but de réduire la mortalité par cancer du sein, une détection précoce est nécessaire (Duffy *et al.*, 2003) motivant ainsi des campagnes de dépistage chez les femmes à partir d'un certain âge variant entre 40 et 50 ans selon les pays. Le type d'imagerie actuellement utilisée pour cette tâche est la mammographie. Dans ce contexte, la quantité de clichés que le radiologue doit interpréter en un temps limité est importante. De plus, la complexité de la tâche de détection est particulièrement élevée en mammographie. Ces raisons motivent l'utilisation d'outils d'aide à la détection par ordinateur (CAD) afin d'améliorer les performances du lecteur, en minimisant en particulier le nombre de signes radiologiques associés à des lésions malignes qui pourraient ne pas être détectées.

Depuis peu, une nouvelle technique d'imagerie, la tomosynthèse numérique du sein, a fait son apparition. Elle peut potentiellement pallier les défauts de la mammographie standard. Dans ce type d'examen, une reconstruction 3D du sein est effectuée permettant potentiellement de mieux discerner les signes radiologiques caractérisant un cancer. Avec une quantité de données grandement augmentée par rapport à un examen standard, le besoin d'un outil automatique d'aide à la détection reste d'actualité, tant pour préserver un taux de détection élevé que pour limiter le temps de lecture.

Dans ce chapitre, nous introduirons brièvement la mammographie conventionnelle dans le cadre de la détection du cancer du sein, puis nous continuerons avec les principes généraux de la tomosynthèse. Nous ferons ensuite un tour d'horizon sur les travaux traitant de la détection automatique dans des mammographies conventionnelles, pour finir par l'état de l'art du CAD en tomosynthèse.

1.1 Le cancer du sein

Dans le but de comprendre la problématique du dépistage du cancer du sein par mammographie, nous allons brièvement rappeler comment les images sont formées pour cette modalité ainsi que l'apparence que prennent les lésions dans ces images.

1.1.1 La mammographie

La mammographie est la modalité actuellement utilisée pour dépister le cancer du sein. Une mammographie est une radiographie du sein qui est un objet tridimensionnel. L'intérêt de ce type d'acquisition est le phénomène d'accumulation induit par la physique des rayons X. En effet, on peut observer le contenu de l'organe radiographié en le projetant sur un plan 2D grâce aux rayons X. De manière simplifiée, on peut modéliser la formation d'une telle image en utilisant l'équation de Beer-Lambert :

$$I = I_0 e^{-\int_{v \in L} \mu(v) dv} \quad (1.1)$$

avec L le chemin entre la source de rayonnement et la cellule du détecteur considérée, $\mu(v)$ les coefficients d'atténuation des matériaux traversés et I_0 l'intensité initiale émise par la source. Cette équation ne modélise pas fidèlement ce qui se passe réellement lors de l'acquisition dans la mesure où elle fait l'hypothèse d'un rayonnement mono-énergétique en omettant des phénomènes tels que le rayonnement diffusé ou le durcissement de faisceau. Néanmoins, elle permet de comprendre ce que l'on observe dans les images, c'est-à-dire les différences de propriétés d'atténuation des objets composant le sein comme la glande, la graisse, les fibres ou encore des lésions.

En mammographie, comme on peut le voir sur la figure 1.1, le sein est comprimé entre le support patient, sous lequel se trouve le détecteur, et la pelote de compression. La compression est notamment appliquée pour limiter la dose délivrée à la patiente, pour réduire la quantité de rayonnement diffusé, pour étaler les tissus et pour immobiliser le sein. Un autre élément important du système de mammographie réside dans sa source de rayonnement, ou tube à rayons X, dont le foyer est positionné à la verticale du bord du détecteur le plus proche de la patiente afin d'assurer une couverture complète du sein par le faisceau de rayons X qui en est issu.

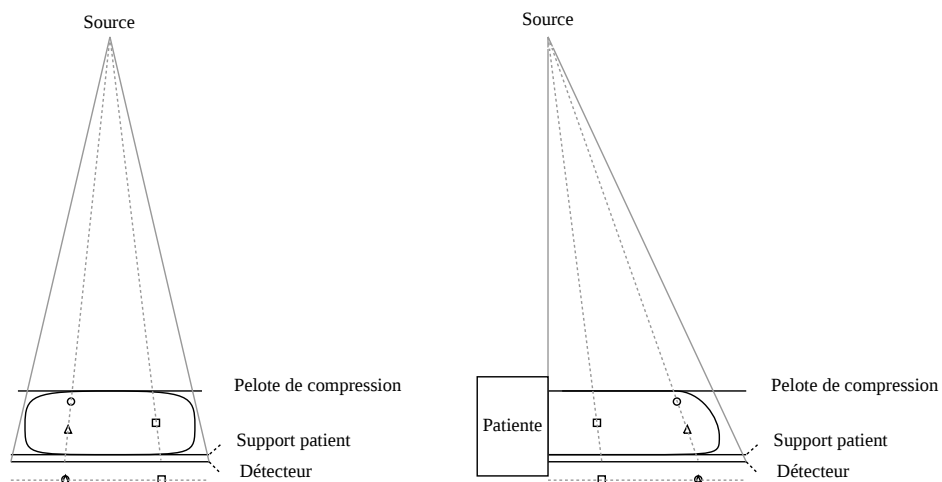


FIG. 1.1 – Géométrie d'acquisition d'une mammographie (vue de face et vue de profil).

Un examen typique de dépistage en mammographie comprend généralement l'acquisition d'images 2D sous deux incidences par sein. Ces mammographies 2D sont réalisées selon différents angles d'incidence pour minimiser les défauts d'identification de structures dus à la superposition de tissus qui pourrait masquer des lésions ou au contraire en faire percevoir de fictives.

1.1.2 Caractérisation des différents types de cancer

Il existe plusieurs signes radiologiques associés à des cancers. Le premier correspond à de petits amas de calcium aussi connus sous le nom de microcalcifications. Elles ne sont pas forcément signe de malignité. Ainsi des grosses calcifications rondes à bord lisse sont bien souvent bénignes. Malheureusement, les microcalcifications malignes sont généralement moins visibles (ACR, 2003). Un exemple d'amas de microcalcifications est présenté à la figure 1.5(a).

Une seconde famille de signes radiologiques se traduit par des sur-densités ou opacités dans les images. Ces dernières présentent une variabilité importante du point de vue de leur forme, de leur taille et de leur contour. Cette dernière caractéristique aide le radiologue dans son diagnostic différentiel. On distingue généralement quatre classes de formes pour les opacités : les opacités rondes, ovales, lobulaires ou irrégulières. Ces différentes formes sont illustrées de manière schématique à la figure 1.2.

Comme indiqué précédemment, les opacités varient aussi au niveau de la définition de leurs contours. On dénombre cinq grandes classes de contours (ACR, 2003). La première correspond aux contours circonscrits qui sont des contours bien définis où la frontière entre lésion et fond est franche et qui correspondent en général à des lésions bénignes (c.f. figure 1.3(a)). Une seconde classe de contours est distinguée lorsqu'une

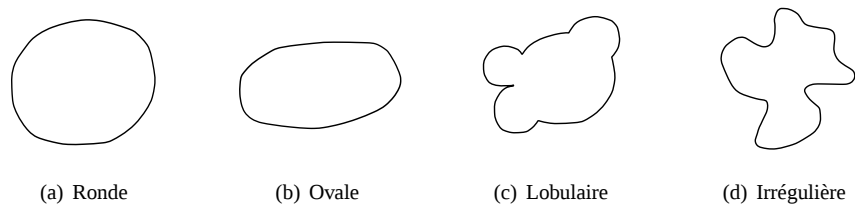


FIG. 1.2 – Différentes formes pour les opacités.

partie du contour est masquée par la superposition ou l'adjacence de tissus normaux laissant penser à une lésion circonscrite dont une partie du contour est cachée (c.f. figure 1.3(b)). Une troisième possibilité pour le contour est d'être micro-lobulé, c'est-à-dire qu'il comporte de petites ondulations (c.f. figure 1.3(c)). Le quatrième type de contours regroupe les contours mal définis (c.f. figure 1.3(d)) qui peuvent laisser penser à la présence d'infiltrations. Enfin une dernière classe regroupe les opacités dont les contours sont spiculés, c'est-à-dire comportant des structures filiformes qui rayonnent en s'éloignant du centre de l'opacité, et qui sont hautement suggestives de malignité.

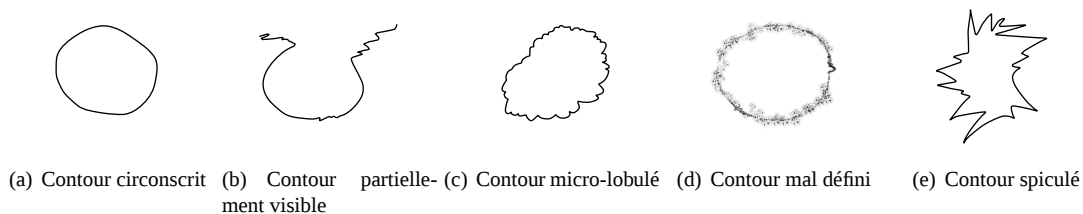


FIG. 1.3 – Différents types de contours pour les opacités.

La troisième forme de signes radiologiques traduisant la présence de cancer correspond aux distorsions architecturales. Ces dernières ne comportent pas de densité centrale comme c'est le cas pour les opacités précédemment décrites. Elles se traduisent par des structures linéaires ou des spicules qui convergent vers une même zone focale. Un exemple schématique de ce genre de signes, qui sont hautement suspects, est présenté à la figure 1.4.

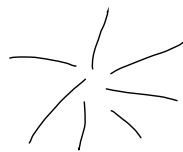


FIG. 1.4 – Forme schématique d'une distorsion architecturale.

Des exemples d'amas de microcalcifications, d'opacités et de distorsions architecturales en mammographie sont présentés à la figure 1.5.

1.2 La tomosynthèse du sein

La tomosynthèse du sein est une nouvelle méthode d'imagerie qui peut potentiellement pallier les différentes limitations de la mammographie conventionnelle notamment en produisant une représentation 3D du sein. Après avoir récapitulé les limitations de l'imagerie 2D, nous rappellerons le principe général de

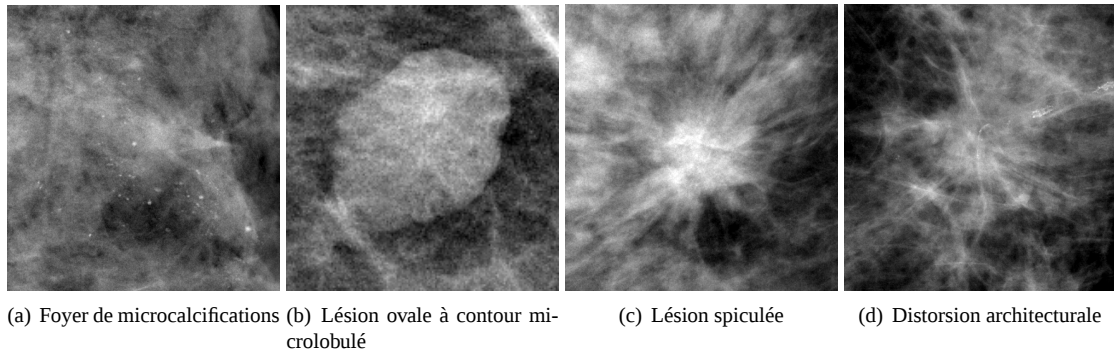


FIG. 1.5 – Différents exemples de signes radiologiques suspects en mammographie standard.

la tomosynthèse ainsi que les méthodes de reconstruction couramment décrites dans la littérature. Nous finirons enfin par un descriptif des limitations de ce type d'imagerie.

1.2.1 Motivation

Comme on l'a vu précédemment, la mammographie 2D est une imagerie de projection. En effet on projette une structure 3D (le sein) sur une surface 2D (le détecteur), résultant en une perte potentielle d'information. Plus concrètement, il existe deux problèmes majeurs qui peuvent être rencontrés. Le premier se traduit par la non détection d'une lésion qui est masquée par les tissus présents dans la trajectoire du faisceau passant par cette dernière. Ce cas est très problématique dans la mesure où ne pas détecter un cancer implique de ne pas le traiter de manière précoce. Le second problème susceptible d'être rencontré est la détection d'un signe radiologique construit par accumulation de différentes structures normales et indépendantes, alignées sur une même trajectoire de rayons X et formant, une fois projetées, un motif laissant penser à la présence d'une lésion. Bien que ce type de fausse alarme ait des conséquences moins graves que la non détection, le fait de devoir subir d'autres examens voire même une biopsie pour infirmer la présence d'un cancer est une procédure stressante pour la patiente qu'il est souhaitable d'éviter.

Ces problèmes sont donc essentiellement causés par la projection de l'information lors de la formation d'une image de mammographie. L'idée la plus simple pour contourner ce problème est de travailler sur une représentation tridimensionnelle du sein. Des méthodes de formation d'image 3D sont couramment utilisées en médecine. Ainsi, le scanner par tomographie est utilisé pour imager divers organes du corps. Néanmoins son utilisation pour le sein peut poser problème essentiellement pour des raisons de dose délivrée à la patiente (Chang *et al.*, 1979; Gisvold *et al.*, 1979; Chang *et al.*, 1982; John et Ewen, 1989). En effet pour être capable d'imager le tronc de la patiente, il faut des rayons X capables de le traverser de part en part ce qui impose une dose beaucoup plus importante que pour une mammographie standard (Muller *et al.*, 1983; Sibala *et al.*, 1981). Des scanners dédiés où seul le sein est imagé peuvent aussi être considérés (Boone *et al.*, 2001; Kwan *et al.*, 2004; Yang *et al.*, 2007). Dans ce type d'imageurs, le tube fait le tour complet du sein et non plus du tronc permettant de pallier les précédentes limitations. Cependant, la région axillaire du sein n'est généralement pas radiographiée à cause de contraintes physiques. D'autres types d'acquisition s'appuyant sur d'autres modalités existent et sont quant à eux réellement employés dans des routines de diagnostic. Ainsi on peut mentionner l'imagerie par résonance magnétique IRM du sein. Cependant pour des raisons de coûts et de disponibilité, ce type d'imagerie est généralement réservé aux patientes à risque dans le cadre du dépistage aux Etats Unis (Saslow *et al.*, 2007).

D'autres causes de non détection sont possibles comme par exemple la non atténuation d'une opacité par les rayons X. La cause de cette limitation est principalement due à la similarité entre les propriétés radiométriques de ces lésions et de celles des structures saines plutôt qu'au processus de formation d'image. Dans le cas de lésions faiblement discernables, des méthodes d'imagerie nécessitant l'injection d'un produit de contraste existent comme l'angiomammographie (Jeunehomme, 2005; Puong, 2008) ou l'IRM (Rankin, 2000) qui permettent la mise en évidence de lésions hypervascularisées.

Dans ce paysage d'outils d'imagerie, la tomosynthèse du sein (Wu *et al.*, 2003b; Ren *et al.*, 2005)

semble pouvoir jouer un rôle important dans le dépistage et le diagnostic des cancers du sein, notamment en se positionnant comme une alternative à la tomographie axillaire permettant l'accès à l'information volumique pour un budget dose maîtrisé.

1.2.2 Principe général

Le principe de la tomosynthèse est de reconstituer l'information tridimensionnelle de l'objet imagé à partir de plusieurs radiographies à faible dose acquises sous différents points de vue (Dobbins III et Godfrey, 2003). En effet comme on peut le voir sur la figure 1.6, en prenant différents clichés à des positions diverses d'un même objet, la position relative des structures qui le composent, une fois projetées sur le détecteur, change. Pour accéder à une information 3D de l'objet imagé, ces différentes projections sont utilisées comme entrée d'un algorithme de reconstruction comme on le verra dans la section suivante. L'étape de reconstruction fournit une série de coupes parallèles au détecteur à différentes hauteurs par rapport à ce dernier qui représentent l'atténuation radiologique des structures qui composent le sein.

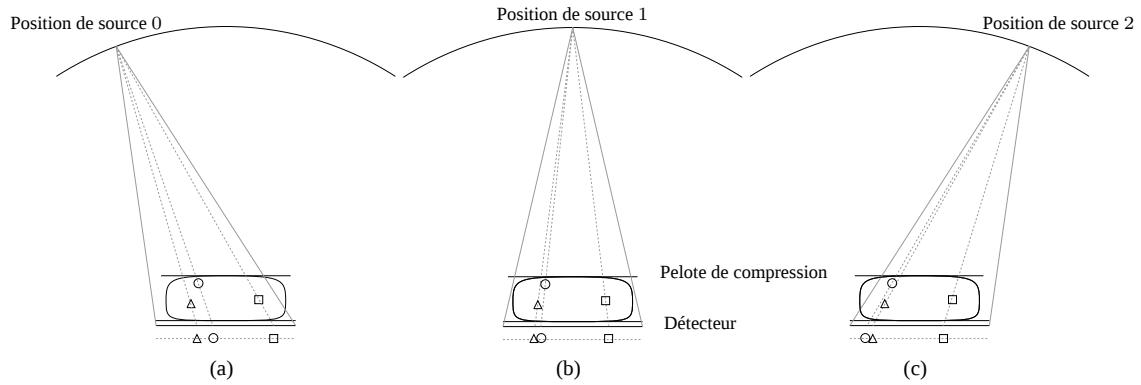


FIG. 1.6 – Principe de la tomosynthèse.

De manière générale, l'apparence des systèmes de tomosynthèse du sein est assez similaire à celle des appareils de mammographie. Le sein de la patiente est positionné en contact avec le support patient et comprimé à l'aide d'une pelote de compression. La position du tube évolue selon une trajectoire généralement circulaire dont le centre de rotation se situe au niveau du détecteur. Comme évoqué précédemment, une série de projections 2D à faible dose est acquise pour différentes positions du tube à rayon X autour du sein. Le nombre de projections varie généralement entre 9 et 45 sur une ouverture angulaire variant de 30 à 90 degrés. La dose allouée à chaque projection est telle que la dose totale délivrée à la patiente est comprise entre une et deux fois la dose d'une mammographie standard. Pour cette raison, la qualité des images servant de données d'entrée à la reconstruction est beaucoup plus médiocre que celle d'une mammographie conventionnelle, notamment en termes de rapport contraste à bruit. Cependant l'étape de reconstruction permet généralement d'obtenir des coupes reconstruites avec de meilleures propriétés.

1.2.3 Méthodes de reconstruction

En reconstruction tomographique, on cherche, à partir de la connaissance d'un ensemble de projections d'un paramètre caractéristique l'objet d'intérêt, à remonter à une cartographie de ce paramètre en tout point de l'objet. Une projection P_θ s'exprime sous une forme intégrale qui, pour ce qui nous concerne correspond à l'intégrale des coefficients d'atténuation linéaires de l'objet selon le trajet des rayons X. Les images I délivrées par le système d'acquisition de tomosynthèse sont donc pré-traitées de telle sorte que l'on ait accès à l'information de projection P_θ :

$$P_\theta = -\ln\left(\frac{I}{I_0}\right) = \int_{v \in L} \mu(v) dv \quad (1.2)$$

Ces images sont connues sous le nom d'images d'épaisseurs radiologiques et elles correspondent à l'intégrale contenue dans l'équation 1.1.

On peut distinguer deux approches principales pour la reconstruction tomographique : les méthodes fondées sur la rétroprojection filtrée des projections et les méthodes de reconstruction itératives.

Rétroprojection filtrée

La reconstruction tomographique d'un volume de coefficients d'atténuation à partir de différentes projections peut être effectuée à l'aide du théorème de Radon (1917). De manière théorique, une reconstruction dans un plan peut se faire par rétroprojection des projections acquises de ce plan préalablement filtrées par une filtre rampe. Cette méthode est connue sous le nom de rétroprojection filtrée. En pratique, l'étape de filtrage est généralement effectuée par le résultat de la convolution entre un filtre passe-bas et le filtre rampe. Cela permet de limiter l'impact du bruit au détriment de la résolution spatiale.

Cette méthode de reconstruction repose sur plusieurs hypothèses qui ne sont pas complètement vérifiées dans le cadre de la tomosynthèse. On peut citer notamment la plus importante qui est de pouvoir faire le tour de la patiente. L'utilisation d'une telle méthode de reconstruction est étudiée en tomosynthèse du sein dans les travaux de Wu *et al.* (2004b).

Méthodes itératives

La seconde classe d'algorithmes de reconstruction regroupe les méthodes itératives. Ces méthodes consistent généralement en une modélisation du problème par une énergie qu'il faut minimiser. C'est en fait cette étape de minimisation qui est réalisée de manière itérative.

Plus concrètement, on définit un opérateur de projection d'un volume sur le plan du détecteur qui est censé être fidèle à l'équation 1.2. Grâce à cet opérateur, on compare, lors d'une itération donnée, les projections obtenues à partir du volume courant et les projections que le vrai volume a engendrées. En rétroprojetant la différence entre ces images, on est capable de se rapprocher d'un volume qui minimise l'énergie considérée. Les différentes méthodes varient dans la façon dont cette étape de mise à jour est incorporée dans les différentes itérations.

Dans la littérature, on trouve quatre grandes méthodes de reconstruction. La première est connue sous le nom de technique de reconstruction algébrique ou *ART* (Gordon *et al.*, 1970; Gordon et Herman, 1974). Elle consiste à mettre à jour le volume en utilisant uniquement l'information de projection obtenue pour un couple (source, pixel) donné. La seconde méthode de reconstruction couramment usitée est connue sous le nom de technique de reconstruction algébrique simultanée ou *SART* (Andersen et Kak, 1984). Dans cette méthode, chaque mise à jour fait appel à toute l'information d'une image de projection. La convergence de ce genre d'approche est discutée dans les travaux de Jiang et Wang (2003). Une troisième variante qui met à jour le volume à chaque étape en utilisant toutes les projections est connue sous le nom *SIRT* qui signifie technique de reconstruction itérative simultanée (Gilbert, 1972). Enfin, une dernière méthode de reconstruction utilisée en tomosynthèse du sein est une méthode fondée sur le maximum de vraisemblance ou *MLEM* (Wu *et al.*, 2003a). Contrairement aux précédentes méthodes, cette modélisation du problème de reconstruction est abordée de manière probabiliste résultant en une procédure de mise à jour du volume qui se fait de manière multiplicative.

Bien que les méthodes de reconstruction itératives soient plus coûteuses en temps de calcul que les méthodes de type rétroprojection filtrées, elles sont particulièrement intéressantes pour aborder les problèmes de reconstruction en présence de données manquantes, comme c'est le cas pour une acquisition de type tomosynthèse. Leur mise en œuvre devient envisageable dans des utilisations pratique notamment grâce au niveau élevé de performances des calculateurs modernes. De plus, ces approches peuvent facilement être parallélisées (Wu *et al.*, 2004c).

1.2.4 Limitations et bénéfices

L'avantage principal de la tomosynthèse est de limiter la superposition de tissus. Ainsi, on peut espérer rendre les signes radiologiques associés à des lésions suspectes plus facilement discernables (Gennaro *et al.*, 2008). La figure 1.7 présente un exemple où la tomosynthèse permet une meilleure caractérisation sur une zone suspecte de l'image.

De plus, la tomosynthèse du sein étant effectuée à plus faible énergie que les examens de tomographie standard, le contraste entre les différentes structures du sein est meilleur. En effet, les courbes d'atténuation linéaires entre tissus adipeux et graisseux par exemple sont plus distinctes à faible énergie.

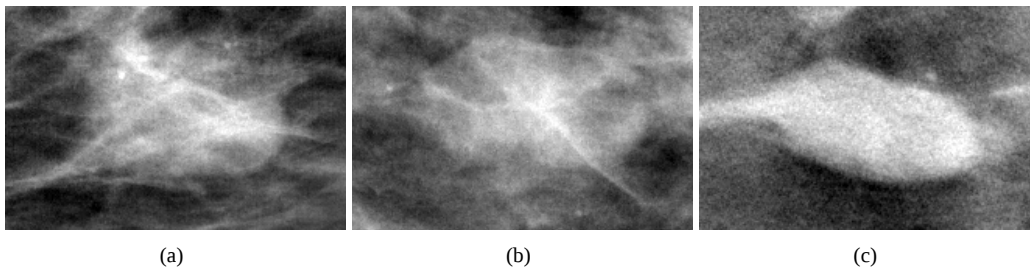


FIG. 1.7 – Exemple de réduction de la superposition de structures rendant une lésion plus visible en tomosynthèse. Dans les vues de mammographie standard CC (a), et MLO (b), le contour de la lésion n'est pas bien défini. Dans la coupe du volume reconstruit où la lésion apparaît le plus clairement (c), le contour de la lésion devient beaucoup plus facile à décrire.

Néanmoins, comme on l'a vu précédemment, la création du volume de tomosynthèse peut être obtenue de différentes manières. L'algorithme de reconstruction joue donc un rôle important sur l'apparence des données reconstruites. Ainsi le type de filtre utilisé dans les approches de type rétroprojection filtrée ou encore le nombre d'itérations ainsi que la stratégie de mise à jour dans les approches itératives influent sur le bruit ou la quantité de flou contenus dans les coupes reconstruites (Zhang *et al.*, 2006).

Un dernier point important concerne les artefacts de reconstruction. En tomosynthèse, on ne peut pas faire le tour de la patiente pour acquérir des données, ainsi d'un point de vue théorique, il manque de l'information pour obtenir une reconstruction parfaite. D'un point de vue pratique cela se traduit par une fonction d'étalement ponctuelle en forme d'étoile et allongée dans la direction de la hauteur. Des études plus approfondies de cette dernière sont proposées par Avinash *et al.* (2006) et par Blessing *et al.* (2006). Ainsi dans les volumes reconstruits, les objets apparaissent déformés dans la même direction comme on peut le voir pour la macrocalcification présentée à la figure 1.8.

Une autre limitation de la tomosynthèse est le pas d'échantillonnage important selon la hauteur des coupes reconstruites (environ 1mm). Cela a un impact sur la reconstruction des microcalcifications qui ont une taille typiquement comprise entre $100\mu\text{m}$ et 1mm et qui peuvent ainsi se situer entre deux coupes reconstruites. Dans ce cas, la visibilité de ces dernières peut être affectée. Des travaux ont montré qu'une ouverture angulaire moins importante pouvait limiter ce problème (Lau *et al.*, 2008). L'idée est qu'une ouverture angulaire diminuée augmente l'étalement des microcalcifications selon la hauteur des coupes, augmentant ainsi une éventuelle intersection avec les coupes reconstruites.

1.3 La détection automatique de cancers en mammographie conventionnelle

Le passage à la mammographie numérique a ouvert la porte à de nouvelles possibilités en termes d'outils visant à extraire de manière optimale l'information dans l'image. Ainsi, des traitements visant à améliorer l'image, ou encore des outils d'aide à la détection sont couramment utilisés par les radiologues. Dans cette section nous allons nous intéresser à cette deuxième classe d'outils. Après un bref rappel sur les motivations de ces derniers, nous détaillerons les différentes manières dont ils peuvent être construits. Nous insisterons aussi sur les points clés qui sont récurrents dans les approches proposées dans la littérature, à savoir, le conditionnement des images, la détection des calcifications et des opacités ainsi que la prise de décision.

1.3.1 Motivations

La détection de signes radiologiques dans une campagne de dépistage est une tâche délicate de par la quantité de données à observer et la subtilité des signes radiologiques. Ainsi il peut être intéressant de

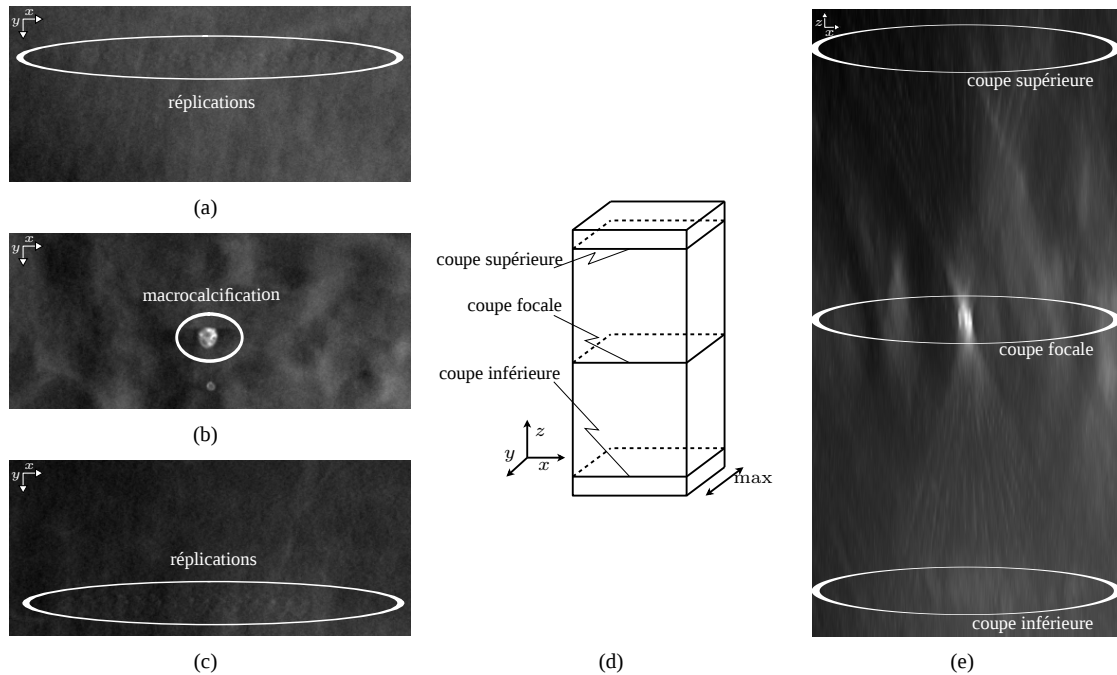


FIG. 1.8 – Illustration de la distorsion d’objets dans un volume reconstruit. (b) Coupe d’un volume contenant une macrocalcification. (a) et (c) Coupes situées respectivement au-dessus et en-dessous de la macrocalcification et présentant des répliques de cette dernière. (d) Visualisation schématique du volume d’intérêt contenant la macrocalcification. (e) Projection ($\max$) selon l’axe des y de ce volume d’intérêt : on voit que la macrocalcification est allongée en z et présente un motif en étoile.

fournir au radiologue un outil qui pourrait avoir une sensibilité de détection élevée et constante dans le temps, lui permettant ainsi de détecter plus de cancers.

Des études récentes confortent l’idée de l’utilité d’un tel outil. En effet, Karssemeijer *et al.* (2003) présentent par exemple une étude comparative entre lecture seule aidée ou non d’un système de détection et la double lecture classique. Bien qu’il en ressorte qu’une double lecture de radiologues expérimentés reste la référence, un système de CAD permet d’améliorer les performances d’un lecteur seul. Ce dernier point était déjà exposé dans des travaux plus anciens (Freer et Ulissey, 2001).

Une étude plus récente (Gilbert *et al.*, 2006) a montré sur un plus grand nombre de cas, qu’une lecture simple associée à l’utilisation d’un CAD augmentait la détection de cancers de manière statistiquement significative par rapport à une double lecture. Néanmoins le taux de rappel se voit, dans cette étude, lui aussi augmenté.

1.3.2 Différents composants d’un système de détection automatique

De manière globale un système de détection automatique de cancers en mammographie se compose de deux branches : une dédiée à la détection des microcalcifications et l’autre à la détection des opacités. Chacun de ces modules peut se décomposer comme une étape de marquage suivie d’une prise de décision. Le marquage peut selon les cas être composé d’une détection rapide suivie d’une segmentation. La prise de décision se compose quant à elle d’une étape d’extraction de caractéristiques suivie d’une étape de classification. Certains de ces éléments peuvent apparaître de manière plus ou moins implicite. La figure 1.9 illustre la décomposition de haut niveau des processus de détection automatique. Dans certains cas, un pré-traitement des données permettant de mettre en évidence les signes recherchés peut être utilisé. Néanmoins, cette étape montre généralement rapidement ses limitations dans la mesure où, pour vraiment mettre en évidence un motif, il faut être capable de le détecter, or c’est le but de cette étape de pré-traitement.

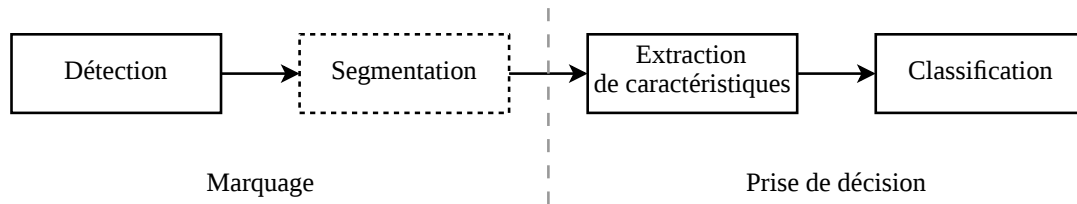


FIG. 1.9 – Schéma général d'une chaîne de détection de microcalcifications ou d'opacités en mammographie 2D.

1.3.3 Pré-traitement des images

Une première étape souvent utilisée est donc la préparation des images avant la détection. En effet les structures que l'on recherche n'étant pas toujours facilement discernables, une étape de pré-traitement destinée à les mettre en évidence peut faciliter leur détection.

Une approche couramment utilisée en traitement d'images consiste à travailler sur l'histogramme dans le but de définir une fonction de transfert sur les niveaux de gris permettant de mettre en valeur les détails présents dans l'image. Cette voie a aussi été explorée en mammographie numérique (Karssemeijer, 1993; Wilson *et al.*, 1997; Gavrielides *et al.*, 2002; Wilson *et al.*, 1997). Néanmoins, le problème d'une telle approche réside en sa limitation au niveau du traitement des textures, ce qui est gênant puisque ces dernières portent une information importante pour la détection de signes radiologiques.

D'autres approches plus locales existent et se comportent mieux vis-à-vis des textures. Ainsi, Braccialarghe et Kaufmann (1996) proposent une transformation des niveaux de gris en fonction d'éléments locaux comme les contours ou les statistiques locales.

D'autres méthodes ont aussi été proposées comme l'utilisation de filtres d'amélioration de la netteté (Chan *et al.*, 1987), où encore la suppression du fond de l'image (Zhou et Gordon, 1989). Cette dernière approche consiste en la soustraction d'une version filtrée passe bas de l'image originale à cette dernière. Certains auteurs proposent quant à eux d'utiliser l'information d'orientation (Chang et Laine, 1997) tandis que d'autres utilisent une décomposition en ondelettes hexagonales (Pfisterer et Aghdasi, 1999).

D'autres formalismes sont aussi utilisés pour accomplir la tâche d'amélioration du contraste. Ainsi Hassanien et Ali (2004) proposent d'utiliser les ensembles flous et l'analyse de l'histogramme de l'image pour atteindre ce but.

L'idée d'un pré-traitement, bien qu'intuitive, peut poser quelques problèmes. En effet, dans l'idéal on voudrait mettre en évidence seulement les zones potentiellement suspectes dans le but de faciliter leur détection ultérieurement. Or pour pouvoir accomplir cette tâche, il faudrait savoir quelles zones de l'image on doit améliorer, c'est-à-dire connaître les zones suspectes, ce qui est difficile puisque l'on cherche à améliorer l'image pour justement détecter ces structures. D'autre part, l'étape de pré-traitement peut aussi modifier certaines structures et les rendre faussement suspectes, ce qui peut être problématique pour l'étape de détection. De manière similaire, un pré-traitement peut modifier substantiellement les propriétés de l'image, rendant la modélisation de l'étape de détection délicate.

1.3.4 Détection des microcalcifications

La première classe de signes radiologiques qui sont recherchés sont les foyers de microcalcifications. La littérature comporte un nombre important de méthodes de traitement à bases d'ondelettes (Yoshida *et al.*, 1996; Strickland et Hahn, 1997; Chen et Lee, 1997; Bernard *et al.*, 2006). En effet les amas de calcium étant constitués d'objets de petite taille assez contrastés, ce genre d'approche semble avoir un réel potentiel. Par exemple, Strickland et Hahn (1997) proposent de modéliser le problème comme la détection d'objets gaussiens dans un bruit markovien en utilisant une décomposition par banc de filtres. Une autre approche (Bernard *et al.*, 2006) consiste à effectuer un filtrage avec un chapeau mexicain dans le but de mesurer le contraste des structures de dimension compatible avec le noyau du filtre. En ajoutant une contrainte dérivée du bruit de Poisson que l'on peut généralement estimer dans les mammographies, et en faisant varier l'échelle de l'ondelette, on peut obtenir des résultats de détection satisfaisants. D'autres formalismes sont

aussi possibles, ainsi on peut trouver l'utilisation de la logique floue pour la détection des calcifications (Bothorel, 1996; Cheng *et al.*, 1998; Rick, 1999) ou encore des approches utilisant des filtrages de la dynamique de l'image à l'aide d'outils de morphologie mathématique (Grimaud, 1991).

1.3.5 Détection des opacités

La détection des opacités est une tâche un peu plus complexe que pour les microcalcifications, notamment à cause de la variabilité qui existe entre différents types de lésions comme on a pu le voir à la section 1.1.2.

Extraction de marqueurs

La détection des zones de sur-densité est une première étape pour la détection de zones suspectes dans le sein. Ainsi Heath et Bowyer (2000) proposent d'utiliser une mesure de détection des sur-densités en calculant la proportion de pixels situés autour d'une lésion qui ont une intensité inférieure au minimum de l'intensité à l'intérieur. Dans la même optique de détection de sur-densités, d'autres approches tentent de détecter les zones contrastées à l'aide d'un chapeau mexicain de la même manière que pour les microcalcifications. Néanmoins ce genre d'approche ne semble pas être bien adapté à la taille des objets à détecter (te Brake et Karssemeijer, 1999; Peters, 2007). En effet, une opacité est généralement beaucoup plus étendue qu'une microcalcification, avec une variabilité en termes de forme plus importante. Ce dernier point met à mal le noyau de filtre utilisé qui fait généralement l'hypothèse d'un objet circulaire à détecter. À cela s'ajoute un contraste pouvant être moins élevé rendant de telles étapes de marquage délicates. Des variantes où le contraste est calculé à partir des cercles, l'un tournant autour de l'autre, ont été proposées (Peters *et al.*, 2006a) permettant de pallier certaines des précédentes limitations.

Vachier (1995) a proposé l'utilisation d'outils de morphologie mathématique pour l'extraction automatique d'opacités. L'idée utilisée est de se servir d'algorithmes de ligne de partage des eaux combiné à un critère de dynamique modélisant le contraste des structures contenues dans l'image.

D'autres approches reposent sur le fait que les lésions malignes présentent souvent une forme irrégulière voire des spicules. Ainsi, Bornefalk (2005) propose l'utilisation de filtres en quadrature pour la détection de ces spicules. Zou *et al.* (2008) quant à eux proposent une décomposition en ondelettes sur deux niveaux permettant le calcul de mesures dérivées du flot de gradient (Xu et Prince, 1997) combinées à d'autres mesures reposant sur l'analyse locale des histogrammes d'orientations de contour. D'autres travaux se placent dans un autre domaine de représentation pour pouvoir travailler plus facilement sur les spicules. Ainsi, Sampat et Bovik (2003) proposent de se placer dans le domaine de Radon où les lignes ont la propriété d'être représentées par des points pour rehausser les spicules. De retour dans le domaine classique de représentation de l'image le filtrage à l'aide de filtres dont le noyau est un anneau où les intensités sont calculées à partir de sinus ou de cosinus, permet de détecter les zones suspectes qui sont centres de convergence de spicules. D'autres travaux définissent des critères de convergence à partir de cercles emboîtés (Karssemeijer et te Brake, 1996; te Brake et Karssemeijer, 1999). En pratique la considération des spicules est un critère pertinent qu'il est intéressant de considérer lors de l'élaboration d'un système de détection car il est très fréquemment associé à des lésions malignes.

Les lésions variant aussi du point de vue de la taille, des approches multi-résolution sont souvent proposées (Liu *et al.*, 2001; Zou *et al.*, 2008). Néanmoins une étude sur la pertinence de telles approches (te Brake et Karssemeijer, 1999) a mené à la conclusion que le choix de l'approche multi-échelle se justifiait surtout pour des méthodes destinées à la détection de lésions dans un faible intervalle de taille. En effet, si la réponse d'un filtre décroît rapidement lorsque l'on s'éloigne de la taille a priori de la structure recherchée, il est nécessaire de considérer plusieurs échelles. D'un autre côté, si l'approche est peu sensible aux variations par rapport à la taille supposée de l'objet, alors l'intérêt de plusieurs échelles est plus relatif.

D'autres méthodes plus marginales existent aussi. Ainsi, l'analyse fractale a été proposée comme première étape pour la détection de marqueurs (Zhen et Chan, 2001). La granulation qui consiste globalement à sous-quantifier les niveaux de gris pour extraire des composantes connexes emboîtées vérifiant un critère de concentricité a aussi été proposée comme étape de marquage par Eltonsy *et al.* (2007).

Dans tous les cas, les méthodes proposées permettent de sélectionner de manière grossière, c'est-à-dire avec beaucoup de faux positifs, les zones qui seront analysées plus finement par la suite.

Segmentation

Les approches précédentes sont généralement conçues pour la détection de zones potentiellement suspectes. Néanmoins bien souvent pour vérifier si ces dernières contiennent une lésion maligne, il est nécessaire d'extraire la forme de lésion potentiellement détectée. Les approches les plus simples reposent sur des seuillages globaux de l'image (Martinez *et al.*, 1999; Brzakovic *et al.*, 1990). Une autre approche souvent utilisée fait appel aux méthodes de croissance de région (Kallergi *et al.*, 1992). L'idée est de faire croître un masque représentant la lésion en agrégeant des pixels voisins de la région qui ont des propriétés semblables à cette dernière (Lee *et al.*, 2000). La difficulté avec ce type d'approche est la définition d'un critère d'arrêt pertinent pour obtenir des régions segmentées de manière optimale.

De manière plus générale des méthodes de segmentation générique utilisant de la morphologie mathématique (Vachier, 1995; Bruynooghe, 2006) ou des outils de détection de contours (Torre et Poggio, 1986; Canny, 1986; Mallat et Zhong, 1992) peuvent être appliquées en mammographie (Abdel-Mottaleb *et al.*, 1996). Cependant ces outils assez simples ne permettent bien souvent d'obtenir des résultats satisfaisants qu'en combinaison avec des approches plus évoluées.

De manière similaire, les méthodes de segmentation par contour actifs (Kass *et al.*, 1988) ou ensembles de niveaux (Osher et Sethian, 1988) ont aussi été utilisées pour la segmentation de lésions. L'idée est de formuler une énergie sur une représentation d'un contour (ensemble de points, forme paramétrique ou fonction implicite d'ordre supérieur) de telle manière qu'elle prenne en compte des informations a priori sur ce que l'on cherche à segmenter (la régularité du contour par exemple) ainsi que des informations d'attache aux données (la segmentation doit correspondre aux zones de transition présentes dans l'image par exemple). Certains travaux (Ball et Bruce, 2007a,b,c) présentent une utilisation intensive des modèles déformables. Ainsi ils proposent une première segmentation du noyau de la lésion, puis une segmentation de sa périphérie et le cas échéant une troisième segmentation des spicules (Ball et Bruce, 2007a) en utilisant des ensembles de niveaux. Dans d'autres méthodes, les contours actifs sont utilisés pour affiner une première segmentation assez grossière (Sahiner *et al.*, 1996b, 2001). La qualité des segmentations obtenues avec ce type d'approche dépend beaucoup de l'énergie utilisée pour juger de la qualité des contours. Cette dernière comporte souvent beaucoup de paramètres et sa minimisation est bien souvent locale.

D'autres approches reposant sur l'expression du problème encore sous forme d'énergie, mais cette fois avec un formalisme de chemins optimaux ont aussi été présentées (Timp et Karssemeijer, 2004). L'idée est d'utiliser des méthodes de programmation dynamique pour trouver le contour optimal de la lésion. L'avantage de ce type d'approches est d'offrir des outils efficaces pour optimiser l'énergie, cette fois ci de manière globale. Cependant l'espace de recherche des contours est quant à lui limité, interdisant ainsi de considérer certains contours que les modèles déformables peuvent trouver.

Comme bien souvent dans d'autres champs applicatifs, la piste des méthodes de segmentation par régions a aussi été investiguée. Ainsi les champs de Markov aléatoires peuvent aussi être utilisés pour obtenir une segmentation des lésions (Zhen et Chan, 2001). Dans ces derniers travaux, les champs de Markov sont combinés à une segmentation multirésolution.

Des méthodes plus spécifiques du problème de segmentation de lésions dans le sein ont aussi été proposées. Ainsi Kupinski et Giger (1998) proposent de considérer des partitions de la zone à segmenter et de considérer que la meilleure segmentation correspond à celle qui maximise un indice de gradient radial. Les mêmes auteurs proposent aussi une modélisation statistique pour trouver la partition optimale. La difficulté dans l'utilisation de ces approches est d'obtenir un ensemble de partitions de qualité indépendamment de la structure à segmenter.

Ici aussi, l'utilisation de la logique floue pour le partitionnement de mammographies a été investiguée (Sameti et Ward, 1996). De même, des travaux relatant l'utilisation des méthodes par croissance de régions dans le cadre du flou ont été proposées (Gulati *et al.*, 1998). On remarquera que l'utilisation des extensions floues des différentes méthodes de segmentation courantes (Kanzaki, 1992) se fondent généralement sur ce formalisme dans leur fonctionnement et non dans leur décision finale. Ainsi si ambiguïté il y a, cette dernière risque d'être perdue une fois l'image segmentée.

Des travaux un peu plus marginaux proposent de traiter des lésions qui ne sont visibles qu'en partie dans l'image à cause, par exemple, d'une localisation trop proche du bord de cette dernière (Hatanaka *et al.*, 2001). Ce genre d'approche peut constituer une amélioration d'un système de détection en traitant des cas identifiés comme difficiles. De manière plus générale, une revue assez complète des outils utilisés pour la segmentation ainsi que pour les autres parties d'un système de détection a été proposée par Cheng *et al.*

(2006).

1.3.6 Prise de décision

La prise de décision se fait généralement après une étape de caractérisation. D'un point de vue haut niveau, on fait des mesures à partir des résultats de détection et/ou de segmentation (selon que l'on travaille sur des microcalcifications ou des opacités) dans le but de prendre dans un second temps une décision à l'aide de méthodes de classification standard.

Extraction de caractéristiques

Dans le cas des microcalcifications, une étape préliminaire d'extraction de caractéristiques est donc requise (Dengler *et al.*, 1993; Bankman *et al.*, 1994; Cheng *et al.*, 2003). Parmi ces caractéristiques, on peut en trouver qui portent sur la texture comme l'analyse des régions avoisinantes, la dépendance spatiale des niveaux de gris, les statistiques de longueurs sur les niveaux de gris ou encore la différence de niveaux de gris (Weszka *et al.*, 1976; Galloway, 1975; Kim et Park, 1999).

En ce qui concerne les opacités, la première classe de caractéristiques utilisées repose sur la forme de l'objet (Vachier, 1995; Sahiner *et al.*, 2001; Bruynooghe, 2006). On peut citer par exemple la compacité, le rapport de Ferret, ou encore le périmètre. D'autres mesures s'appuient sur l'analyse des distributions de la longueur radiale normalisée (distance entre le centroïde de la lésion et les points du contour) et des longueurs de cordes normalisées (Kilday *et al.*, 1993; Rangayyan *et al.*, 1997). Ces mesures ne permettent que de caractériser le résultat de la segmentation. Pour cette raison leur intérêt dépend essentiellement de la qualité de cette dernière et elles ont généralement besoin d'être combinées à d'autres descripteurs pour permettre une caractérisation satisfaisante.

Des mesures sur l'intensité peuvent aussi être utilisées (Huo *et al.*, 1998). On peut par exemple analyser la distribution des niveaux de gris dans la lésion supposée, sur son contour ou dans son voisinage. On peut aussi utiliser des mesures plus robustes qui sont invariantes par changement de contraste (Bruynooghe, 2006). Ce type de mesures apporte une information complémentaire par rapport aux précédentes mesures.

Une caractérisation importante de certaines lésions est la présence de spicules au niveau de leur périphérie. Pour prendre cet élément en compte, une série de mesures permettant d'évaluer à quel point la lésion est spiculée ont aussi été proposées (Viton *et al.*, 1996). Elles reposent essentiellement sur l'analyse des orientations locales des contours (Kegelmeyer *et al.*, 1994). Ces orientations peuvent par exemple être obtenues à partir d'une décomposition en ondelettes (Li *et al.*, 1997). Sur le même principe, des travaux proposent d'évaluer la complexité d'un contour par analyse fractale (Kim et Kim, 2005).

L'information fournie par l'association du contour et du contenu de l'image peut aussi être exploitée en extrayant et en analysant une bande plus ou moins large qui suit le contour de la lésion (Sahiner *et al.*, 1998; Mudigonda *et al.*, 2001). Cette bande est généralement étirée de manière à rendre le contour linéaire avant analyse. L'analyse se fait par l'extraction de mesures sur cette bande, comme l'étude des orientations du contour (Domínguez et Nandi, 2009). Ce type d'approche est donc dépendant du contour, ainsi les mesures obtenues vont être fortement liées aux caractéristiques des segmentations considérées. Par exemple, si l'on utilise un contour qui suit bien les spicules, ces dernières apparaîtront de manière linéaire dans l'image de bande. Inversement, si le contour coupe les spicules, on les verra apparaître de manière plus prononcée dans la même image.

Une autre classe de mesures repose sur l'analyse de texture (Haralick *et al.*, 1973). Ainsi on peut citer les mesures dérivées de matrices de cooccurrence des niveaux de gris (Wei *et al.*, 1997; Sahiner *et al.*, 2001; Llobet *et al.*, 2005; Mudigonda *et al.*, 2001). On remarquera néanmoins que le calcul de tels attributs est assez lourd.

Sélection de caractéristiques

Un trop grand nombre de caractéristiques pour les individus à classer peut nuire à la performance du classifieur. En effet, cela peut augmenter de manière non nécessaire l'espace de travail empêchant par la même occasion une généralisation efficace lors de l'apprentissage de la partie classification. Des heuristiques, comme l'aire sous la courbe ROC du classifieur (Kupinski et Giger, 1999), permettent de modéliser la qualité d'un ensemble de caractéristiques. Cela permet d'utiliser des méthodes d'optimisation comme les

algorithmes génétiques (Holland, 1992) pour obtenir l'ensemble optimal. Ce genre d'approche a été mis en œuvre avec succès en mammographie (Dhawan *et al.*, 1996; Kupinski et Giger, 1997).

Classification

L'étape de classification a pour but de donner la réponse finale sur ce qui est détecté et considéré comme un signe de lésion. L'idée est de combiner les informations extraites précédemment pour obtenir une décision. On peut voir cela de manière similaire à ce que fait le radiologue pour prendre sa décision. Le problème de classification n'est pas propre à la détection de lésions dans le sein, en effet c'est un champ de recherche à part entière qui peut être utilisé dans des applications très différentes.

Une première famille de classifieurs comprend les méthodes de type plus proches voisins (Cover et Hart, 1967). L'idée pour classer un individu est de regarder les individus dont on connaît la classe et qui sont assez semblables. En mammographie, des travaux reposant sur cette technique assez simple dans son fonctionnement ont été proposés pour la classification d'opacités (Llobet *et al.*, 2005). Veldkamp et Karssemeijer (1999) ainsi que Zadeh *et al.* (2001) utilisent des méthodes de K plus proches voisins pour la classification de microcalcifications. Ce type d'approche a l'avantage de ne faire aucune hypothèse sur les données à classer. Néanmoins, son pouvoir de généralisation sur des données éloignées de celles utilisées pour l'apprentissage est assez limité et peut poser problème.

Une autre méthode de classification regroupe les réseaux de neurones (Rumelhart *et al.*, 1986; Lau, 1991). Le plus communément utilisé est le perceptron où les neurones sont reliés de manière causale. Dans le cas où le problème n'est pas linéairement séparable, plusieurs couches de neurones sont nécessaires. Ce type de réseaux a été utilisé pour la classification d'opacités (Floyd Jr. *et al.*, 1994; Cheng *et al.*, 1994; Fogel *et al.*, 1998; Zou *et al.*, 2008; Arbach *et al.*, 2003). Dans le cas des microcalcifications, les réseaux de neurones sont aussi souvent utilisés (Bourrelly et Muller, 1989; Chitre *et al.*, 1994; Bankman *et al.*, 1994; Kim et Park, 1999; Kramer et Aghdasi, 1999; Zheng et Regentova, 2006; Verma et Zakos, 2001). Une étude propose même de comparer les performances d'un réseau de neurones avec la classification de cinq radiologues (Jiang *et al.*, 1997) pour montrer que ce type de classifieur peut faire mieux que l'expert en termes d'aire sous la courbe de performances de détection (ROC). Dans certains cas, l'utilisation de réseaux de neurones de convolution a été proposée, ainsi Sahiner *et al.* (1996a) s'en servent pour les opacités.

Des approches reposant plus sur une modélisation statistique, connues sous le nom de réseaux de croyances bayésiens (Pearl, 1988), ont aussi été employées en mammographie. Woods *et al.* (1993) a montré le potentiel d'un tel classifieur en le comparant à d'autres méthodes comme les réseaux de neurones. Néanmoins Zheng *et al.* (1999) a mis en évidence que la façon dont étaient appris les deux types de classifieurs par l'intermédiaire de bases d'apprentissage était l'élément le plus crucial pour obtenir de meilleures performances.

Une autre approche classique de classification est connue sous le nom de machines à vecteurs de support. Dans certains cas pour correspondre à l'acronyme anglais SVM, on emploie le terme de séparateurs à vaste marge (Vapnik, 1995). L'idée est de trouver un hyperplan qui sépare deux classes en maximisant la marge de séparation des éléments les plus proches des deux classes. Bien entendu, bien souvent les classes ne sont pas linéairement séparables. Dans de tels cas, les éléments à classer sont transférés dans un espace de dimension supérieure où une séparation par un hyperplan aura plus de sens. Plusieurs travaux utilisent les SVM pour la classification des masses (Cao *et al.*, 2004; Campanini *et al.*, 2004; Bornefalk, 2005).

Les arbres de décision (Safavian et Landgrebe, 1991; Zighed et Rakotomalala, 2000) se composent de tests à effectuer sur les différentes composantes des éléments à classer. Ces différents tests prennent la forme de seuils et sont combinés entre eux de manière conjonctive lorsque l'on suit une branche de l'arbre. L'avantage premier de ce type de méthodes est la sémantique qui se cache derrière l'arbre. En effet on peut y voir une série de règles d'experts combinées avec des *et* et des *ou* logiques. Kegelmeyer *et al.* (1994) ont montré la possibilité d'utiliser un tel processus de classification pour la détection de lésions spiculées. L'utilisation de la logique floue a aussi été proposée dans le cadre des arbres de décision. L'idée est d'assouplir les tests faits aux différents nœuds pour être capable de considérer plusieurs possibilités pour les cas qui sont proches de la frontière de décision (Janikow, 1998). Dans certains cas, les arbres de décision sont aussi capables de manipuler des données imprécises (S. Bothorel, 1997). Le point délicat avec ce genre d'approche est l'apprentissage, notamment avec le dernier type d'arbres.

1.3.7 Autres approches

Des travaux sur des variantes ou des raffinements des méthodes précédemment évoquées existent. On peut citer par exemple la correspondance de motifs, où l'exploitation de plusieurs images soit en considérant la symétrie des seins droit ou gauche soit en utilisant plusieurs incidences du même sein. Certains types de lésions ne présentant pas de sur-densité, leur détection est aussi abordée dans la littérature.

Correspondance de motifs

Ainsi la correspondance de motifs (Lai *et al.*, 1989) peut être utilisée pour détecter des opacités. L'idée est de construire une base de données de régions d'intérêt contenant des lésions dans le but de faire des requêtes lorsque l'on désire savoir si une zone est suspecte ou non. Cette étape de requête consiste à calculer un critère de similarité, comme par exemple l'information mutuelle, entre la zone considérée et toutes celles contenues dans la base de données (Tourassi *et al.*, 2003). Un des points clés pour la mise en pratique de telles approches est la constitution d'une base de référence optimale. En effet, il faut que cette dernière contienne assez d'éléments représentatifs pour pouvoir espérer traiter correctement une nouvelle requête, et en même temps ne pas être trop redondante pour ne pas rendre l'approche inutilisable car trop longue.

Ce genre d'approches peut aussi avoir un intérêt dans des utilisations interactives, où le radiologue désire voir une liste de lésions similaires à celle qu'il regarde dans le but d'utiliser les résultats de biopsie de ces dernières pour diriger son diagnostic.

Exploitation de l'asymétrie

De manière générale les structures apparaissant sur les images de mammographie du sein droit et du sein gauche d'une même patiente s'ordonnent selon une certaine symétrie. Cette propriété peut être en partie invalidée en cas d'apparition d'une masse dans l'un des deux seins. Des travaux proposent l'alignement et la soustraction des images des deux latéralités pour mettre en évidence leur dissimilarité, et donc les lésions potentielles. Plus récemment, des méthodes reposant sur l'analyse de la dissimilarité des zones suspectes entre sein droit et sein gauche a été proposée (Tahmoush et Samet, 2006a,b).

Utilisation de plusieurs vues

Un examen de mammographie standard se compose généralement de deux incidences par sein (une vue cranio caudale et une vue médio latérale oblique). Comme on l'a remarqué au début de ce chapitre, l'utilisation de ces deux clichés permet de limiter les problèmes de superposition de tissus en aidant le radiologue à prendre une décision. Des travaux récents proposent d'utiliser ces deux vues dans le but d'améliorer les performances de la chaîne de détection pour les opacités (van Engeland *et al.*, 2006; van Engeland et Karssemeijer, 2007). L'idée est de trouver la correspondance entre les différentes vues pour affiner la décision. Huo *et al.* (2001) proposent quant à eux de combiner les deux vues classiques à une vue non standard telle qu'un cliché centré agrandi. Ce type de raisonnement ne peut cependant être perçu que comme une amélioration d'un système existant. En effet, avant l'appariement, une détection dans chaque image est généralement nécessaire. De plus, dans certains cas, on peut ne pas avoir accès à toutes les images.

Distorsions architecturales

Les distorsions architecturales ne présentant pas de sur-densité, les étapes précédentes ne peuvent généralement pas être utilisées pour les détecter. Dans la mesure où ce type de signes radiologiques traduit le plus souvent une contraction des structures laissant apparaître une forme stellaire, le problème est généralement posé de manière à mettre en évidence les zones fortement convergentes. Ainsi Karssemeijer et te Brake (1996) proposent une mesure de convergence d'un anneau vers un disque inclus dans le centre de ce dernier. L'analyse des champs d'orientations par portrait de phase peut aussi être utilisée (Ayres et Rangayyan, 2005; Ayres et Rangayyan, 2007). L'idée de ce genre d'approche est d'étudier analytiquement l'orientation des textures dans l'image pour différencier les motifs correspondant à des convergences suspectes des autres. D'autres approches utilisant l'indice de concentration (Matsubara *et al.*, 2005) permettent aussi d'analyser les orientations des structures composant l'image (Hara *et al.*, 2006).

1.4 Aide à la détection pour des volumes de tomosynthèse

Les travaux présentés dans la section précédente portaient essentiellement sur l'aide à la détection en mammographie standard. Certains de ces éléments peuvent plus ou moins directement être adaptés pour le traitement de volumes de tomosynthèse. Dans un premier temps nous reviendrons sur les motivations d'un outil de détection en 3D, puis nous présenterons de manière haut niveau les schémas de détection proposés dans la littérature. Nous poursuivrons par une description des approches de détection de microcalcifications et de masses. Nous finirons par une présentation sommaire de notre approche.

1.4.1 Motivations

Bien que la tomosynthèse ait le potentiel de réduire les problèmes de superpositions de structures que l'on retrouve en mammographie standard, les motivations pour la création d'un outil de détection automatique de signes radiologiques qui étaient valables en 2D sont toujours d'actualité. De plus, comme la quantité de données que le radiologue doit observer augmente de manière non négligeable (un volume 3D comptant dix à cinquante fois plus d'images à lire que deux clichés 2D), un tel système est plus que jamais utile afin d'assister le spécialiste dans sa tâche de lecture.

1.4.2 Différents designs de système de détection

Dans le cadre de la tomosynthèse, le choix de la structure de la chaîne de traitement n'est pas direct. En effet, différentes possibilités peuvent être envisagées. Ainsi on peut choisir de travailler dans le volume reconstruit (c.f. figure 1.10(a)), ou dans les projections (c.f. figure 1.10(b)), ou encore de faire des allers retours entre 2D et 3D (c.f. figure 1.10(c)).

Dans le premier cas, la chaîne de traitement peut se présenter d'une manière haut niveau comme la détection de marqueurs dans le volume, suivie d'une analyse plus en détails reposant par exemple sur la segmentation et la caractérisation des opacités, pour finir avec la prise de décision. Chacune de ces étapes peut se faire en 3D ou en pseudo-3D, c'est-à-dire en considérant tous les voisins des voxels dans l'espace ou en travaillant sur les coupes reconstruites de manière presque indépendante (on ne s'occupe pas des coupes adjacentes). En effet, contrairement aux images 3D plus classiques comme celles acquises en CT, la limitation de la géométrie du système menant à la distorsion en Z des objets peut rendre difficile la modélisation d'algorithmes de détection. Cette approche a le défaut majeur d'être dépendante de la reconstruction. Néanmoins, les techniques de reconstruction en tomosynthèse du sein semblent se stabiliser ce qui permet d'envisager ce type de traitements.

La deuxième approche consiste à travailler sur les projections s'affranchissant ainsi de la méthode de reconstruction. Ce genre d'approche permet d'exploiter directement les méthodes conçues pour les mammographies standard. Néanmoins, deux problèmes se posent lorsque l'on choisit une telle approche. Premièrement, la dose de chacune des projections en tomosynthèse est inférieure à celle de mammographies standard. Cela a pour effet principal de diminuer le rapport signal à bruit dans les images, rendant les tâches de détection et de caractérisation plus difficiles. Le deuxième point à prendre en compte est la fusion d'informations. En effet, dans un tel schéma, chacune des projections est traitée de manière indépendante, ainsi si on prend une décision trop tôt dans une de ces dernières, il peut être difficile voire impossible de rattraper l'erreur dans la suite du traitement.

La troisième approche peut être vue comme une approche hybride des deux premières. L'idée est de commencer l'analyse dans les projections en faisant le moins de prise de décision possible. Une fois cela fait, une étape d'agrégation permet de regrouper tous les éléments que l'on a pu rassembler de manière indépendante jusque là. Ensuite, une analyse 3D peut avoir lieu, soit à partir du volume, soit à partir de la fusion des informations 2D. Éventuellement on peut retourner dans les domaines des projections dans le but de rechercher de l'information que l'on aurait pu manquer, comme par exemple une lésion non détectée dans une des projections car elle y apparaît de manière subtile. Les itérations entre 2D et 3D peuvent permettre d'obtenir toutes les informations nécessaires pour prendre la *meilleure décision*.

1.4.3 Microcalcifications

Wheeler *et al.* (2006) proposent une méthode de détection de microcalcifications qui modélise à la fois

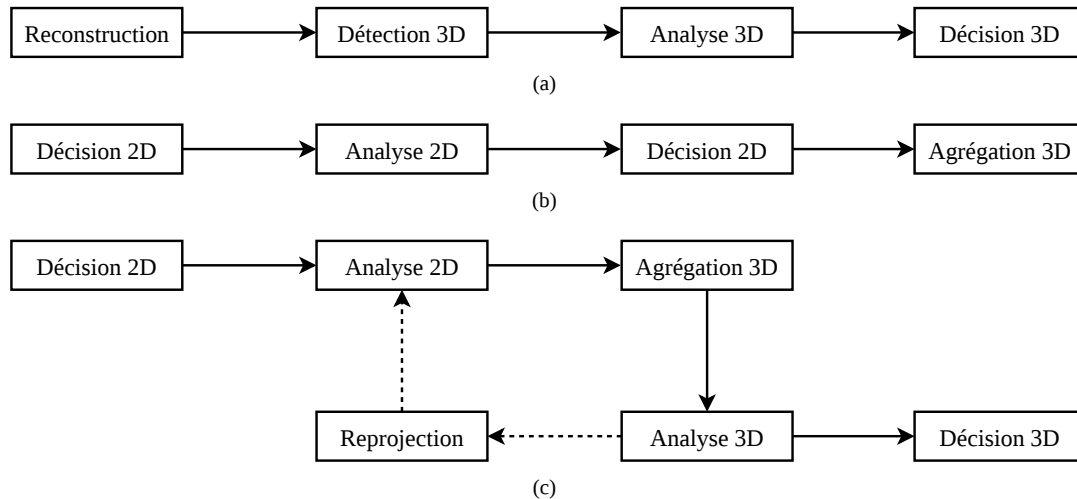


FIG. 1.10 – Différents schémas de détection en tomosynthèse.

le contenu du sein ainsi que le bruit présent dans les images acquises. Cette approche repose essentiellement sur le traitement des projections avant agrégation.

Une autre approche travaillant cette fois-ci en 3D a été proposée par Bernard *et al.* (2006). L'idée est de modéliser le bruit ainsi que la quantité de calcium présent en chaque voxel dans le cas d'une reconstruction par rétroprojection simple et de détecter les calcifications par convolution avec un chapeau mexicain. En effet dans ce cas on peut montrer que le rapport contraste à bruit dans le volume est bien meilleur que dans les projections. Des travaux plus récents ont montré que cette approche peut se reformuler de manière presque équivalente à un filtrage des projections avant rétroprojection, ce qui a pour avantage principal d'accélérer le traitement des données (Bernard *et al.*, 2007). La gestion des artefacts de rétroprojection ainsi que la validation de ce type d'approche sur une base de données plus conséquente ont aussi été proposées (Bernard *et al.*, 2008).

Des modélisations de plus haut niveau, notamment pour l'étape de fusion après détection dans les projections ont aussi été proposées. On peut citer les travaux de Peters *et al.* (2005) qui reposent sur la notion de contours flous, c'est-à-dire un ensemble de contours possibles, pour chaque microcalcification détectée. L'utilisation de la logique floue permet de fournir un cadre de travail avantageux lors de la prise de décision en 3D, dans la mesure où il permet d'être assez laxiste sur les décisions prises en 2D.

1.4.4 Densités

La détection des opacités est comme en mammographie conventionnelle toujours une étape délicate. Chan *et al.* (2004) ont proposé des travaux préliminaires pour traiter le problème reposant sur le traitement des projections de manière indépendante (c.f. figure 1.10(b)). Les mêmes auteurs (Chan *et al.*, 2005) ont aussi investigué la détection à partir de coupes reconstruites de manière itérative par maximum de vraisemblance (Wu *et al.*, 2003a).

D'autres équipes ont investigué les deux approches (volume/projections). Ainsi, Reiser *et al.* (2006) proposent de détecter les opacités dans les projections de manière indépendante puis de recombinaison ces détectations dans un volume portant aussi une information sur l'ouverture angulaire dans laquelle la lésion est visible. Leur investigation de la piste 3D repose sur l'utilisation d'un filtre d'analyse radiale, pour mettre en évidence les structures sphériques, associé à une projection par maximum d'intensité et un chapeau haut de forme (Reiser *et al.*, 2005). Ces travaux se limitent néanmoins juste à une détection grossière de zones suspectes à forte sensibilité mais produisant beaucoup de faux positifs. Ainsi les auteurs n'abordent pas les parties d'analyse et de décision que l'on peut voir sur la figure 1.10.

Peters (2007) a proposé un système défini dans le cadre de la logique floue faisant des boucles entre 2D et 3D pour la détection et la caractérisation des masses (c.f. figure 1.10(c)). Plus concrètement une première

détection dans les projections est faite de manière indépendante suivie d'une extraction de contours flous (Peters *et al.*, 2007) pour chacune des zones suspectes. Pour chacune de ces segmentations, des attributs flous sont calculés (Peters *et al.*, 2006b). Ces derniers correspondent à des mesures telles que définies dans la section 1.3.6 mais portant l'imprécision induite par les contours flous. Une fois ces traitements effectués pour toutes les projections, les différentes détections correspondant aux mêmes zones dans le sein sont associées et leurs caractéristiques sont agrégées dans le but d'utiliser un arbre de décision flou pour juger de la malignité de la lésion sous-jacente à la structure détectée. Ce genre d'approche a l'avantage de pouvoir manipuler l'imprécision et l'incertitude provenant des données et des filtres de détection utilisés, et ce jusqu'à la fin de la chaîne de détection où toute l'information est disponible pour prendre une décision.

Plus récemment, des études sur l'influence du nombre de projections et de la dose sur la détection ont été menées. Chan *et al.* (2008b) proposent de comparer les performances entre des examens contenant 21 projections et des examens contenant 11 projections. Les mêmes cas/images étant utilisés pour l'étude, les séries ne contenant que 11 projections ont une dose environ deux fois plus faible par rapport aux séries contenant 21 projections. La conclusion à laquelle arrivent les auteurs est qu'une meilleure détection est obtenue dans le cas où l'on utilise 21 projections. Néanmoins on peut s'interroger sur la part de responsabilité de la dose dans ce résultat.

La comparaison des différentes approches de détection a aussi été abordée dans la littérature. Ainsi, Chan *et al.* (2008a) ont fait une comparaison entre détection dans les volumes reconstruits, détection dans les projection et approches hybrides. La méthode purement 2D en est ressortie la moins performante alors que l'approche hybride a démontré le meilleur comportement. Ces travaux sont les plus aboutis notamment pour la validation sur une base de données assez conséquente (69 lésions malignes) compte tenu de la maturité de la modalité d'imagerie utilisée. Une comparaison de notre approche avec, entre autres, ces travaux sera largement détaillée au chapitre 8.

On peut enfin citer une dernière famille de travaux reposant sur l'utilisation d'outils de correspondance de motifs. Cette dernière a été proposée pour aider à la prise de décision, et notamment pour réduire le nombre de faux positifs générés par des systèmes de détection automatique d'opacités (Singh *et al.*, 2008b). Ces travaux seront aussi discutés au chapitre 8.

1.4.5 Approche proposée

L'approche que nous proposons ne considère que le volume reconstruit (voir figure 1.10(a)). L'hypothèse sous-jacente de ce choix est de considérer que les méthodes de reconstruction ont atteint une certaine maturité dans le domaine de la tomosynthèse numérique du sein.

Plus concrètement, notre approche se focalise non seulement sur la détection des opacités, au même titre que les travaux précédemment cités, mais aussi sur les distorsions architecturales. Pour cette raison, nous proposons un schéma à deux canaux respectivement dédiés aux deux précédents signes radiologiques. Ce schéma est présenté dans sa globalité à la figure 1.11.

Dans le cas du canal dédié à la détection des opacités, l'étape préliminaire de détection se fait à l'aide d'un nouveau type de filtres, à savoir les filtres connexes flous. Ces derniers permettent de prendre en compte les imperfections des images ainsi que des critères utilisés pour la détection. Le chapitre 2 les introduit de manière formelle tout en les étudiant de manière théorique. Les chapitres 3 et 4 traitent respectivement de leur mise en œuvre et du conditionnement des données pour de tels filtres. Les étapes d'analyse et de prise de décision dans ce canal sont quant à elles plus classiques : méthode de segmentation par programmation dynamique (Timp et Karssemeijer, 2004) et utilisation d'un classifieur fondé sur les SVM.

Le second canal utilise des principes similaires à ceux proposés par Karssemeijer et te Brake (1996) dans une modélisation a contrario (Desolneux *et al.*, 2000, 2001, 2003; Moisan, 2003; Moisan et Stival, 2004). Cette méthode sera présentée au chapitre 7. Le processus proposé permet de détecter les zones suspectes de convergence dans le volume. Pour chacune de ces zones, on extrait des caractéristiques servant d'entrée à un classifieur, encore une fois de type SVM, pour obtenir les zones vraiment suspectes.

Les volumes de tomosynthèse souffrant de défauts assez marqués notamment dans la direction perpendiculaire au détecteur (réplications et allongement de structures dans cette direction), les algorithmes proposés fonctionnent essentiellement coupe par coupe. En effet la prise en compte des défauts que subissent les objets reste assez délicate. Néanmoins l'aspect 3D n'est pas complètement écarté, ainsi les différentes parties des algorithmes s'assurent de la cohérence des résultats obtenus pour chacune des coupes, et ce à plusieurs

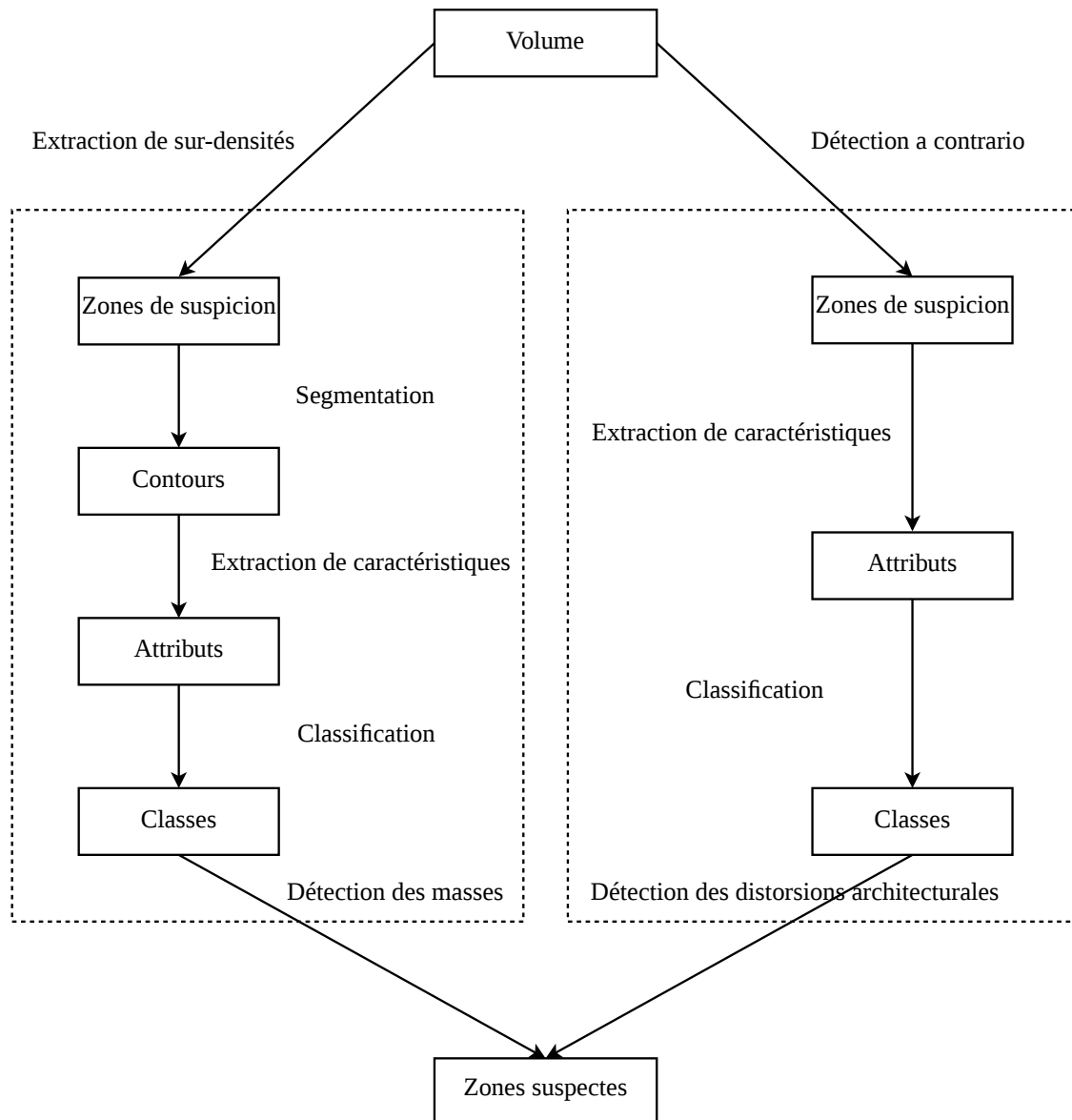


FIG. 1.11 – Schéma global de l'approche de détection proposée.

reprises dans toute la chaîne de détection.

Les deux canaux ayant des populations cibles de signes radiologiques distinctes, les zones détectées par chacun d'entre eux sont donc agrégées de manière disjonctive : pour qu'une zone soit suspecte, il faut qu'au moins un des deux canaux la considère comme telle.

De manière parallèle à cette chaîne de détection, nous illustrons aussi des expérimentations préliminaires sur l'utilisation de méthodes de segmentation et de classification floues dans le cas d'opacités à noyaux denses. Cette contribution est présentée au chapitre 6.

1.5 Conclusion

La détection automatique de cancers permet de fournir au radiologue des outils très utiles dans sa tâche de dépistage du cancer du sein. Parmi les signes qu'il est intéressant de détecter on peut citer les microcal-

cifications, les masses et les distorsions architecturales. Une grande panoplie de techniques existe pour les détecter en mammographie conventionnelle.

Avec l'arrivée de la tomosynthèse, l'intérêt d'outils de détection automatique reste d'actualité. Comme on l'a vu, la quantité de méthodes proposées dans la littérature pour détecter des signes suspects dans ce nouveau type d'imagerie est moins importante que pour la mammographie conventionnelle. Cependant des travaux d'une certaine maturité proposent déjà des solutions prometteuses. Dans ce contexte nous proposons une chaîne à deux canaux permettant la détection d'opacités et de distorsions architecturales. Pour cela nous introduisons plusieurs outils reposant sur la logique floue. Ces derniers seront étudiés dans les différents chapitres de ce manuscrit. Nous montrerons aussi que notre approche est compétitive par rapport aux méthodes de référence.

Chapitre 2

Filtres connexes flous

La segmentation de structures dans les images médicales est souvent exprimée comme l'extraction de régions connexes de radiométrie *stable* (Mumford *et al.*, 1988; Leclerc, 1989; Zhu et Yuille, 1996). Les filtres connexes que l'on retrouve dans la littérature entre autres dans les travaux de Klein (1985); Vincent (1993); Serra et Salembier (1993); Heijmans (1997) ont été introduits à cet effet et ont montré leur intérêt notamment pour accomplir des tâches de filtrage et de segmentation (Salembier et Serra, 1995; Salembier *et al.*, 1998; Vincent, 1993).

Cependant ces filtres semblent moins bien adaptés pour répondre à un problème de reconnaissance tel que l'extraction d'objets supposés connexes, de radiométrie stable et répondant à certains critères (tels qu'un cardinal *a priori*). En effet ne sont considérées comme zones plates que des régions de radiométrie constante, ce qui n'est pas le cas des objets recherchés en raison, en particulier, du bruit qui affecte l'image. Il en va de même pour les filtres connexes construits à partir de seuillages successifs de l'image dans la mesure où le bruit joue un rôle sur la qualité des composantes connexes extraites par leur intermédiaire.

Plutôt que de filtrer les images avant de réaliser la reconnaissance, ce qui constitue l'approche classique, nous intégrons dans la représentation de l'image l'imprécision affectant les intensités observées. Les ensembles flous fournissent un cadre approprié pour la représentation de cette imprécision, ainsi que pour la représentation des régions. L'image est alors représentée sous forme d'image floue où le niveau de gris en chaque point est modélisé par une quantité floue (ensemble flou défini sur l'ensemble des niveaux de gris possibles). En couplant ce type de représentation à des définitions classiques de connexité floue, on peut proposer une approche innovante par rapport à ce qui existe dans la littérature pour assouplir la notion de zone plate (Wilkinson, 2007; Braga-Neto et Goutsias, 2003a), mais aussi plus généralement aux définitions de filtres connexes.

Dans un premier temps, nous donnerons quelques rappels sur ce que sont les filtres connexes et comment ils sont utilisés avec des images à niveaux de gris. Ensuite, nous introduirons un nouveau formalisme pour représenter l'imprécision contenue dans les images. Ce dernier nous permettra dans un troisième temps de proposer une extension dans le domaine des ensembles flous des filtres connexes. Nous finirons enfin par présenter une utilisation concrète de ce nouveau type de filtrage sur des images de tomosynthèse du sein.

2.1 Les filtres connexes

Nous notons Ω l'espace discret $\mathbb{Z}^n$ muni d'une connexité discrète c_d et représentant le support des images à niveaux de gris. En notant $\mathcal{G}$ l'ensemble des niveaux de gris d'une image, on définira $\mathcal{I}$ comme l'ensemble des images à niveaux de gris $I : \Omega \rightarrow \mathcal{G}$. On notera 2^Ω l'ensemble des sous-ensembles de Ω .

Les filtres connexes pour les images à niveaux de gris reposent essentiellement sur la capacité d'extraire des composantes connexes binaires de l'image et de les filtrer grâce à un opérateur connexe binaire. Le but de cette section n'étant que de donner une idée de ce que sont les filtres connexes, nous rappellerons brièvement dans un premier temps la définition binaire de ce type d'opérateurs, puis nous parlerons de leurs extensions aux images à niveaux de gris pour enfin finir par un tour d'horizon sur les grandes classes d'opérateurs connexes dédiés à ces dernières images.

L'annexe A regroupe quelques notions qui servent de base aux travaux proposés dans ce chapitre. On pourra aussi se référer à d'autres ouvrages pour approfondir ces notions (Klette et Rosenfeld, 2004).

2.1.1 Opérateurs connexes binaires

Un opérateur connexe binaire repose sur la notion de connexité. Ainsi on dira qu'un ensemble est connexe s'il vérifie la définition suivante.

Définition 2.1.1. Un ensemble $N \in \Omega$ est connexe si pour tout couple de points de cet ensemble, il existe un chemin (au sens de la connexité c_d définie sur Ω) reliant ces deux points tels que tous les points de ce chemin appartiennent à N .

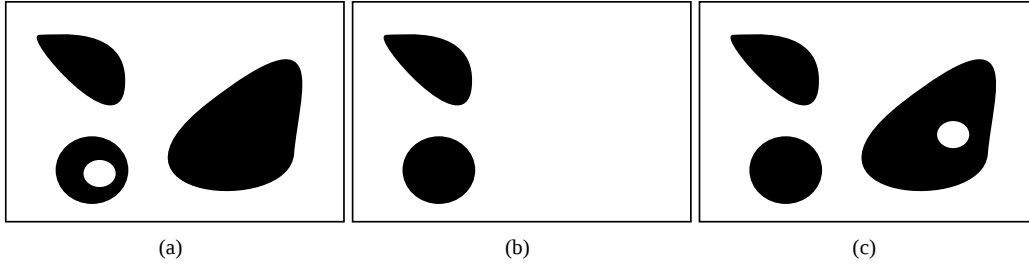


FIG. 2.1 – Filtrage d'un ensemble binaire (a) avec un opérateur ϕ connexe (b) et un opérateur ϕ non connexe (c).

Définition 2.1.2. L'ensemble $\mathcal{C}(N)$ des composantes connexes d'un ensemble $N \subseteq \Omega$ est l'ensemble des ensembles connexes de taille maximale inclus dans N .

Définition 2.1.3. Pour tout ensemble net N , on définit la composante connexe associée à un point p comme l'ensemble des points $p' \in \Omega$ tels que $\hat{\Gamma}_N^p(p') = 1$ avec $\hat{\Gamma}$ défini comme :

$$\forall (p, p') \in \Omega^2 \quad \hat{\Gamma}_N^p(p') = \begin{cases} 1 & \text{si } p \text{ et } p' \text{ sont connexes dans } N \text{ ou } \overline{N} \\ 0 & \text{sinon} \end{cases}$$

Lors de la définition de la connexité discrète c_d sur Ω , on remarquera que les connexités utilisées dans N et son complémentaire $\overline{N}$ ne sont pas forcément les mêmes. Par exemple si on utilise une 4-connexité pour l'objet (N), on devra utiliser une 8-connexité pour le fond ($\overline{N}$). On se reportera à l'annexe A.2 pour plus d'information sur les différents types de connexité.

Enfin on utilisera la définition suivante pour dire qu'un opérateur est connexe :

Définition 2.1.4. Un opérateur $\phi : 2^\Omega \rightarrow 2^\Omega$ est un opérateur connexe si :

$$\forall N \subseteq \Omega \quad (\mathcal{C}(N \cap \overline{\phi(N)}) \subseteq \mathcal{C}(N)) \wedge (\mathcal{C}(\overline{N} \cap \phi(N)) \subseteq \mathcal{C}(\overline{N}))$$

Plus concrètement, cela signifie qu'un opérateur est connexe si l'ensemble des composantes connexes contenues dans l'ensemble d'entrée contient les composantes connexes formées par les éléments qui sont passés de l'état d'objet à fond par le filtrage. De même, un tel opérateur doit s'assurer que les composantes connexes formées par les éléments passant du statut de fond à objet sont bien incluses dans l'ensemble des composantes connexes contenues dans le fond de l'ensemble d'entrée. La définition précédente se traduit par l'idée que chaque composante connexe du fond (resp. de l'objet) peut après filtrage par ϕ soit devenir complètement de l'objet (resp. du fond) ou rester telle quelle. Il est important de remarquer qu'ainsi une composante connexe ne peut être coupée par l'opérateur. Un opérateur connexe ne peut donc pas créer de nouveaux contours : il peut juste en garder ou en supprimer. La figure 2.1 présente un exemple de ce processus. Dans le cas de l'opérateur connexe (b), un objet (en noir) disparaît alors que dans le cas de l'opérateur non connexe (c), un objet se voit modifié par l'ajout d'une nouvelle composante connexe appartenant au fond (en blanc), invalidant ainsi la première partie de la définition 2.1.4.

2.1.2 Extension aux images à niveaux de gris et filtres de nivellement

Selon la définition donnée par Salembier et Serra (1995), les zones plates (*flat-zones*) d'une image $I : \Omega \rightarrow \mathcal{G}$ sont les plus grandes régions connexes de Ω sur lesquelles I est constante. Les zones plates forment une partition de l'image. Leur extraction et leur filtrage peuvent être réalisés efficacement. Cependant la notion d'homogénéité radiométrique considérée est très stricte et peu robuste au bruit qui peut altérer l'image.

Définition 2.1.5. Un opérateur connexe $\delta : \mathcal{I} \rightarrow \mathcal{I}$ pour les images à niveaux de gris est un opérateur qui ne crée pas de nouveaux contours dans l'image filtrée, c'est-à-dire que toute zone plate qui peut être extraite de l'image d'entrée est incluse dans une zone plate extraite de l'image filtrée :

$$\forall I \in \mathcal{I}, \forall p \in \Omega \quad \hat{\Gamma}_{I^p}^p \subseteq \hat{\Gamma}_{\delta(I)^p}^p$$

avec $\forall p \in \Omega$ l'ensemble I^p définit comme suit : $\forall p' \in \Omega \ (p' \in I^p) \Leftrightarrow (I(p') = I(p))$

Ainsi les zones plates contenues dans l'image d'entrée ne peuvent pas être cassées en deux.

Une grande classe d'opérateurs connexes utilisés dans la littérature (Meyer, 2005) sont les filtres de nivellement. Ces derniers sont des opérateurs connexes qui préservent la direction des transitions locales. Ainsi, un minimum local ne peut pas être transformé en maximum local et inversement. Ces filtres suppriment entre autres les extrema mais ne peuvent en aucun cas en créer de nouveaux.

Bien sûr, sans autres contraintes sur les définitions, on ne peut généralement pas parler de filtres au sens morphologique du terme dans la mesure où l'on ne peut généralement pas garantir des propriétés comme l'idempotence ou la croissance de tels opérateurs (c.f. annexe A.3 pour plus de détails). On remarquera que selon les applications, des filtres de nivellement n'ayant pas ces propriétés peuvent s'avérer néanmoins intéressants. On peut citer pour l'exemple l'arasement volumique introduit par Vachier (2001) ou des variantes qui sont utilisées en classification (Géraud *et al.*, 2004).

2.1.3 Ouvertures d'attribut et filtres d'amincissement (thinning)

D'un point de vue plus pratique, Breen et Jones (1996) ont introduit deux grandes classes de filtres : les ouvertures d'attribut et les filtres d'amincissement. Le principe général est de définir un filtre connexe binaire par l'intermédiaire d'un critère C qu'une composante connexe doit respecter pour être conservée dans le but de filtrer tous les seuillages possibles de l'image d'entrée. Plus concrètement, le critère peut être par exemple l'aire de la composante, ou le diamètre d'un disque englobant la composante. De manière générale, on peut écrire un filtre δ de la façon suivante :

$$\forall I \in \mathcal{I}, \forall p \in \Omega \quad \delta(I)(p) = \max \left\{ t/C \left(\hat{\Gamma}_{X_t^+(I)}^p \right) = 1 \right\}$$

avec $X_t^+(I)$ l'opérateur de seuillage qui retourne l'ensemble des points dont l'intensité est supérieure ou égale à t et C un critère à valeur dans $\{0, 1\}$ que les composantes connexes doivent respecter.

La différence entre ces deux types de filtres réside dans la propriété de croissance que C peut avoir. En effet, en cas d'absence de cette propriété (cas des amincissements), on ne se trouve plus dans un cadre où le filtre ne fait que réduire des maxima de l'image. Cela a pour effet de produire généralement des images qui contiennent des contours artificiels. Néanmoins les deux types de filtres peuvent avoir des comportements adaptés à différentes applications (simplification d'image, segmentation, etc.). Ces différents filtres sont discutés plus en détails à l'annexe A.4.1.

2.2 Représentation de l'imprécision dans les images

Nous allons maintenant voir comment on peut représenter l'imprécision contenue dans une image à niveaux de gris. Pour cela, nous allons utiliser le cadre théorique de la logique floue. Ainsi, après un bref rappel sur les notations utilisées pour manipuler les ensembles flous, nous introduirons un formalisme pour représenter des images à niveaux de gris flous. Enfin nous définirons deux classes d'images floues et motiverons le formalisme introduit par rapport aux utilisations du flou proposées dans la littérature.

2.2.1 Ensembles flous

Un ensemble flou sur Ω est représenté par sa fonction d'appartenance $f : \Omega \rightarrow [0, 1]$. Nous nous restreignons à des ensembles flous dont le support est borné. Nous notons f_α une α -coupe de f et $\mathcal{S}$ l'ensemble des ensembles flous définis sur Ω . En considérant la relation d'ordre usuelle $\leq$ sur $\mathcal{S}$, $(\mathcal{S}, \leq)$ est un treillis complet. Le supremum $\vee$ et l'infimum $\wedge$ sont respectivement le max et le min. Nous notons $0_{\mathcal{F}}$ le plus petit élément et $1_{\mathcal{F}}$ le plus grand élément de $\mathcal{S}$. On remarquera que l'ensemble 2^Ω des ensembles inclus dans Ω est inclus dans $\mathcal{S}$. Pour des raisons de simplicité d'écriture, lorsque nous manipulerons un ensemble flou de fonction d'appartenance f , nous ne parlerons que de f . Enfin on définira $\mathcal{K}$ comme l'ensemble des sous-ensembles flous définis sur $\mathbb{R}^+$ ($\mathbb{R}^+ \rightarrow [0, 1]$).

Certaines des notions usuelles sur les ensembles flous qui sont utilisées dans les travaux présentés ici sont rappelées à l'annexe B.

2.2.2 Les image floues

Pour représenter l'imprécision d'une image nous proposons de ne plus représenter les valeurs des pixels par des valeur nettes, mais plutôt d'utiliser des quantités floues. Ainsi une image floue sera définie comme un ensemble flou sur $\Omega \times \mathcal{G}$ ($\Omega \times \mathcal{G} \rightarrow [0, 1]$). On notera $\mathcal{F}$ l'ensemble de ces images floues. On notera $2^{\Omega \times \mathcal{G}}$ les images floues de $\mathcal{F}$ dont les degrés d'appartenance valent 0 ou 1. Bien entendu, selon la nature de ces quantités floues, la sémantique de l'image pourra changer comme on le verra par la suite.

Ω	domaine borné de l'image muni d'une connexité discrète
$\mathcal{G}$	ensemble des niveaux de gris
$\mathcal{I}$	ensemble des images naturelles $I : \Omega \rightarrow \mathcal{G}$
$\mathcal{S}$	ensemble des ensembles flous définis sur Ω ($\Omega \rightarrow [0, 1]$)
$\mathcal{F}$	ensemble des ensembles flous définis sur $\Omega \times \mathcal{G}$ ($\Omega \times \mathcal{G} \rightarrow [0, 1]$)
2^Ω	ensembles des ensembles inclus dans Ω ($\Omega \rightarrow \{0, 1\}$). De manière évidente, on a : $2^\Omega \subset \mathcal{S}$.
$2^{\Omega \times \mathcal{G}}$	ensembles des ensembles inclus dans $\Omega \times \mathcal{G}$ ($\Omega \times \mathcal{G} \rightarrow \{0, 1\}$). De manière évidente, on a : $2^{\Omega \times \mathcal{G}} \subset \mathcal{F}$.
$\mathcal{K}$	Ensemble des sous-ensembles flous définis sur $\mathbb{R}^+$ ($\mathbb{R}^+ \rightarrow [0, 1]$)

TAB. 2.1 – Notations générales sur les images floues.

Dans un souci de clarté dans la lecture de ce chapitre, nous proposons de résumer les notations essentielles qui ont été introduites jusqu'ici dans le tableau 2.1. Pour un $\alpha \in [0, 1]$, et une image floue $F \in \mathcal{F}$, F_α représente une α -coupe de F . On utilisera aussi, pour un niveau de gris donné $g \in \mathcal{G}$, la notation $F(*, g)$ pour définir l'ensemble flou $f \in \mathcal{S}$ qui vérifie $\forall p \in \Omega \ f(p) = F(p, g)$. De même pour un point $p \in \Omega$, $F(p, *)$ représentera l'ensemble flou f sur $\mathcal{G}$ défini comme $\forall g \in \mathcal{G} \ f(g) = F(p, g)$.

2.2.3 Images d'ombres floues

La première grande famille d'images floues que nous allons introduire est celle des image d'ombres floues. Ce type d'images a pour vocation d'étendre les filtres connexes de type ouverture d'attribut et amincissement où les composantes connexes considérées lors du filtrage sont extraites par seuillages multiples de l'image (c.f. annexe A.4.1).

Définition 2.2.1. $F \in \mathcal{F}$ est une image d'ombres floues (IOF) ssi :

$$\forall p \in \Omega, \forall g_1 \in \mathcal{G}, \forall g_2 \in \mathcal{G} \quad g_1 \leq g_2 \Rightarrow F(p, g_1) \geq F(p, g_2)$$

Une image d'ombres floues est une extension directe du concept d'image d'ombres binaire : on augmente la dimension d'une image en niveaux de gris pour obtenir une image binaire de dimension $n + 1$ (c.f. annexe A.1 pour plus de détails). Dans notre cas on considère des images en niveaux de gris avec une imprécision sur ces derniers qui se traduit par des ensembles flous dans l'image d'ombres.

Pour une image d'ombres floues F , pour tout $p \in \Omega$ et $g \in \mathcal{G}$, $F(p, g)$ représente le degré avec lequel l'image représentée est supérieure ou égale à g au point p . L'image d'ombres floues permet donc entre autres de représenter l'imprécision sur les niveaux de gris d'une image classique. Une exemple d'image d'ombres floues est présentée à la figure 2.2.

Les *IOF* étant une extension des images d'ombres, ces dernières peuvent à part entière être considérées comme des images floues. Dans ce formalisme, on peut introduire un opérateur um qui décrit un processus de conversion d'image en niveaux de gris vers des images d'ombres classiques.

Définition 2.2.2. L'opérateur de conversion d'une image classique vers une *IOF* $um : \mathcal{I} \rightarrow 2^{\Omega \times \mathcal{G}}$ est défini comme :

$$\forall I \in \mathcal{I}, \forall p \in \Omega, \forall g \in \mathcal{G} \quad um(I)(p, g) = \begin{cases} 1 & \text{si } g \leq I(p) \\ 0 & \text{sinon} \end{cases}$$

Cette définition est conforme à la définition classique évoquée en annexe A.1 : pour chaque point $p \in \Omega$, seuls les niveaux de gris inférieurs ou égaux à l'intensité de l'image en ce point appartiennent à l'image d'ombres. Cependant, les images appartenant à $\mathcal{I}$ ne disposant pas d'imprécision sur les niveaux des gris, l'image d'ombres obtenue est donc binaire, mais peut être considérée comme un sous-ensemble flou. Des méthodes de construction plus évoluées seront présentées dans le chapitre 4. La nature de l'imprécision représentée par ces images est fortement liée à leur construction. Comme on le verra au chapitre 4, on peut représenter de l'ambiguïté lié au bruit dans les images ou des défauts liés aux conditions d'acquisition (artefacts de reconstruction en imagerie 3D par exemple). Nous détaillerons principalement l'impact du bruit dans ce même chapitre 4.

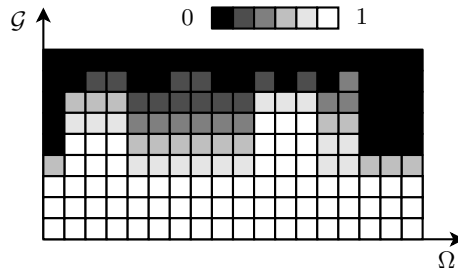


FIG. 2.2 – Exemple d'une image d'ombres floues définie sur un domaine Ω 1D.

De manière similaire, on peut définir un outil simple pour la suppression de l'imprécision d'une *IOF* :

Définition 2.2.3. L'opérateur $im : [0, 1] \times \mathcal{F} \rightarrow \mathcal{I}$ de conversion d'une *IOF* vers une image nette est défini comme :

$$\forall \alpha \in [0, 1], \forall F \in \mathcal{F}, \forall p \in \Omega \\ im_{\alpha}(F)(p) = \max\{g \in \mathcal{G} / F(p, g) \geq \alpha\}$$

Cet opérateur sert à revenir dans le domaine des images à niveaux de gris à partir d'une image d'ombres floues. Comme cette dernière peut présenter une imprécision sur la valeur réelle de niveau de gris pour tous les points de l'espace, il faut lever l'ambiguïté. Cela est fait en considérant l' α -coupe de niveau α pour obtenir une image d'ombres binaire, et ainsi pouvoir revenir de manière directe vers une image à niveaux de gris.

2.2.4 Images d'intervalles flous et images de nombres flous

Nous proposons ici une autre classe d'images floues. Contrairement à la définition précédente des images d'ombres floues où l'imprécision était traduite par une inégalité ($F(p, g)$ correspond au degré avec lequel l'image est supérieure ou égale à g en p), nous allons essayer de considérer une imprécision sur l'égalité en un point donné à un niveau de gris donné. Ainsi on souhaite une image floue $F \in \mathcal{F}$ qui ait la sémantique suivante pour un point p et un niveau de gris g : $F(p, g)$ correspond au degré avec lequel l'image est égale à g au point p . De manière plus pratique, on représentera les niveaux de gris soit par des intervalles flous, soit par des nombres flous.

Définition 2.2.4. Une image $F \in \mathcal{F}$ est une image d'intervalles flous (IIF) ssi :

$$\forall p \in \Omega \quad F(p, *) \text{ est un intervalle flou}$$

De manière similaire, on définit les images de nombres flous (INF) :

Définition 2.2.5. Une image $F \in \mathcal{F}$ est une image de nombres flous ssi :

$$\forall p \in \Omega \quad F(p, *) \text{ est un nombre flou}$$

Ces deux types d'images ont pour but de permettre la formalisation d'extensions floues de filtres connexes travaillant sur des composantes connexes extraites à partir des zones plates de l'image.

Une manière simple de construire de telles images à partir d'une image nette $I \in \mathcal{I}$ est de déployer un patron centré autour de chaque valeur nette de niveau de gris. Les limitations de cette méthode seront discutées dans le chapitre 4, et d'autres méthodes plus évoluées seront proposées. Des exemples d'images d'intervalles et de nombres flous sont présentés à la figure 2.3.

Comme pour les images d'ombres floues, l'imprécision représentée par les degrés d'appartenance des IIF dépend fortement des défauts modélisés lors de l'étape de construction de ces dernières.

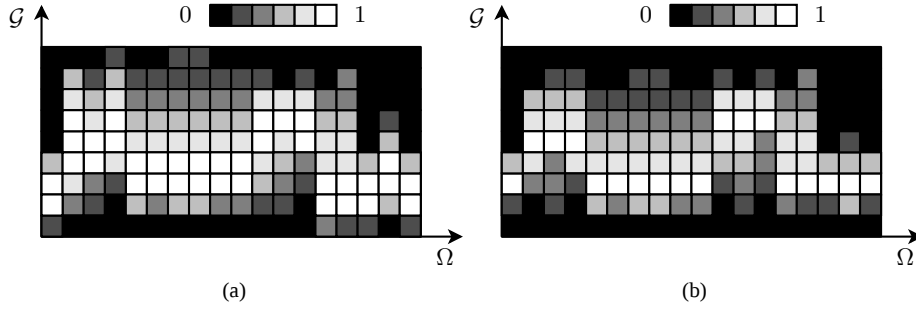


FIG. 2.3 – Exemple d'une image d'intervalles flous (a) et d'une image de nombres flous (b) définies sur un domaine Ω 1D.

La différence principale entre les deux types de quantités floues utilisées dans les images que l'on vient d'introduire se situe au niveau du noyau. Dans le cas d'intervalles flous, ce dernier est un intervalle alors que dans le cas de nombres flous, il est réduit à un point. Ainsi on remarque que les INF sont juste un cas particulier des IIF.

2.2.5 Autres utilisations du flou en traitement d'image

La modélisation de l'imprécision que nous venons de proposer n'est bien évidemment pas la seule possible. Ainsi on peut trouver dans la littérature des méthodes utilisant la théorie des ensembles flous pour résoudre différents problèmes liés à l'image (Nachtegael *et al.*, 2007) comme le débruitage (Schulte *et al.*, 2006b,a) ou la segmentation (Couto *et al.*, 2008). Récemment des travaux proposant d'utiliser des ensembles flous de type 2 ont été introduits (Tizhoosh, 2008; Tehami *et al.*, 2007; Bustince *et al.*, 2007). Ce dernier type d'ensembles permet de modéliser l'imprécision qui peut apparaître sur les degrés d'appartenance que l'on associe aux niveaux de gris d'une image. Malheureusement, ce type d'approches, bien que potentiellement plus puissantes que des approches reposant uniquement sur des ensembles flous classiques (ou de type 1), est aussi plus gourmand en temps de calcul. Leur utilisation est donc contrainte à des simplifications qui sont propres à des applications spécifiques.

Contrairement aux approches existantes, le formalisme que nous proposons dans ce chapitre peut être vu comme un cadre de travail générique permettant de représenter l'imprécision dans les images. Ainsi, on ne se limite pas uniquement à une tâche spécifique comme la segmentation par exemple même si on utilisera cette dernière pour illustrer le potentiel de nos développements.

2.3 Les filtres connexes flous

Dans cette section, nous allons nous intéresser à la définition de filtres connexes sur le nouveau type d'images que nous venons d'introduire. Pour ce faire, nous rappellerons certaines notions de connexité floue qui existent dans la littérature, nous nous attarderons ensuite sur la définition de zones plates floues. Puis nous introduirons la notion d'opérateurs connexes pour les ensembles flous, de filtres d'extinction et d'opérateurs connexes flous pour les images à niveaux de gris. Nous étudierons aussi les propriétés mathématiques de ces nouveaux opérateurs ainsi que les liens existant avec les filtres plus classiques. Nous finirons par une définition plus générale de filtres dans le cadre de la détection de structures.

2.3.1 Connexité des ensembles flous

La notion de connexité des ensembles flous a fait l'objet de différentes définitions comme celle proposée par Rosenfeld (1979) ou encore celle de Nempont *et al.* (2008), la plus répandue étant la première. Ces notions de connexité peuvent être représentées et manipulées de façon appropriée par une hyperconnexion (Serra, 1998; Braga-Neto et Goutsias, 2003b). Nous notons $\mathcal{H} \subseteq \mathcal{F}$ l'ensemble des ensembles flous connexes au sens d'une de ces définitions.

Connexité de Rosenfeld (1979)

La connexité floue la plus simple a été introduite par Rosenfeld (1979). L'idée est d'associer un degré de connexité entre deux points dans un ensemble $f \in \mathcal{S}$ en utilisant la définition suivante :

Définition 2.3.1. Le degré de connexité dans un ensemble flou $f \in \mathcal{S}$ entre deux points p et p' de Ω^2 est (Rosenfeld, 1984, 1992, 1998) :

$$c_f(p, p') = \max_{L \in \{\text{chemin}_{p,p'}\}} \left(\min_{p_i \in L} f(p_i) \right)$$

avec $\{\text{chemin}_{p,p'}\}$ l'ensemble des chemins de p vers p' en utilisant la connexité discrète sur Ω .

En utilisant cette définition, on peut introduire la notion de composante connexe associée à un point en utilisant la définition suivante :

Définition 2.3.2. La composante connexe associée à un point $p \in \Omega$ est exprimée comme :

$$\forall p' \in \Omega \quad \tilde{\Gamma}_f^p(p') = c_f(p, p')$$

La définition 2.3.1 implique que le degré de connexité entre deux points p et p' d'un ensemble flou connexe f est égal à $\min(f(p), f(p'))$. D'un point de vue topologique, si on considère un ensemble flou connexe comme un paysage, la définition 2.3.2 implique qu'en partant de son unique sommet (qui peut être un plateau), on ne peut que descendre ou stagner. On définira de manière plus générale un ensemble connexe de la manière suivante :

Définition 2.3.3. Un ensemble $f \in \mathcal{S}$ est un ensemble connexe flou ssi toutes ses α -coupes sont connexes au sens de la connexité discrète c_d (Rosenfeld, 1979).

L'idée derrière cette définition est que l'on considère qu'un ensemble flou est connexe si chacune de ses α -coupes contient au plus une composante connexe au sens binaire (c.f. figure 2.4). On notera $\mathcal{H}^1$ l'ensemble des composantes connexes de $\mathcal{S}$ selon cette définition.

Hyperconnexion de Braga-Neto et Goutsias (2003b)

Bien que très intuitive, la définition précédente a pour défaut majeur d'être assez stricte. En effet, pour qu'un ensemble flou soit considéré comme connexe, il faut que toutes ses α -coupes le soient. Cela peut poser problème lorsque des petites variations dans l'ensemble flou invalident cette définition. De plus, étant une mesure en tout ou rien, il se pose aussi un problème de continuité. Pour contourner ces problèmes, on peut utiliser les travaux de Braga-Neto et Goutsias (2003b) dans le but d'assouplir la définition précédente. Pour se faire, on définit une hyperconnexion $\mathcal{H}_\tau^1$, qui représente l'ensemble des ensembles flous connexes modulo le paramètre τ qui correspond à la limite de ce que l'on considère comme connexe :

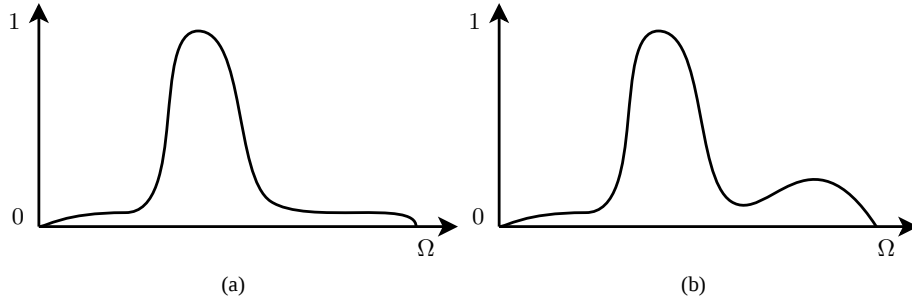


FIG. 2.4 – Illustration de $\mathcal{H}^1$. (a) Ensemble flou connexe selon la définition 2.3.3. (b) Ensemble non connexe.

Définition 2.3.4. Pour tout $\tau \in [0, 1]$, l'hyperconnexion $\mathcal{H}_\tau^1$ s'exprime de la manière suivante :

$$\mathcal{H}_\tau^1 = \{f \in \mathcal{S} / \forall \alpha \leq \tau, f_\alpha \text{ est connexe}\}$$

En d'autres termes on s'autorise à avoir plusieurs composantes connexes dans les α -coupes pour des $\alpha > \tau$ comme on peut le voir à la figure 2.5. Toujours dans l'optique d'assouplir la notion de connexité, on peut définir un degré de connexité à partir de cette hyperconnexion.

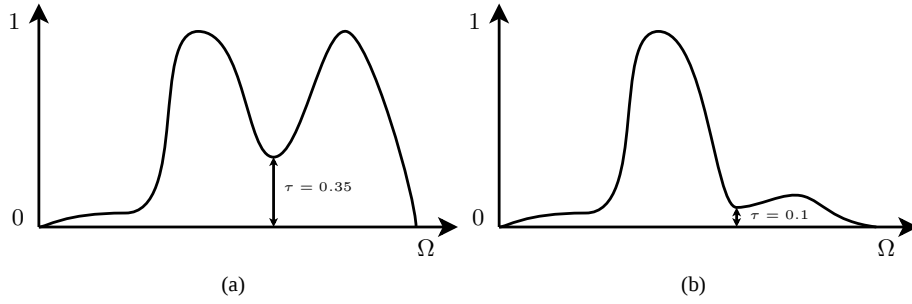


FIG. 2.5 – Illustration de $\mathcal{H}_\tau^1$. (a) Ensemble appartenant à $\mathcal{H}_{0.35}^1$. (b) Ensemble appartenant à $\mathcal{H}_{0.1}^1$. Le degré de connexité du premier ensemble est supérieur à celui du deuxième bien que visuellement, ce dernier semble le plus connexe.

Définition 2.3.5. Le degré de connexité d'un ensemble flou en fonction des hyperconnexions $\{\mathcal{H}_\tau^1, \tau \in [0, 1]\}$ s'exprime comme :

$$\forall f \in \mathcal{S} \quad c^1(f) = \max \{ \tau \in [0, 1] / f \in \mathcal{H}_\tau^1 \}$$

Hyperconnexion proposée par Nempont *et al.* (2008)

Le problème avec l'hyperconnexion précédente est que des faibles variations dans la fonction d'appartenance d'un ensemble flou ne sont pas traitées de la même manière si elles apparaissent pour des niveaux d' α -coupes élevés ou pour des niveaux d' α -coupes faibles. Pour cette raison, Nempont *et al.* (2008) ont introduit une nouvelle hyperconnexion reposant sur un nouveau degré de connexité d'un ensemble flou. Ce dernier repose sur le degré de connexité de deux points dans un ensemble flou.

Définition 2.3.6. Le degré de connexité c_f^2 entre deux points p, p' de Ω dans un ensemble flou $f \in \mathcal{S}$ s'exprime comme :

$$c_f^2(p, p') = 1 - \min(f(p), f(p')) + c_f(p, p')$$

Ce degré de connexité s'obtient en faisant le constat que pour tout ensemble flou f , pour tout couple de points $(p, p') \in \Omega$, on a toujours $c_f(p, p') \leq \min(f(p), f(p'))$, ce qui permet de dire que l'expression $c_f(p, p') = \min(f(p), f(p'))$ que vérifie un ensemble flou au sens de la définition 2.3.3 est équivalente à $c_f(p, p') \geq \min(f(p), f(p'))$. En substituant l'inégalité par une implication de Lukasiewicz (on écrit $a \leq b$ sous la forme $\min(1, b + 1 - a)$ pour $b = c_f(p, p')$ et $a = \min(f(p), f(p'))$), on obtient $c_f^2(p, p')$ comme degré de satisfaction pour cette inégalité (Nempont *et al.*, 2008).

Ce degré de connexité entre deux points permet d'associer un degré de connexité à un ensemble flou :

Définition 2.3.7. Le degré de connexité c^2 d'un ensemble flou $f \in \mathcal{S}$ s'exprime comme :

$$c^2(f) = \min_{p, p' \in \Omega} c_f^2(p, p')$$

En utilisant ce degré de connexité, on peut définir une nouvelle hyperconnexion dont les ensembles qui la composent doivent avoir un degré de connexité supérieur ou égal à un seuil (c.f. figure 2.6).

Définition 2.3.8. Pour tout $\tau \in [0, 1]$, l'hyperconnexion $\mathcal{H}_\tau^2$ est définie comme :

$$\mathcal{H}_\tau^2 = \{f \in \mathcal{S} / c^2(f) \geq \tau\}$$

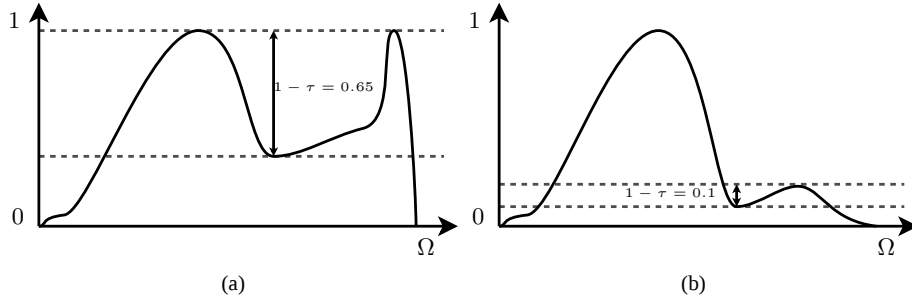


FIG. 2.6 – Illustration de $\mathcal{H}_\tau^2$. (a) Ensemble appartenant à $\mathcal{H}_{0.35}^2$. (b) Ensemble appartenant à $\mathcal{H}_{0.9}^1$.

Composante connexe floues

Le but de ce chapitre étant de fournir une extension dans le cadre du flou des filtres connexes, il nous faut établir la notion de composante connexe floue. De manière générique, on considérera à présent que l'on utilise une hyperconnexion $\mathcal{H}$ du même genre que celles décrites précédemment ($\mathcal{H}^1$, $\mathcal{H}_\tau^1$ ou $\mathcal{H}_\tau^2$) sans s'y limiter. On peut donc déduire la définition suivante :

Définition 2.3.9. L'ensemble des composantes connexes floues incluses dans un ensemble flou $f \in \mathcal{S}$ notée $\mathcal{H}(f)$ est constitué des plus grands éléments de $\mathcal{H}$ plus petits que f .

La figure 2.7 illustre cette définition dans le cas où $\mathcal{H} = \mathcal{H}^1$. Selon l'hyperconnexion choisie, on étend la définition de $\tilde{\Gamma}_f^p$ de la manière suivante :

Définition 2.3.10. $\forall f \in \mathcal{S}, \forall p \in \Omega$ Γ_f^p est la composante connexe incluse dans $\mathcal{H}(f)$ qui maximise le degré d'appartenance au point p .

2.3.2 Régions homogènes et zones plates floues

Pour étendre la notion de zones plates dans le cadre du flou, nous considérerons les images de nombres flous ainsi que les images d'intervalles flous. Dans ces images, le niveau de gris en chaque point est modélisé par une quantité floue pouvant être considérée comme l'extension floue d'une valeur de niveau de gris classique. Si tous les points d'une région donnée peuvent présenter le même niveau de gris nous considérons

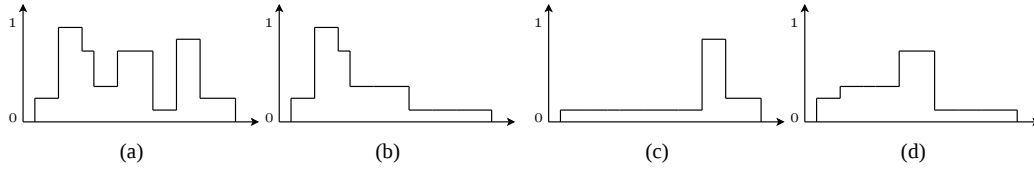
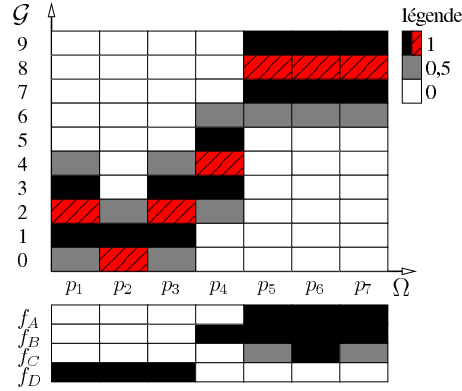


FIG. 2.7 – Composantes connexe floues (b), (c), (d) incluses dans l'ensemble (a).

alors qu'elle est stable radiométriquement. Les régions étant représentées par des ensembles flous, le sous-ensemble χ^F de $\mathcal{S}$ des ensembles flous radiométriquement stables dans une image I peut alors être défini relativement à l'image floue correspondante F :

$$\chi^F = \{f \in \mathcal{S} / \exists g \in \mathcal{G}, f \leq F(*, g)\}. \quad (2.1)$$

L'ensemble $\chi^F \cap \mathcal{H}$ est composé des ensembles flous connexes et radiométriquement stables. Nous définissons les zones plates floues comme les éléments de taille maximale de cet ensemble associés aux différents niveaux de gris $g : \{f \in \mathcal{H}(F(*, g)) / g \in \mathcal{G}\}$.

FIG. 2.8 – Composantes connexes incluses dans une image floue. Les données originales correspondent aux zones hachurées. f_A , f_C et f_D appartiennent à $\chi^F \cap \mathcal{H}$, mais pas f_B .

La figure 2.8 illustre cette définition sur un exemple 1D. Les données originales sont représentées en hachuré et nous choisissons $F(p, g) = 1$ si $|g - I(p)| \leq 1$, $F(p, g) = 0,5$ si $|g - I(x)| = 2$ et $F(p, g) = 0$ sinon. L'ensemble flou f_A appartient à $\chi^F \cap \mathcal{H}$. En effet il est connexe et vérifie $f_A \leq F(*, g)$ pour g valant 7, 8 ou 9. Par contre f_B n'appartient pas à ce dernier. En effet il n'existe aucun niveau g tel que la condition d'inclusion de l'équation 2.1 soit vérifiée. Les ensembles flous f_C et f_D sont deux autres exemples d'ensembles appartenant à $\chi^F \cap \mathcal{H}$. Notons de plus que f_A et f_D sont deux zones plates floues.

Contrairement aux zones plates classiques, l'expression floue de ces dernières ne forme plus nécessairement une partition de l'image I puisqu'elles peuvent se recouvrir. Par ailleurs si aucune imprécision n'est introduite lors de la construction de l'image floue F , c'est-à-dire que pour tout point p et pour tout niveau de gris g nous avons $F(p, g) = 1$ si $I(p) = g$ et $F(x, g) = 0$ sinon, les zones plates floues sont exactement les zones plates de l'image I et la définition proposée est donc cohérente avec le cas classique.

Remarquons que l'expression 2.1 peut aussi avoir du sens dans le cadre d'image d'ombres floues. La différence principale étant que la notion de stabilité radiométrique est remplacée par une notion de *compatibilité radiométrique* définie à partir d'une inégalité ($\geq$) sur les niveaux de gris dans l'image. Ainsi, des points appartenant à un ensemble flou de $\mathcal{S}$ inclus dans un $F(*, g)$ seraient radiométriquement compatibles.

2.3.3 Familles d'opérateurs connexes flous

Commençons par définir un opérateur connexe pour un ensemble flou de $\mathcal{S}$:

Définition 2.3.11. Un opérateur $\phi : \mathcal{S} \rightarrow \mathcal{S}$ est un opérateur connexe flou (CF) si :

$$\forall f \in \mathcal{S}, \forall \alpha \in [0, 1] \quad (\mathcal{C}(f_\alpha \cap \overline{\phi(f)_\alpha}) \subseteq \mathcal{C}(f_\alpha)) \wedge (\mathcal{C}(\overline{f_\alpha} \cap \phi(f)_\alpha) \subseteq \mathcal{C}(\overline{f_\alpha}))$$

En d'autres termes, pour chacune des α -coupes, on peut raisonner comme dans le cas binaire et considérer les composantes connexes de l' α -coupe. Le comportement pour chaque α -coupe est donc le même que dans le cas d'un opérateur connexe binaire classique (c.f. définition 2.1.4).

Soit $\Psi = \{\psi^{g,F} : \mathcal{S} \rightarrow \mathcal{S} / g \in \mathcal{G}, F \in \mathcal{F}\}$ un ensemble d'opérateurs.

Définition 2.3.12. Ψ est un ensemble d'opérateurs connexe flous faibles (EOCFF) si :

$$\forall g \in \mathcal{G}, \forall F \in \mathcal{F} \quad \psi^{g,F} \text{ est CF} \quad (2.2)$$

$$\forall g \in \mathcal{G}, \forall F \in \mathcal{F}, \forall f \in \mathcal{S} \quad \psi^{g,F}(f) \subseteq f \quad (2.3)$$

où $\subseteq$ correspond à l'inclusion classique sur $\mathcal{S}$.

Un opérateur flou connexe faible ($\psi^{g,F}$) est l'équivalent d'un opérateur connexe en binaire avec une propriété d'anti-extensivité. Les trois indices ont des buts assez distincts, et pourront (des exemples seront présentés à la fin de ce chapitre) dans certains cas être supprimés. Les paramètres g et F correspondent à des données extérieures qui seront utilisées dans le cadre de la définition des opérateurs connexes sur les images floues. On s'autorisera à définir grâce à cela des opérateurs qui dépendent du contenu d'une image (celle qui est filtrée par exemple) en plus de l'emboîtement des composantes. Cela correspond à l'extension de filtres tel que l'arasement volumique (Vachier, 2001).

L'équation 2.2 permet d'imposer que l'on travaille avec l'extension précédente d'opérateurs connexes flous. On s'assure ainsi de ne pas introduire ou déplacer des contours dans l'ensemble filtré (en raisonnant sur les α -coupes). L'équation 2.3 impose l'anti-extensivité de l'opérateur. Cette condition permettra d'assurer l'anti-extensivité de l'opérateur déduit de Ψ destiné aux images floues.

Notons enfin qu'aucune croissance par rapport à g n'est imposée pour l'instant. Cela se traduira, lorsque l'on filtrera des images d'ombres floues, par la non assurance de rester dans le domaine des images d'ombres floues telles que définies à l'équation 2.2.1. Cela sera détaillé après.

Définition 2.3.13. Ψ est un ensemble d'opérateurs connexe flous simples (EOCFS) si :

$$\Psi \text{ est un EOCFF} \quad (2.4)$$

$$\begin{aligned} \forall g \in \mathcal{G}, \forall (F, F') \in \mathcal{F}^2, \forall (f, h) \in \mathcal{S}^2 \\ (f \subseteq h) \wedge (F \subseteq F') \Rightarrow \psi^{g,F}(f) \subseteq \psi^{g,F'}(h) \end{aligned} \quad (2.5)$$

La propriété de décroissance (équation 2.5) permettra d'assurer la croissance d'un opérateur dédié aux images d'ombres construit à partir d'un Ψ EOCFS comme il sera montré plus tard.

Définition 2.3.14. Ψ est un ensemble d'opérateurs connexe flous ordonnés (EOCFO) si :

$$\Psi \text{ est un EOCFS} \quad (2.6)$$

$$\begin{aligned} \forall (g_1, g_2) \in \mathcal{G}^2, \forall F \in \mathcal{F}, \forall f \in \mathcal{S} \\ g_1 \geq g_2 \Rightarrow \psi^{g_1,F}(f) \subseteq \psi^{g_2,F}(f) \end{aligned} \quad (2.7)$$

On raffine ici la définition de Ψ par l'introduction de la propriété de croissance par rapport à g . Comme il le sera montré plus tard, cette condition est suffisante pour s'assurer que le filtrage d'une IOF donne bien une IOF.

Soit Ψ_λ un ensemble de fonctions paramétrées par $\lambda \in \mathbb{R}^+ : \Psi_\lambda = \{\psi_\lambda^{g,F} : \mathcal{S} \rightarrow \mathcal{S} / g \in \mathcal{G}, F \in \mathcal{F}\}$.

Définition 2.3.15. Ψ_λ est un ensemble d'opérateurs connexe flous étendus (EOCFE) si :

$$\forall \lambda \quad \Psi_\lambda \text{ est un EOCFO} \quad (2.8)$$

$$\begin{aligned} \forall (g_1, g_2) \in \mathcal{G}^2, \forall (\lambda, \mu) \in \mathbb{R}^{+2}, \forall F \in \mathcal{F}, \forall f \in \mathcal{S} \\ (\lambda \geq \mu) \wedge (g_1 \geq g_2) \Rightarrow \psi_\lambda^{g_1,F}(f) \subseteq \psi_\mu^{g_2,F}(f) \end{aligned} \quad (2.9)$$

$$\forall g \in \mathcal{G}, \forall F \in \mathcal{F}, \forall f \in \mathcal{S} \quad \psi_0^{g,F}(f) = f \quad (2.10)$$

On raffine ici la définition de Ψ en ajoutant un paramètre λ avec une contrainte de décroissance par rapport à ce dernier. L'idée est ici de proposer les briques pour une extension des filtres d'extinction. En effet, la sémantique de λ est que ce paramètre représente un seuil sur un attribut (aire, volume, etc.) calculé sur les ensembles flous. En jouant avec ce seuil on pourrait apprécier la disparition graduelle d'une composante connexe floue. Remarquons enfin que l'équation 2.9 entraîne de manière directe l'équation 2.7 pour ψ_λ .

2.3.4 Persistance d'un ensemble flou

Un opérateur connexe classique simplifie l'ensemble qui lui est passé en entrée en supprimant certaines composantes connexes. Si cet opérateur est paramétré par un paramètre scalaire, il peut être intéressant de voir pour quelles valeurs critiques de ce dernier les différentes composantes connexes de l'ensemble disparaissent. Cela est généralement réalisé à l'aide d'opérateurs d'extinction (c.f. annexe A.4.2 pour plus de détails sur ce type d'opérateurs).

Comme c'est le cas pour les opérateurs connexes classiques, les opérateurs flous ont pour vocation de faire disparaître certaines composantes connexes de l'ensemble d'entrée. La différence principale est que ces composantes connexes peuvent disparaître partiellement. Dans le même esprit que les opérateurs d'extinction, nous proposons un outil pour apprécier la conservation des différentes composantes connexes floues en fonction de la force du filtrage (par l'intermédiaire d'un paramètre scalaire). Nous introduisons cette notion sous le nom de persistance floue. L'idée est de construire pour une composante connexe floue une quantité floue qui associe à chaque valeur de paramètre le degré de similitude entre la composante avant et après filtrage. Nous allons tout d'abord rappeler la définition d'une mesure de similitude nécessaire pour définir la persistance.

Définition 2.3.16. $S : \mathcal{S} \times \mathcal{S} \rightarrow [0, 1]$ est une mesure de similarité si (Dubois et Prade, 1980) :

$$\forall f \in \mathcal{S} \quad S(f, f) = 1 \quad (2.11)$$

$$\forall (f, h) \in \mathcal{S}^2 \quad S(f, h) = S(h, f) \quad (2.12)$$

$$\forall (f, h, j) \in \mathcal{S}^3 \quad \min(S(f, h), S(h, j)) \leq S(f, j) \quad (2.13)$$

Ces propriétés signifient que la mesure S doit être réflexive (équation 2.11), symétrique (équation 2.12) et min-max transitive (équation 2.13).

En considérant un Ψ_λ EOCFE (λ étant le scalaire permettant d'influer sur la force du filtrage), on peut introduire la notion de persistance floue d'une composante connexe associée à un point de Ω dans une image IOF :

$$pers : \mathcal{F} \times E \times \Omega \rightarrow \mathcal{K}$$

qui est définie comme :

$$\forall F \in \mathcal{F}, \forall g \in \mathcal{G}, \forall p, \forall \lambda \in \mathbb{R}^+ \in \Omega \quad pers(F, g, p)(\lambda) = S\left(\Gamma_{F(*,g)}^p, \Gamma_{\psi_\lambda^{g,F}(F(*,g))}^p\right)$$

Pour une image F et un niveau de gris g , $pers(F, g, p)(\lambda)$ correspond au degré avec lequel l'objet représenté par la composante connexe floue associée à p persiste au filtrage par ψ_λ (c.f. figure 2.9). La quantité floue résultante $pers(F, g, p)$ traduit donc la disparition en fonction de λ de cette composante. Cette formulation a donc pour but de traduire le même concept que les fonctions classiques d'extinction (Matheron, 1974; Serra, 1982).

2.3.5 Opérateurs connexes flous pour les images à niveaux de gris flous

Nous allons maintenant utiliser les définitions introduites précédemment pour définir un moyen de filtrer les images à niveaux de gris.

Définition 2.3.17. Soit $\Psi = \{\psi^{g,F} : \mathcal{S} \rightarrow \mathcal{S} / g \in \mathcal{G}, F \in \mathcal{F}\}$ un EOCCF, l'opérateur $\delta_\Psi : \mathcal{F} \rightarrow \mathcal{F}$ associé à Ψ dédié aux images floues est défini de la manière suivante :

$$\forall F \in \mathcal{F}, \forall p \in \Omega, \forall g \in \mathcal{G} \quad \delta_\Psi(F)(p, g) = \psi^{g,F}(F(*, g))(p)$$

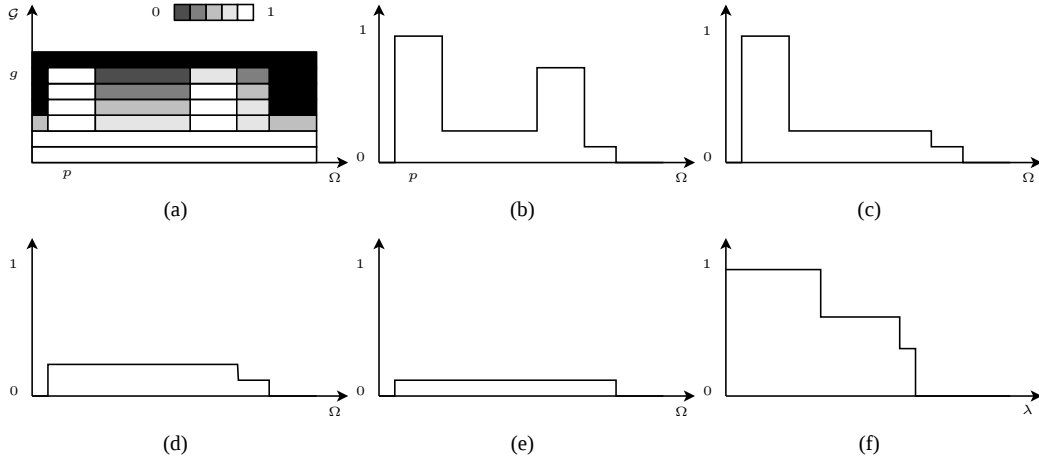


FIG. 2.9 – Persistance (f) d’une composante connexe floue $\Gamma_{F(*,g)}^p$ (c) à partir de l’ensemble $F(*,g)$ (b) extrait de l’image d’ombres floues (a). (d) et (e) représentent $\Gamma_{\psi_{\lambda}^{g,F}(F(*,g))}^p$ pour des valeurs croissantes de λ .

De par sa construction, on remarque qu’un tel filtre va traiter une image en entrée en utilisant les sous-ensembles flous contenus dans chaque niveau de gris de manière pseudo indépendante (grâce aux indices g et F , on peut toujours avoir accès à l’information globale de l’image).

Grâce aux propriétés que peut avoir l’ensemble d’opérateurs Ψ , on va déduire les propriétés que peut posséder l’opérateur δ_{Ψ} . Cette étape d’étude de propriétés est intéressante pour pouvoir facilement, et de manière abstraite, connaître les possibilités d’un filtre et son utilisation potentielle. Il serait aussi intéressant de définir un autre type d’opérateur de la forme $\psi^{g,F'}(F(*,g))$ avec F' une image différente de l’image (F) que l’on cherche à filtrer.

On cherche ici à définir les conditions pour qu’un filtre associé à un ensemble d’opérateurs Ψ retourne une IOF lorsqu’il est appliqué à une IOF, en d’autres termes que l’on travaille toujours avec des images d’ombres floues.

Théorème 2.3.1. *Un filtre associé à un Ψ EOCFO est une opération interne dans les IOF :*

$$\forall \Psi \text{ EOCFO}, \forall F \in \mathcal{F} \quad F \text{ IOF} \Rightarrow \delta_{\Psi}(F) \text{ IOF}$$

Démonstration. Soit Ψ un EOCFO, soit F une IOF. En utilisant la définition 2.2.1, on a :

$$\begin{aligned} \forall p \in \Omega, \forall (g_1, g_2) \in \mathcal{G}^2 \quad g_1 \leq g_2 &\Rightarrow F(p, g_1) \geq F(p, g_2) \\ &\Rightarrow F(*, g_2) \subseteq F(*, g_1) \end{aligned}$$

de plus on a :

$$\begin{aligned} \delta_{\Psi}(F)(p, g_1) &= \psi^{g_1, F}(F(*, g_1))(p) \\ \delta_{\Psi}(F)(p, g_2) &= \psi^{g_2, F}(F(*, g_2))(p) \end{aligned}$$

en utilisant l’équation 2.5 on a donc :

$$\psi^{g_1, F}(F(*, g_2)) \subseteq \psi^{g_1, F}(F(*, g_1))$$

or en utilisant l’équation 2.7, on a :

$$\psi^{g_2, F}(F(*, g_2)) \subseteq \psi^{g_1, F}(F(*, g_2))$$

d’où :

$$\psi^{g_2, F}(F(*, g_2)) \subseteq \psi^{g_1, F}(F(*, g_1))$$

donc :

$$\delta_{\Psi}(F)(p, g_2) \leq \delta_{\Psi}(F)(p, g_1)$$

ce qui implique que $\delta_{\Psi}(F)$ est une *IOF*. □

Remarquons néanmoins qu'il peut être tout de même intéressant de travailler avec des ensembles d'opérateurs qui ne sont pas des *EOCFO* (et donc obtenir des images qui ne sont pas des *IOF*) : cela peut permettre de travailler seulement sur les composantes emboîtées d'une image qui vérifient un certain critère (e.g. compacité). Il reste dans ce cas à trouver un moyen autre que l'opérateur im_{α} pour défuzzifier et ainsi interpréter le résultat du filtrage. Nous aborderons cela à la définition 2.3.18.

Théorème 2.3.2. *Un filtre δ produit à partir d'un Ψ EOCFS est croissant :*

$$\forall \Psi \text{ EOCFS}, \forall (F, F') \in \mathcal{F}^2 \quad F \subseteq F' \Rightarrow \delta_{\Psi}(F) \subseteq \delta_{\Psi}(F')$$

Démonstration. Soit Ψ un EOCFS, soit $(F, F') \in \mathcal{F}^2$ tel que $F \subseteq F'$

$$\begin{aligned} \forall g \in \mathcal{G} \quad F(*, g) \subseteq F'(*, g) &\Rightarrow \psi^{g, F}(F(*, g)) \subseteq \psi^{g, F'}(F'(*, g)) \quad (\text{c.f. équation 2.5}) \\ &\Rightarrow \forall p \in \Omega \quad \psi^{g, F}(F(*, g))(p) \leq \psi^{g, F'}(F'(*, g))(p) \\ &\Rightarrow \forall p \in \Omega \quad \delta_{\Psi}(F)(p, g) \leq \delta_{\Psi}(F')(p, g) \quad (\text{c.f. définition 2.3.17}) \\ &\Rightarrow \delta_{\Psi}(F) \subseteq \delta_{\Psi}(F') \end{aligned}$$

□

Théorème 2.3.3. *Un filtre δ construit à partir d'un Ψ EOCFF vérifiant $\psi^{g, F}(f) = \psi^{g, F'}(\psi^{g, F}(f))$ est idempotent :*

$$\left\{ \begin{array}{l} \forall \Psi \text{ EOCFF tq } \forall (F, F') \in \mathcal{F}^2, \forall g \in \mathcal{G}, \forall f \in \mathcal{S} \\ \psi^{g, F}(f) = \psi^{g, F'}(\psi^{g, F}(f)) \Rightarrow \delta_{\Psi} \text{ idempotent} \end{array} \right. \quad (2.14)$$

Démonstration. Soit Ψ un EOCFF tel que $\forall (F, F') \in \mathcal{F}^2, \forall g \in \mathcal{G}, \psi^{g, F}(f) = \psi^{g, F'}(\psi^{g, F}(f))$.
Soit $J \in \mathcal{F}$

$$\begin{aligned} \forall p \in \Omega, \forall g \in \mathcal{G} \quad \delta_{\Psi}(\delta_{\Psi}(J))(p, g) &= \psi^{g, \delta_{\Psi}(J)}(\delta_{\Psi}(J)(* , g))(p) \\ &= \psi^{g, \delta_{\Psi}(J)}(\psi^{g, J}(J(*, g)))(p) \\ &= \psi^{g, J}(J(*, g))(p) \\ &= \delta_{\Psi}(J)(p, g) \\ &\Rightarrow \delta_{\Psi} \text{ idempotent} \end{aligned}$$

□

On remarquera que l'hypothèse faite sur l'absence de lien entre F et F' est assez forte. En fait, les filtres qui sont idempotents sont généralement ceux qui ne dépendent ni de g ni de F . Un exemple reposant sur le cardinal des composantes sera détaillé plus loin.

Théorème 2.3.4. *Un δ_{Ψ} construit à partir d'un Ψ EOCFS qui vérifie $\forall (F, F') \in \mathcal{F}^2, \forall g \in \mathcal{G}, \forall f \in \mathcal{S} \quad \psi^{g, F}(f) = \psi^{g, F'}(\psi^{g, F}(f))$ est un filtre morphologique.*

Démonstration. Comme δ_{Ψ} est croissant et idempotent, c'est bien un filtre morphologique. □

Théorème 2.3.5. *Un δ_{Ψ} construit à partir d'un Ψ EOCFF est anti-extensif.*

Démonstration.

$$\begin{aligned} \forall \Psi \text{ EOCFF}, \forall F \in \mathcal{F}, \forall g \in \mathcal{G} \\ \psi^{g, F}(F(*, g)) \subseteq F(*, g) &\Rightarrow \forall p \in \Omega \quad \psi^{g, F}(F(*, g))(p) \leq F(p, g) \quad (\text{c.f. équation 2.3}) \\ &\Rightarrow \delta_{\Psi}(F)(p, g) \leq F(p, g) \\ &\Rightarrow \delta_{\Psi} \text{ anti-extensif} \end{aligned}$$

□

L'anti-extensivité est donc transmise des ψ vers l'opérateur δ_Ψ . On en déduit donc le résultat suivant.

Théorème 2.3.6. *Un δ_Ψ construit à partir d'un Ψ EOCFS qui vérifie $\forall (F, F') \in \mathcal{F}^2, \forall g \in \mathcal{G}, \forall f \in \mathcal{S} \ \psi^{g, F'}(f) = \psi^{g, F'}(\psi^{g, F}(f))$ est une ouverture algébrique.*

Dans le but d'interpréter les images filtrées qui ne sont ni des *IOF* ni des *IIF* ou des *INF*, une étape d'agrégation peut être utilisée. Cet opérateur peut être utile lorsque l'on veut passer du domaine $(\Omega \times \mathcal{G})$ des images floues vers le domaine de définition des images classique (Ω) . Cela permet de s'abstraire des problèmes d'interprétation qui peuvent apparaître lorsque l'on travaille avec des filtres qui ont un comportement similaire à celui des amincissements (Breen et Jones, 1996) comme par exemple la création de contours artificiels.

Définition 2.3.18. L'opérateur $agg : \mathcal{F} \rightarrow \mathcal{S}$ est défini comme :

$$\forall F \in \mathcal{F}, p \in \Omega \quad agg(F)(p) = \bigwedge_{g \in \mathcal{G}} F(p, g)$$

avec $\bigwedge : [0, 1] \times [0, 1] \rightarrow [0, 1]$ une t-conorme quelconque.

Le choix d'une t-conorme se justifie par la sémantique qu'ont les filtres qui nécessitent l'emploi de l'opérateur agg . En général, ils ont pour but de détecter les composantes connexes qui vérifient un critère donné. Un opérateur de disjonction permet donc de voir les endroits dans Ω où au moins une composante connexe a été conservée, et ce indépendamment de son niveau de gris.

La figure 2.10 présente un exemple où l'opérateur agg peut être utilisé. Dans ce cas précis seules les composantes contenues dans une certaine fourchette de tailles sont gardées dans l'image filtrée, l'opérateur d'agrégation permet donc de voir où sont (dans Ω) les objets qui ont cette propriété.

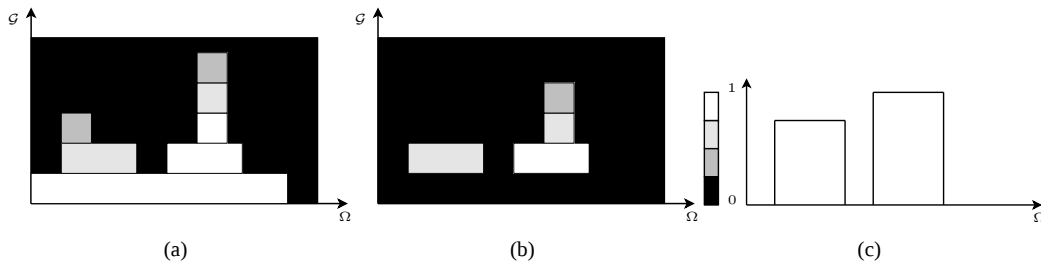


FIG. 2.10 – Agrégation (c) du résultat du filtrage (b) d'une image d'ombres floues (a).

2.3.6 Lien avec les opérateurs connexes classiques

Nous allons maintenant nous intéresser à la façon d'exprimer les filtres classiques (e.g. ouverture d'attribut (Vincent, 1993)) dans ce nouveau formalisme. Les développements proposés ont pour but de fournir un cadre générique pour l'étude d'opérateurs classiques dans le formalisme des images floues. Bien que cela se fasse au coût d'une certaine lourdeur d'écriture, cette étude permet d'établir un pont solide d'un point de vue théorique entre le formalisme flou proposé et le formalisme classique.

Définition 2.3.19. $\delta : \mathcal{F} \rightarrow \mathcal{F}$ est une extension floue (EF) d'un opérateur $G : \mathcal{I} \rightarrow \mathcal{I}$ si :

$$\forall I \in \mathcal{I} \quad im_1(\delta(um(I))) = G(I)$$

Cette définition d'extension floue peut se traduire comme le simple fait qu'un opérateur est une extension floue d'un filtre classique à la condition qu'il se comporte comme ce dernier sur le noyau (partie où les degrés d'appartenance sont égaux à un) des images floues passées en entrée.

Nous allons maintenant introduire un formalisme permettant d'étudier de manière générique les extensions de filtres classiques.

Soit $A : \Omega \times \mathcal{G} \times \mathcal{I} \times 2^\Omega \rightarrow \mathbb{R}^+$ vérifiant :

$$\forall g \in \mathcal{G}, \forall p \in \Omega, \forall I \in \mathcal{I}, \forall N \subseteq \Omega, \forall p' \in \Omega \quad p' \in \hat{\Gamma}_N^p \Rightarrow A_p^{g,I}(N) = A_p^{g,I}(N) \quad (2.15)$$

$$\forall (g_1, g_2) \in \mathcal{G}^2, \forall p \in \Omega, \forall I \in \mathcal{I}, \forall N \subseteq \Omega \quad g_1 \leq g_2 \Rightarrow A_p^{g_1,I}(N) \geq A_p^{g_2,I}(N) \quad (2.16)$$

$$\forall g \in \mathcal{G}, \forall p \in \Omega, \forall (I, I') \in \mathcal{I}^2, \forall N \subseteq \Omega, M \subseteq \Omega \quad (I \leq I') \wedge (N \subseteq M) \Rightarrow A_p^{g,I}(N) \leq A_p^{g,I'}(M) \quad (2.17)$$

A peut être interprété comme une mesure faite sur la composante connexe dans un ensemble binaire N ou son complémentaire contenant le point p en s'appuyant sur les données apportées par l'image I et le niveau de gris g .

Soit $A' : \Omega \times \mathcal{G} \times \mathcal{I} \rightarrow \mathbb{R}^+$ défini comme :

$$\forall g \in \mathcal{G}, \forall p \in \Omega, \forall I \in \mathcal{I} \quad A_p^{g,I}(I) = A_p^{g,I}(X_g^+(I)) \quad (2.18)$$

avec X_g^+ l'opérateur de seuillage au niveau g .

A' représente ici une version *allégée* de A . En effet on force un lien entre l'ensemble précédemment noté N , le niveau de gris g et l'image I (N est un seuillage de I au niveau g).

Finalement, l'opérateur $G : \mathbb{R}^+ \times \mathcal{I} \rightarrow \mathcal{I}$ est introduit et défini comme :

$$\forall \lambda \in \mathbb{R}^+, \forall I \in \mathcal{I}, \forall p \in \Omega \quad G_\lambda(I)(p) = \max\{g \in \mathcal{G} / A_p^{g,I}(I) \geq \lambda\}$$

De manière évidente, en utilisant l'équation 2.18, G peut être réécrit de la manière suivante :

$$G_\lambda(I)(p) = \max\{g \in \mathcal{G} / A_p^{g,I}(X_g^+(I)) \geq \lambda\}$$

G représente quant à lui un filtre connexe qui va réduire les maxima de l'image I tant qu'ils ne satisfont pas un certain critère induit par l'opérateur A . Il sera montré plus loin que cette forme permet d'exprimer des filtres tels que l'ouverture d'aire ou encore l'arasement volumique.

Théorème 2.3.7. Un opérateur δ_{Ψ_λ} construit à partir d'un $\Psi_\lambda = \{\psi_\lambda^{g,F} : \mathcal{S} \rightarrow \mathcal{S} / g \in \mathcal{G}, F \in \mathcal{F}, f \in \mathcal{S}\}$, $\lambda \in \mathbb{R}^+$ avec $\forall p \in \Omega$:

$$\psi_\lambda^{g,F}(f)(p) = \max\left\{\alpha \in [0; f(p)] / A_p^{g,im_1(F_\alpha)}(f_\alpha) \geq \lambda\right\}$$

est une extension floue de l'opérateur G_λ .

Démonstration. Soit $p \in \Omega, I \in \mathcal{I}$

$$\begin{aligned} im_1(\delta_\Psi^\lambda(um(I)))(p) &= \max\{g \in \mathcal{G} / \delta_\Psi^\lambda(um(I))(p, g) \geq 1\} \\ &= \max\{g \in \mathcal{G} / \delta_\Psi^\lambda(um(I))(p, g) = 1\} \\ &= \max\{g \in \mathcal{G} / \psi_\lambda^{g,um(I)}(um(I)(*, g))(p) = 1\} \\ &\text{or } um(I)(*, g) = X_g^+(I) \\ &= \max\{g \in \mathcal{G} / \psi_\lambda^{g,um(I)}(X_g^+(I))(p) = 1\} \\ &= \max\{g / \max\{\alpha \in [0; X_g^+(I)(p)] / A_p^{g,im_1(um(I)_\alpha)}(X_g^+(I)_\alpha) \geq \lambda\} = 1\} \\ &= \max\{g / \max\{\alpha \in [0, 1] / A_p^{g,im_1(um(I)_\alpha)}(X_g^+(I)) = \lambda\} \geq 1\} \text{ car } X_g^+(I) \text{ est net} \end{aligned}$$

or $\forall \alpha \in]0; 1], um(I)_\alpha = um(I)$ et $um(I)_0 = \Omega$:

$$\begin{aligned} &\Rightarrow um(I)_\alpha \subseteq um(I)_0 \\ &\Rightarrow im_1(um(I)_\alpha) \leq im_1(um(I)_0) \\ &\Rightarrow \forall p \in \Omega, \forall g \in \mathcal{G} \quad A_p^{g,im_1(um(I)_\alpha)}(X_g^+(I)) \leq A_p^{g,im_1(um(I)_0)}(X_g^+(I)) \\ &\quad \text{(c.f. équation 2.17)} \end{aligned}$$

on a donc (car $\forall \alpha \in]0; 1] \quad um(I)_\alpha = um(I)_1$, puisque $um(I)$ est binaire) :

$$\begin{aligned} im_1(\delta_\Psi^\lambda(um(I)))(p) &= \max\{g / A_p^{g,I}(X_g^+(I)) \geq \lambda\} \\ &= G_\lambda(I) \end{aligned} \quad (2.19)$$

□

Cet ensemble Ψ peut être interprété comme le filtrage de chacune des tranches de la décomposition en α -coupes de f et de F . Comme on peut le voir ce n'est pas F_α qui est utilisée directement mais $im_1(F_\alpha)$. Cela se traduit par la génération d'une image dans $\mathcal{I}$ à partir de l' α -coupe de F qui va être utilisée comme donnée extérieure (au même titre que le niveau de gris g) pour effectuer une mesure sur l'ensemble f . La quantité $\psi_\lambda^{g,F}(f)(p)$ représente donc le degré d'appartenance maximal tel que la mesure sur f_α soit supérieure ou égale à un seuil λ . La figure 2.11 présente un exemple simple expliquant comment peut fonctionner un $\psi_\lambda^{g,F}$ avec une fonction A ne calculant que l'aire des composante et un $\lambda = 3$.

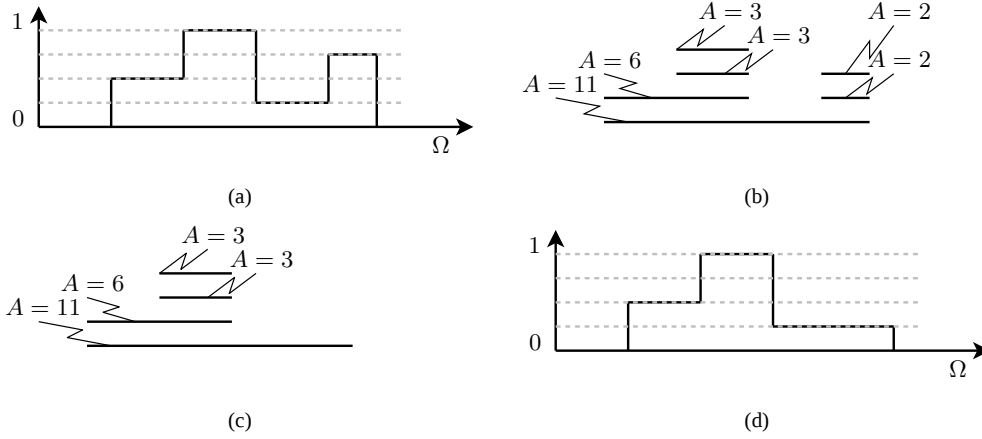


FIG. 2.11 – Filtrage d'un ensemble flou (a) par décomposition en α -coupes avec un ψ tel que défini à l'équation 2.36 qui est un cas particulier de l'équation 2.20. Une première séparation en α -coupes est faite (b), puis chacune d'entre elles est filtrée par une ouverture d'aire binaire ($\lambda = 3$) (c), et enfin le sous-ensemble résultant est reconstruit à partir des α -coupes filtrées (d).

Nous allons maintenant nous intéresser à la classification de la famille d'opérateurs permettant d'obtenir des extensions floues d'opérateurs comme proposé au théorème 2.3.7. Les résultats qui vont être exposés auront pour but de faire le lien avec les propriétés introduites dans la sous-section précédente concernant les opérateurs de filtrage d'images floues.

Théorème 2.3.8. Soit $A : \Omega \times \mathcal{G} \times \mathcal{I} \times \{\Omega\} \rightarrow \mathbb{R}^+$ vérifiant les équations 2.15, 2.16 et 2.17. Un ensemble Ψ^A d'opérateurs défini de la manière suivante :

$$\Psi^A = \left\{ \begin{array}{l} \psi_\lambda^{g,F} / \forall \lambda \in \mathbb{R}^+, \forall g \in \mathcal{G}, \forall F \in \mathcal{F}, \forall f \in \mathcal{S}, \forall p \in \Omega \\ \psi_\lambda^{g,F}(f)(p) = \max\{\alpha \in [0; f(p)] / A_p^{g, im_1(F_\alpha)}(f_\alpha) \geq \lambda \} \end{array} \right\} \quad (2.20)$$

est un EOCFE.

Démonstration. Montrons que Ψ^A est un EOCFE :

Montrons l'équation 2.2 ($\psi_\lambda^{g,F}$ est CF) :

$$\forall f \in \mathcal{S}, \forall F \in \mathcal{F}, \forall p \in \Omega, \forall \lambda \in \mathbb{R}^+, \forall g \in \mathcal{G}, \forall \beta \in [0, 1]$$

$$\psi_\lambda^{g,F}(f)_\beta(p) = \begin{cases} 1 & \text{si } A_p^{g, im_1(F_\beta)}(f_\beta) \geq \lambda \\ 0 & \text{sinon} \end{cases}$$

En utilisant l'équation 2.15 sur les α -coupes, on obtient :

$$\begin{aligned} \forall (p, p') \in \Omega^2 \quad p' \in \hat{\Gamma}_{f_\beta}^p &\Rightarrow A_p^{g, im_1(F_\beta)}(f_\beta) = A_{p'}^{g, im_1(F_\beta)}(f_\beta) \\ &\Rightarrow \hat{\Gamma}_{f_\beta}^p \in \mathcal{C} \left(\psi_\lambda^{g,F}(f)_\beta \right) \vee \hat{\Gamma}_{f_\beta}^p \in \mathcal{C} \left(\overline{\psi_\lambda^{g,F}(f)_\beta} \right) \end{aligned}$$

donc :

$$\text{l'équation 2.2 est vraie} \quad (2.21)$$

Montrons l'équation 2.3 (anti-extensivité) :

$$\begin{aligned} \forall f \in \mathcal{S}, \forall F \in \mathcal{F}, \forall p \in \Omega, \forall \lambda \in \mathbb{R}^+, \forall g \in \mathcal{G} \\ \psi_{\lambda}^{g,F}(f)(p) &= \max\{\alpha \in [0; f(p)] / A_p^{g,im_1(F_{\alpha})}(f_{\alpha}) \geq \lambda\} \\ &\leq f(p) \\ &\Rightarrow \text{équation 2.3 vraie} \end{aligned} \quad (2.22)$$

Montrons l'équation 2.5 (croissance) :

$$\forall (f, h) \in \mathcal{S}^2, \forall (F, F') \in \mathcal{F}^2, \forall \lambda \in \mathbb{R}^+, \forall g \in \mathcal{G}$$

on a :

$$F \subseteq F' \Rightarrow \forall p \in \Omega, \forall \alpha \in [0, 1] \quad im_1(F_{\alpha})(p) \leq im_1(F'_{\alpha})(p) \quad (2.23)$$

et

$$f \subseteq h \Rightarrow \forall \alpha \in [0, 1] \quad f_{\alpha} \subseteq h_{\alpha} \quad (2.24)$$

$$\begin{aligned} \psi_{\lambda}^{g,F}(f)(p) &= \max\{\alpha \in [0; f(p)] / A_p^{g,im_1(F_{\alpha})}(f_{\alpha}) \geq \lambda\} \\ \psi_{\lambda}^{g,F}(f)(p) &\leq \max\{\alpha \in [0; f(p)] / A_p^{g,im_1(F'_{\alpha})}(h_{\alpha}) \geq \lambda\} \text{ (c.f. équations 2.17, 2.23 et 2.24)} \\ &\leq \max\{\alpha \in [0; h(p)] / A_p^{g,im_1(F'_{\alpha})}(h_{\alpha}) \geq \lambda\} \\ &\leq \psi_{\lambda}^{g,F'}(h) \end{aligned}$$

$$\Rightarrow \text{équation 2.5 vraie} \quad (2.25)$$

On a donc bien :

$$\text{équations 2.21, 2.22 et 2.25} \Rightarrow \Psi^A \text{EOCFS} \quad (2.26)$$

Montrons enfin que Ψ^A est un EOCFE :

$$\begin{aligned} \forall (\lambda, \mu) \in \mathbb{R}^2, \forall (g_1, g_2) \in \mathcal{G}^2, \forall F \in \mathcal{F}, \forall f \in \mathcal{S}, \forall p \in \Omega \\ (\lambda \geq \mu) \wedge (g_1 \geq g_2) \Rightarrow \psi_{\lambda}^{g_1,F}(f)(p) &= \max\{\alpha \in [0; f(p)] / A_p^{g_1,im_1(F_{\alpha})}(f_{\alpha}) \geq \lambda\} \\ &\leq \max\{\alpha \in [0; f(p)] / A_p^{g_1,im_1(F_{\alpha})}(f_{\alpha}) \geq \mu\} \\ &\leq \max\{\alpha \in [0; f(p)] / A_p^{g_2,im_1(F_{\alpha})}(f_{\alpha}) \geq \mu\} \\ &\leq \psi_{\lambda}^{g_2,F}(f)(p) \\ &\Rightarrow \text{équation 2.9 vraie} \end{aligned} \quad (2.27)$$

Montrons l'équation 2.10 : Par définition, on a :

$$\forall p \in \Omega, \forall g \in \mathcal{G}, \forall F \in \mathcal{F}, \forall N \subseteq \Omega \quad A_p^{g,F}(N) \geq 0$$

donc

$$\begin{aligned} \forall p \in \Omega, \forall g \in \mathcal{G}, \forall F' \in \mathcal{F}, \forall f \in \mathcal{S} \\ \psi_0^{g,F'}(f) &= \max\{\alpha \in [0; f(p)] / A_p^{g,im_1(F_{\alpha})}(f_{\alpha}) \geq 0\} \\ &= f(p) \\ &\Rightarrow \text{équation 2.10 vraie} \end{aligned} \quad (2.28)$$

Comme on l'a déjà dit, l'équation 2.9 implique l'équation 2.7 donc Ψ^A est un EOCFO. En utilisant les équations 2.27 et 2.28, on peut donc conclure que Ψ^A est un EOCFE. $\square$

Cette dernière propriété est intéressante car en plus de définir une extension à une famille de filtres classiques en construisant un Ψ^A à partir d'un attribut A , on construit par la même occasion des ψ qui peuvent être utilisés pour le calculs de persistances floues sur des sous-ensembles flous dans $\mathcal{S}$.

Nous allons maintenant utiliser les théorèmes précédents pour introduire des extensions floues de deux filtres couramment utilisés, à savoir l'arasement volumique et l'ouverture d'aire. Ces résultats vont permettre de faire le lien entre le formalisme générique introduit et utilisé dans cette section avec des exemples pratiques, illustrant ainsi son intérêt.

On défini $\text{Vol} : \Omega \times \mathcal{G} \times \mathcal{I} \times 2^\Omega \rightarrow \mathbb{R}^+$ comme :

$$\forall p \in \Omega, \forall g \in \mathcal{G}, \forall I : \mathcal{I}, \forall N \subseteq \Omega \quad \text{Vol}_p^{g,I}(N) = \sum_{p' \in \hat{\Gamma}_N^p} \max(0, (I(p') - g))$$

Cet attribut correspond au calcul du volume dans $\max(0, I - g)$ sur un support fourni par la composante connexe dans N contenant p (c.f. figure 2.12).

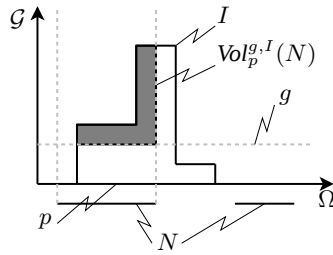


FIG. 2.12 – Exemple de calcul de $\text{Vol}_p^{g,I}(N)$.

Définissons aussi $\text{Vol}' : \Omega \times E \times \mathcal{I} \rightarrow \mathbb{R}^+$ comme :

$$\begin{aligned} \text{Vol}' : \Omega \times E \times \mathcal{I} &\rightarrow \mathbb{R}^+ / \\ \forall g \in E, \forall p \in \Omega, \forall I \in \mathcal{I} \quad \text{Vol}'_p^g(I) &= \text{Vol}_p^{g,I}(X_g^+(I)) \end{aligned} \quad (2.29)$$

Cette définition est une restriction de la précédente. En effet on impose un lien fort entre N et I . Ainsi l'exemple de la figure 2.12 n'est plus valide : les composantes connexes de N doivent correspondre à un seuillage de I . Un exemple de calcul de Vol' est présenté à la figure 2.13.

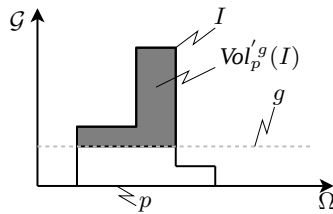


FIG. 2.13 – Exemple de calcul de $\text{Vol}'_p^g(I)$.

Théorème 2.3.9. $\delta_{\Psi^{\text{Vol}}}$ est une extension floue de l'arasement volumique $G_\lambda^{\text{Vol}} : \mathcal{I} \rightarrow \mathcal{I}$, $\lambda \in \mathbb{R}^+$ défini comme :

$$\forall I \in \mathcal{I}, \forall p \in \Omega \quad G_\lambda^{\text{Vol}}(I)(p) = \max\{g \in \mathcal{G} / \text{Vol}'_p^g(I) \geq \lambda\}$$

Démonstration. Vérifions les propriétés de l'opérateur Vol pour montrer qu'il répond bien aux critères définis aux équations 2.15, 2.16 et 2.17.

De manière évidente (somme sur $\hat{\Gamma}$), on a :

$$\text{équation 2.15 vraie} \quad (2.30)$$

$$\begin{aligned} \forall (g_1, g_2) \in \mathcal{G}^2, \forall I \in \mathcal{I}, \forall p \in \Omega, \forall f \subseteq \Omega \\ g_1 \leq g_2 \Rightarrow \sum_{p' \in \hat{\Gamma}_f^p} (\max(0, I(p') - g_1)) &\geq \sum_{p' \in \hat{\Gamma}_f^p} (\max(0, I(p') - g_2)) \\ \text{car } \forall p' \in \Omega \quad I(p') - g_1 &\geq I(p') - g_2 \end{aligned}$$

donc :

$$\text{équation 2.16 vraie} \quad (2.31)$$

$$\begin{aligned} \forall (I, I') \in \mathcal{I}^2, \forall N \subseteq \Omega, \forall M \subseteq \Omega \\ I \leq I' \Rightarrow \forall g \in \mathcal{G}, \forall p \in \Omega \quad \max(0, I(p) - g) &\leq \max(0, I'(p) - g) \\ N \subseteq M \Rightarrow \forall p \in \Omega \quad \hat{\Gamma}_N^p &\subseteq \hat{\Gamma}_M^p \end{aligned}$$

on a donc :

$$\forall p \in \Omega \quad \sum_{p' \in \hat{\Gamma}_N^p} (\max(0, I(p') - g)) \leq \sum_{p' \in \hat{\Gamma}_M^p} (\max(0, I'(p') - g))$$

donc :

$$\text{équation 2.17 vraie} \quad (2.32)$$

On a donc un opérateur de type “A”.

Soit :

$$\Psi^{\text{Vol}} = \left\{ \begin{array}{l} \psi_{\lambda}^{g, F} / \forall \lambda \in \mathbb{R}^+, \forall g \in \mathcal{G}, \forall F \in \mathcal{F}, \forall f \in \mathcal{S} \\ \psi_{\lambda}^{g, F}(f) = \max\{\alpha \in [0; f(p)] / \text{Vol}_p^{g, \text{im}_1(F_{\alpha})}(f_{\alpha}) \geq \lambda\} \end{array} \right\}$$

équations 2.30, 2.31, 2.32, théorème 2.3.7 et définition 2.3.19 $\Rightarrow \forall \lambda \in \mathbb{R}^+ \quad \delta_{\Psi^{\text{Vol}}}^{\lambda}$ est une EF de G_{λ}^{Vol}

De plus :

$$\text{équations 2.30, 2.31, 2.32 et théorème 2.3.8} \Rightarrow \Psi^{\text{Vol}} \text{ est un EOCFE}$$

□

La formalisme jusque là proposé nous a permis de définir relativement simplement une extension possible d'un filtre évolué comme l'arasement volumique. En utilisant la même approche, un résultat similaire peut être obtenu pour un filtre d'expression plus simple comme l'ouverture d'aire.

Définissons $\text{Card} : \Omega \times \mathcal{G} \times \mathcal{I} \times 2^{\Omega} \rightarrow \mathbb{R}^+$ comme suit :

$$\forall p \in \Omega, \forall g \in \mathcal{G}, \forall I \in \mathcal{I}, \forall N \subseteq \Omega \quad \text{Card}_p^{g, I} = \sum_{y \in \hat{\Gamma}_N^p} 1$$

Théorème 2.3.10. $\delta_{\Psi^{\text{Card}}}^{\lambda}$ est une extension floue de l'ouverture d'aire $G_{\lambda}^{\text{Card}} : \mathcal{I} \rightarrow \mathcal{I}, \lambda \in \mathbb{R}^+$ définie comme :

$$\forall I \in \mathcal{I}, \forall p \in \Omega \quad G_{\lambda}^{\text{Card}}(I)(p) = \max\{g \in \mathcal{G} / \text{Card}_p^{g, I}(I) \geq \lambda\}$$

avec :

$$\begin{aligned} \text{Card}' : \Omega \times \mathcal{G} \times \mathcal{I} &\rightarrow \mathbb{R}^+ / \forall g \in \mathcal{G}, \forall p \in \Omega, \\ \forall I \in \mathcal{I}, \quad \text{Card}'_p(g) &= \text{Card}_p^{g, I}(X_g^+(I)) \end{aligned}$$

Démonstration. Vérifions que l'opérateur $Card$ a bien les propriétés d'un opérateur de type "A" :

De manière évidente (somme sur $\hat{\Gamma}$), on a :

$$\text{équation 2.15 vraie} \quad (2.33)$$

$Card$ est indépendant de g donc :

$$\text{équation 2.16 vraie} \quad (2.34)$$

$Card$ est indépendant de I donc :

$$\forall N \subseteq \Omega, \forall M \subseteq \Omega \quad N \subseteq M \Rightarrow Card_p(N) \leq Card_p(M)$$

donc :

$$\text{équation 2.17 vraie} \quad (2.35)$$

Soit :

$$\Psi^{Card} = \left\{ \begin{array}{l} \psi_{\lambda}^{g,I} / \forall \lambda \in \mathbb{R}^+, \forall g \in \mathcal{G}, \forall F \in \mathcal{F}, \forall f \in \mathcal{S} \\ \psi_{\lambda}^{g,F}(f) = \max\{\alpha \in [0; f(p)] / Card_p^{g,im_1(F_{\alpha})}(f_{\alpha}) \geq \lambda \} \end{array} \right\} \quad (2.36)$$

équations 2.33, 2.34, 2.35, théorème 2.3.7 et définition 2.3.19 $\Rightarrow \forall \mathbb{R}^+ \quad \delta_{\Psi^{Card}}^{\lambda}$ est une EF de G_{λ}^{Card}

De plus, on peut aussi déduire que Ψ^{Card} est un EOCFE grâce aux équations 2.33, 2.34, 2.35 et au théorème 2.3.8. $\square$

Les deux théorèmes 2.3.9 et 2.3.10 ne sont que des exemples d'utilisation du formalisme introduit dans ce chapitre, illustrant ainsi la simplification rendue possible par ce dernier pour l'étude générale des filtres existants.

2.3.7 Définition de filtres dans un contexte de reconnaissance

Intéressons-nous maintenant à l'écriture générique d'un filtre dans le cadre de la reconnaissance. Pour faciliter le raisonnement, nous nous placerons dans le cadre du filtrage d'une image d'intervalles flous en utilisant la notion de zone plate floue. Néanmoins le même raisonnement peut être appliqué lorsque l'on part de la notion de compatibilité radiométrique associée aux images d'ombres floues.

Nous voulons donc extraire d'une image une structure B connexe, stable radiométriquement et dont nous connaissons certaines propriétés représentées par une fonction caractéristique $C : \mathcal{S} \rightarrow \{0, 1\}$. L'ensemble B appartient donc à $\{h \in \chi^F \cap \mathcal{H} / C(h) = 1\}$ et la borne supérieure de cet ensemble est une surestimation de $B : \xi^I = \bigvee \{h \in \chi^F \cap \mathcal{H} / C(h) = 1\}$. Si nous disposons d'une première surestimation f de B (obtenue par un traitement différent), nous définissons : $\xi^F(f) = \bigvee \{h \in \chi^F \cap \mathcal{H} / h \leq f \text{ et } C(h) = 1\}$. Cet opérateur est idempotent, croissant et antiextensif. Il s'agit donc d'une ouverture.

Nous reformulons ce filtre sur les différents niveaux $F(*, g)$ de l'image floue F associée à χ^F (on remarquera l'utilisation de l'opérateur d'agrégation pour interpréter le résultat) :

$$\xi^F(f) = \bigvee_{g \in \mathcal{G}} \bigvee \{h \in \mathcal{H} / h \leq \min(F(*, g), f) \text{ et } C(h) = 1\}. \quad (2.37)$$

Le calcul direct d'un tel filtre n'est pas possible en pratique. En effet $\mathcal{H}$ qui contient tous les ensembles flous connexes (et non toutes les composantes connexes) est de taille exponentielle (par rapport au cardinal de Ω). De plus, comme on ne s'intéresse pas seulement aux composantes connexes de $F(*, g)$ (pas de contrainte sur la maximalité des ensembles), on définit un filtre (avant agrégation) qui n'est pas forcément connexe.

Si le critère C est croissant, nous obtenons la simplification :

$$\xi^I(f) = \bigvee_{g \in \mathcal{G}} \bigvee \{h \in \mathcal{H}(\min(F(*, g), f)) / C(h) = 1\}$$

Le calcul de ce filtre peut donc être réalisé par la manipulation des zones plates floues de F' , avec $\forall g \in \mathcal{G} \ F'(*, g) = \min(F(*, g), f)$. Un algorithme permettant l'extraction efficace de ces zones plates floues sera présenté au chapitre 3. Dans ce cas on a bien avant agrégation un filtre connexe au sens de la définition 2.3.11.

Si le critère C n'est pas croissant, cette simplification n'est pas valable. Cependant puisque la manipulation directe de $\chi^F \cap \mathcal{H}$ n'est pas possible, nous nous contentons de l'approximation fournie par le calcul de $\xi^I(f)$ sur les zones plates floues de F' .

En pratique nous appliquons successivement différents filtres avec des critères différents. Si ces critères ne sont pas croissants, appliquer successivement ces filtres n'est pas équivalent à appliquer un filtre dont le critère est la conjonction des différents critères (en raison de l'approximation réalisée dans le cas non croissant). Dans ce cas il est intéressant de conserver le résultat des filtres sous forme d'une image floue qui sera traitée par le filtre suivant. En procédant ainsi nous conservons l'information perdue par l'agrégation réalisée sur les niveaux de gris par l'opérateur de l'équation 2.37. Nous écrivons donc les filtres comme des opérateurs sur les images floues :

$$\delta(F)(*, g) = \bigvee \{h \in \mathcal{H}(F(*, g)) / C(h) = 1\}. \quad (2.38)$$

Le résultat final est obtenu en réalisant l'agrégation sur les différents niveaux de l'image floue issue du dernier filtre en utilisant l'opérateur agg . On généralise ainsi l'étape d'agrégation des précédents filtres ξ .

On peut remarquer que le critère ici défini est strict dans la mesure où il fournit une mesure en tout ou rien. On peut assouplir cette contrainte en proposant une formulation de ce critère qui retourne un degré entre 0 et 1. De plus pour plus de généralité, on peut aussi le faire dépendre du niveau de gris de la composante ainsi que de l'image floue. On notera un tel critère $C^{g, F}$. La formulation d'une famille $\Psi = \{\psi^{g, F}\}$ de filtres utilisant ce critère est définie comme :

$$\forall f \in \mathcal{S}, \forall p \in \Omega \quad \psi^{F, g}(f)(p) = \max_{h \in \mathcal{H}(f)} (\min(h(p), C^{g, F}(h))) \quad (2.39)$$

Dans cette section, nous avons vu comment tirer parti de la représentation de l'imprécision contenue dans les images qui fut proposée à la section 2.2, pour définir des opérateurs connexes flous. Les propriétés de ces derniers étudiées jusqu'ici permettent d'appréhender leur comportement tout en comprenant la façon dont ils sont construits. A titre d'exemple, l'expression générique proposée à l'équation 2.39 va être utilisée dans la section 2.4 pour illustrer le potentiel de ce type de filtrage sur des images de tomosynthèse du sein.

2.4 Utilisation pratique des filtres connexes

Nous allons à présent présenter un filtre simple pour détecter les lésions circonscrites dans des volumes de tomosynthèse du sein. Pour ce faire nous allons introduire des outils permettant de construire une mesure discriminante pour ce type de lésions, puis nous définirons et étudierons ses propriétés. Enfin, nous nous attarderons sur son utilisation sur des *IOF* et des *IIF*.

2.4.1 Outils

Définition 2.4.1. En utilisant la définition proposée par Bloch et Maitre (1995), la dilatation floue $D : \mathcal{S} \times \mathcal{S} \rightarrow \mathcal{S}$ peut s'exprimer de la façon suivante :

$$\forall (f, h) \in \mathcal{S}^2, \forall p \in \Omega \quad D_h(f)(p) = \max_{p' \in \Omega} \{h(p - p') \top f(p')\}$$

Définition 2.4.2. La moyenne floue $fmean_{IOF} : \mathcal{F} \times \mathcal{S} \rightarrow \mathbb{R}^+$ se calcule de la façon suivante dans une *IOF* :

$$\forall F \in \mathcal{F}, \forall f \in \mathcal{S} \quad fmean_{IOF}(F, f) = \frac{\sum_{p \in \Omega} f(p) \sum_{g \in \mathcal{G}} F(p, g)}{fcard(f)}$$

avec $fcard$ le cardinal flou d'un ensemble flou défini par Rosenfeld et Haber (1985).

Cela peut s'interpréter comme le comptage du nombres d'éléments à l'intérieur F pondérés par leur degrés d'appartenance à f (c.f. figure 2.14).

Définition 2.4.3. Dans le cas d'une *IIF* ou d'une *INF*, la moyenne floue se calcule comme :

$$\forall F \in \mathcal{F}, \forall f \in \mathcal{S} \quad fmean_{IIF}(F, f) = \frac{\sum_{p \in \Omega} f(p) \frac{\sum_{g \in \mathcal{G}} g F(p, g)}{\sum_{g \in \mathcal{G}} F(p, g)}}{fcard(f)} \quad (2.40)$$

Ici, on pondère les valeurs des niveaux de gris défuzzifiées de F par les degrés d'appartenance de f (c.f. figure 2.14).

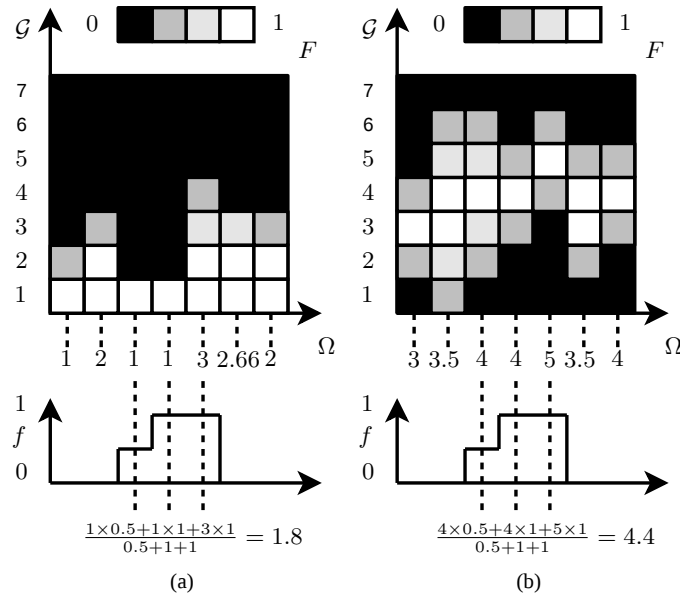


FIG. 2.14 – Exemple de calcul d'une moyenne floue pour une IOF (a) et une IIF (b).

Définition 2.4.4. Le contraste flou $fctrast : \Omega \times \mathcal{F} \times \mathcal{S} \rightarrow \mathbb{R}$ est défini de la manière suivante :

$$\forall F \in \mathcal{F}, \forall f \in \mathcal{S}, \forall h \in \mathcal{S} \\ fctrast_h(F, f) = fmean(F, f) - fmean(F, D_h(f) \cap \bar{f})$$

avec $\forall f \in \mathcal{S}, \forall p \in \Omega \quad \bar{f}(p) = 1 - f(p)$ et $fmean$ l'opérateur $fmean_{IOF}$ ou $fmean_{IIF}$ selon le type d'image à traiter.

L'idée est ici de soustraire la moyenne à l'intérieur et à l'extérieur de la composante. L'intérieur de la composante est modélisée par f , et l'extérieur est décrit par l'intersection de $\bar{f}$ et de la dilatation de f .

Théorème 2.4.1. L'union de deux opérateurs binaires connexes est un opérateur binaire connexe.

Démonstration. Soient o_1 et o_2 deux opérateurs connexes binaires.

Par définition, on a $\forall N \in 2^\Omega$:

$$\mathcal{C}(N \cap \overline{o_1(N)}) \subseteq \mathcal{C}(N) \quad (2.41)$$

$$\mathcal{C}(\overline{N} \cap o_1(N)) \subseteq \mathcal{C}(\overline{N}) \quad (2.42)$$

$$\mathcal{C}(N \cap \overline{o_2(N)}) \subseteq \mathcal{C}(N) \quad (2.43)$$

$$\mathcal{C}(\overline{N} \cap o_2(N)) \subseteq \mathcal{C}(\overline{N}) \quad (2.44)$$

De manière directe :

$$\begin{aligned} (2.41) \wedge (2.43) &\Rightarrow \mathcal{C}(N \cap \overline{o_1(N)}) \cap \mathcal{C}(N \cap \overline{o_2(N)}) \subseteq \mathcal{C}(N) \\ \mathcal{C}(N \cap \overline{o_1(N)} \cap \overline{N \cap o_2(N)}) &\subseteq \mathcal{C}(N) \text{ car } o_1 \text{ et } o_2 \text{ sont connexes} \\ \mathcal{C}(N \cap \overline{o_1(N)} \cap \overline{o_2(N)}) &\subseteq \mathcal{C}(N) \\ \mathcal{C}(N \cap (\overline{o_1(N)} \cup \overline{o_2(N)})) &\subseteq \mathcal{C}(N) \end{aligned}$$

De même :

$$\begin{aligned} (2.42) \wedge (2.44) &\Rightarrow \mathcal{C}(\overline{N} \cap o_1(N)) \cup \mathcal{C}(\overline{N} \cap o_2(N)) \subseteq \mathcal{C}(\overline{N}) \\ \mathcal{C}(\overline{N} \cap o_1(N) \cup \overline{N} \cap o_2(N)) &\subseteq \mathcal{C}(\overline{N}) \text{ car } o_1 \text{ et } o_2 \text{ sont connexes} \\ \mathcal{C}(\overline{N} \cap (o_1(N) \cup o_2(N))) &\subseteq \mathcal{C}(\overline{N}) \end{aligned}$$

Donc la réunion de deux opérateurs connexes est bien un opérateur connexe : $\square$

Théorème 2.4.2. *Le max de deux opérateurs CF est un opérateur CF.*

Démonstration. Soient o_1 et o_2 deux opérateurs CF. Le max de ces deux opérateurs CF forme un opérateur CF car :

$$\forall \alpha \in [0, 1], \forall f \in \mathcal{S} \quad \max(o_1(f), o_2(f))_\alpha = o_1(f)_\alpha \cup o_2(f)_\alpha$$

permet d'utiliser le théorème 2.4.1.

On a donc $o_1 \text{ CF} \wedge o_2 \text{ CF} \Rightarrow \max(o_1, o_2) \text{ CF}$. $\square$

2.4.2 Filtre de détection de lésions circonscrites

Soit $\Psi_{mass} = \{\psi_{mass}^F : \mathcal{S} \rightarrow \mathcal{S}/F \in \mathcal{F}\}$ un ensemble d'opérateurs tels que $\forall f \in \mathcal{S}, \forall p \in \Omega$:

$$\begin{aligned} \psi_{mass}^F(f)(p) = & \max_{h \in \mathcal{H}(f)} (\min(h(p), t_{u_1}(fcard(h)))) \\ & \top \max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_2}(fcomp(h)))) \\ & \top \max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_3}(fctrast(F, h)))) \end{aligned}$$

avec $fcomp$ la compacité floue d'un ensemble flou telle que définie dans Rosenfeld et Haber (1985), $t : \mathbb{R}^4 \times \mathbb{R} \rightarrow [0, 1]$ la fonction trapèze :

$$\forall (a, b, c, d) \in \mathbb{R}^4, \forall x \in \mathbb{R} \quad t_{a,b,c,d}(x) = \begin{cases} 0 & \text{si } x \leq a \vee x \geq d \\ 1 & \text{si } b \leq x \leq c \\ \frac{x-a}{b-a} & \text{si } a \leq x \leq b \\ \frac{x-d}{c-d} & \text{si } c \leq x \leq d \end{cases}$$

et $r : \mathbb{R}^2 \times \mathbb{R} \rightarrow [0, 1]$ la fonction rampe :

$$\forall (a, b) \in \mathbb{R}^2, \forall x \in \mathbb{R} \quad r_{a,b}(x) = \begin{cases} 0 & \text{si } x \leq a \\ 1 & \text{si } x \geq b \\ \frac{x-a}{b-a} & \text{sinon} \end{cases}$$

et où $u_1 \in \mathbb{R}^4$, $u_2 \in \mathbb{R}^2$ and $u_3 \in \mathbb{R}^2$ sont des constantes définissant ces fonctions. Les paramètres de ces dernières représentent l'information a priori de ce que l'on cherche à détecter, c'est-à-dire qu'ils permettent de définir de manière non stricte à partir de quelle valeurs on considère une composante connexe assez grande, assez compacte et assez contrastée.

Comme on peut le voir, cet ensemble Ψ_{mass} est composé de ψ_{mass} indépendants de g . Chacun de ces opérateurs peut être interprété comme l'agrégation entre le degré d'appartenance dérivé de mesures faites sur chacun des objets (composantes connexes floues de f) dans l'ensemble f en utilisant des données de F . On cherche donc ici des objets contrastés, compacts et d'une certaine taille.

L'opérateur $\delta_{\Psi_{mass}}$ dédié aux images à niveaux de gris associé à Ψ_{mass} peut être défini comme :

$$\forall F \in \mathcal{F}, \forall p \in \Omega, \forall g \in \mathcal{G} \quad \delta_{\Psi_{mass}}(F)(p, g) = \psi_{mass}^F(F(*, g))(p)$$

Théorème 2.4.3. Ψ_{mass} est un EOCCF.

Démonstration. Montrons tout d'abord que Ψ_{mass} est un EOCCF.

Montrons que les ψ_{mass} sont CF : De manière évidente $\mathcal{H}$ est CF. De plus l'opérateur qui à $f \in \mathcal{S}$ associe $\min(f, c)$ où c est une constante est lui aussi CF car il ne fait que supprimer le contenu des α -coupes de niveau supérieur à c . De plus le max de deux opérateurs (o_1 et o_2) CF forme un opérateur CF (c.f. théorème 2.4.2). On peut donc conclure que les trois opérateurs $\max_{h \in \mathcal{H}(f)} (\min(h(p), t_{u_1}(fcard(h))))$, $\max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_2}(fcomp(h))))$ et $\max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_3}(fctrast(F, h))))$ sont CF. Enfin en utilisant la propriété de monotonie de la t-norme ($(a \leq c) \wedge (b \leq d) \Rightarrow a \top b \leq c \top d$) on peut déduire que les ψ_{mass}^F sont CF.

Montrons que les ψ_{mass}^F sont anti-extensifs :

$$\forall p \in \Omega, \forall f \in \mathcal{S}, \forall h \in \mathcal{H}(f) \quad \min(h(p), t_{u_1}(fcard(h))) \leq f(p) \Rightarrow \max_{h \in \mathcal{H}(f)} (\min(h(p), t_{u_1}(fcard(h)))) \leq f(p)$$

de manière similaire, on a :

$$\forall p \in \Omega, \forall f \in \mathcal{S} \quad \max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_2}(fcomp(h)))) \leq f(p)$$

et

$$\forall p \in \Omega, \forall f \in \mathcal{S} \quad \max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_3}(fctrast(F, h)))) \leq f(p)$$

En utilisant encore la propriété de monotonie de la t-norme, on obtient que :

$$\forall p \in \Omega, \forall f \in \mathcal{S}, \forall F \in \mathcal{F} \quad \psi_{mass}^F(f)(p) \leq f(p)$$

donc Ψ_{mass} est un EOCCF. □

Cette famille d'opérateurs propose des propriétés faibles par rapport aux autres familles (EOCFS, EOCCFO, ou EOCCFE), néanmoins elle permet de définir un opérateur travaillant sur des images floues ($\mathcal{F}$) qui a l'intérêt d'aller chercher les objets qui ont une certaine propriété géométrique dans l'image sous l'hypothèse que les composantes connexes contenues dans l'image floue représentent ces objets.

2.4.3 Extraction à partir d'une IOF

Dans le cas où l'image floue est une IOF, on peut faire le lien entre le filtre $\delta_{\psi_{mass}}$ et les filtres classiques de type *amincissement* (Breen et Jones, 1996) qui fonctionnent de façon similaire. On remarquera que pour le traitement d'une IOF, on doit utiliser la mesure de moyenne (et donc de contraste) dédiée aux IOF.

Dans un volume tridimensionnel, la valeur des pixels est censée être une valeur proportionnelle au coefficient d'atténuation du matériau représenté. Ainsi, si on recherche des lésions caractérisées par une opacité de radiométrie plus importante que celle des structures environnantes, certaines des composantes connexes issues de la décomposition par multi-seuillages de l'image correspondent aux objets recherchés. On peut étendre le même raisonnement pour les composantes connexes floues extraites à partir d'une IOF. Les lésions circonscrites répondent à ces critères radiométriques, ainsi l'utilisation du filtre $\delta_{\psi_{mass}}$ peut potentiellement les détecter.

Une chaîne de traitement d'une mammographie I en 3D (coupe du volume reconstruit ou volume en entier) en utilisant ce filtre pourrait être décomposée comme suit :

- transformation de l'image I en image d'ombres floues F (utilisation de l'opérateur um , puis ajout possible de l'imprécision par exemple),
- filtrage de l'image F à l'aide de l'opérateur $\delta_{\Psi_{mass}}$,
- défuzzification partielle à l'aide de l'opérateur agg . Le résultat de $agg(\delta_{\Psi_{mass}}(F))$ indiquant le degré de possibilité des zones où se trouvent des lésions.

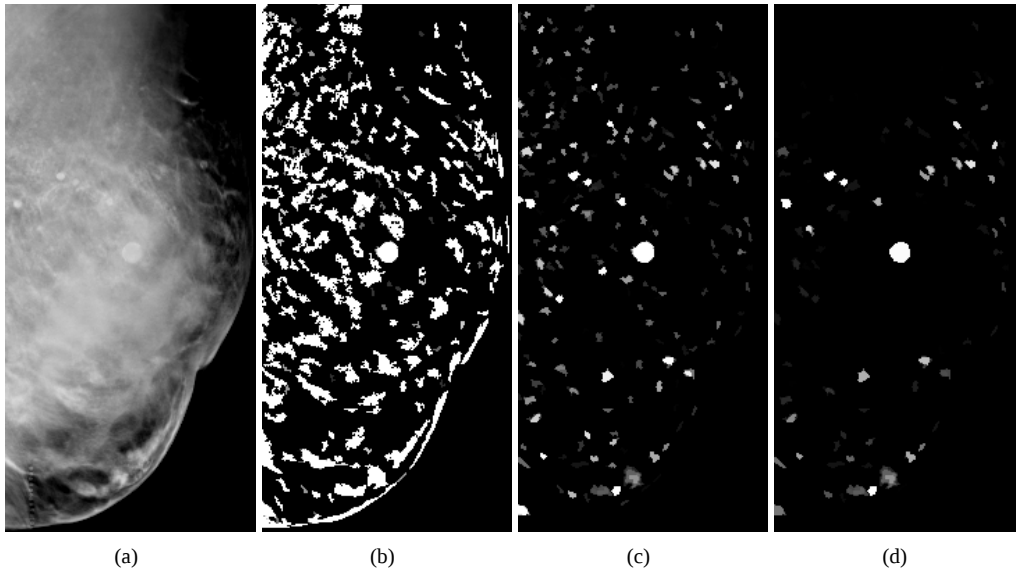


FIG. 2.15 – (d) Filtrage d’une coupe de volume de tomosynthèse du sein (a) par $\delta_{\Psi_{mass}}$. (b) Résultat de $\delta_{\Psi_{mass}}$ en ne prenant en compte que le critère d’aire. (c) Résultat en ne prenant en compte que les critères d’aire et de compacité.

La figure 2.15 illustre le résultat d’un tel procédé. Ici, une coupe d’un sein reconstruit avec des techniques itératives (Andersen et Kak, 1984) est filtrée par $\delta_{\Psi_{mass}}$. L’image est transformée en *IOF* par l’opérateur *um*. L’image de la figure 2.15(d) résultante du filtrage peut être interprétée pour chaque pixel p de Ω comme un degré d’appartenance à la classe *objet circonscrit*. Dans la mesure où, dans cet exemple, aucune imprécision n’est introduite sur les niveaux de gris avant filtrage, l’utilisation des ensembles flous sert uniquement à modéliser l’ambiguïté au niveau de la détection. Les degrés obtenus permettent d’évaluer à quel point les objets vérifient le critère de détection alors qu’une formulation nette du même filtre ne le permettrait pas (réponse en tout ou rien). Ainsi la sortie de ce filtre peut être un bon point de départ pour des traitements ultérieurs, comme de la classification (Peters, 2007), pour la détection des lésions.

On remarquera que l’hypothèse sur l’extraction des lésions par seuillages de l’image peut dans certains cas être invalidée dans la mesure où la limitation de la géométrie du système introduit des artefacts de reconstruction comme présentés au chapitre 1. Ainsi en cas de superposition de la lésion avec d’autres objets, la détection peut échouer. La distorsion en Z peut aussi reconnecter des structures disjointes dans la réalité, ainsi en pratique le traitement coupe par coupe donne de meilleurs résultats. Remarquons enfin que ce filtre ne peut généralement pas être utilisé en mammographie 2D standard, dans la mesure où dans une image de projection, l’hypothèse de seuillage n’est pas vérifiée.

2.4.4 Extraction à partir d’une *IIF*

Si maintenant on considère que l’image floue est une *IIF*, on ne considère que les zones de niveaux de gris à peu près constant dans l’image. On définit ainsi un $\delta_{\Psi_{mass}}$ qui correspond à un filtrage par zones plates (seule la définition de *fmean* change).

Pour illustrer le comportement de ce filtre pour des *IIF*, on peut se référer à la figure 2.16 qui présente le traitement d’une région d’intérêt contenant une lésion d’intensité à peu près constante. L’image d’intervalles flous est construite en remplaçant les intensités nettes de l’image originale par des trapèzes centrés en ces mêmes intensités. Le filtrage est effectué en utilisant $\delta_{\Psi_{mass}}$ avec de formulations nettes des fonctions d’appartenance des différents critères (les fonctions trapèze et rampe sont respectivement remplacées par des fonctions porte et de Heaviside). En ne considérant que le critère d’aire utilisé dans ce filtre, on obtient le résultat de la figure 2.16(b). Comme nous pouvons le voir ce seul critère d’aire n’est pas suffisant. En ajoutant les critères de compacité et de contraste (tel que défini pour les *INF* et *IIF*), on obtient le résultat

de la figure 2.16(c).

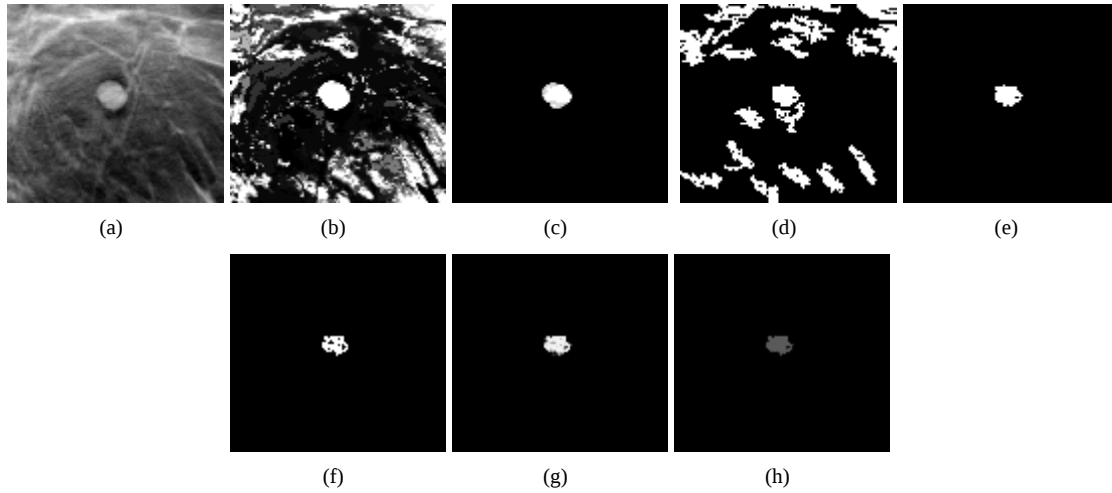


FIG. 2.16 – Filtrage d’une région d’intérêt d’une coupe de volume de tomosynthèse du sein et contenant une lésion circonscrite (a) en utilisant un critère net d’aire (b), et d’aire + compacité + contraste (e). Filtrage de la même région d’intérêt en utilisant les mêmes outils/critères mais dans une version non floue (les niveaux de gris sont soit complètement possibles, soit complètement impossibles). (d) Aire. (e) Aire + compacité + contraste. (f) Détection limite dans le cas net (si l’intervalle net déployé autour des niveaux de gris diminue, la lésion n’est plus détectée). (g) Détection limite dans le cas flou. (h) Détection limite dans le cas flou où on durcit les contraintes sur les critères utilisés.

Les figures 2.16(d) et 2.16(e) illustrent l’intérêt du flou pour représenter l’imprécision liée aux niveaux de gris. Le filtrage réalisé est en effet similaire à celui des figures 2.16(b) et 2.16(c) mais l’imprécision sur le niveau de gris est cette fois représentée par un intervalle classique plutôt que par un ensemble flou. Le filtre obtenu est plus sensible au paramètre de tolérance (représentant la largeur de l’intervalle, donc l’imprécision) qui en pratique sera restreint à des valeurs assez faibles pour éviter le regroupement de *zones plates*. Cela explique l’aspect morcelé du résultat obtenu. Pour illustrer cela plus en détail, on peut étudier les valeurs limites de largeur d’intervalles représentant les niveaux de gris pour lesquelles on commence à détecter la lésion. Dans le cas net qui est présenté à la figure 2.16(f), cette dernière apparaît partiellement avec des trous qui devraient intuitivement être supprimés. Si on considère maintenant le cas flou (c.f. figure 2.16(g)), on peut remarquer que ces derniers sont en partie remplis avec des degrés d’appartenance proches de 1 laissant apparaître une détection un peu plus plausible. Cette remarque est plus vraie encore pour des cas plus normaux comme celui de la figure 2.16(c) où une partie non négligeable de la lésion à un degré d’appartenance légèrement inférieur à 1. On remarquera cependant que le résultat peut être amélioré en utilisant des méthodes de construction plus évoluées de l’image floue comme celles proposées au chapitre 4. Contrairement à l’exemple précédent d’utilisation du filtre $\delta_{\Psi_{mass}}$ sur des images d’ombres floues où l’ambiguïté ne venait que du résultat du filtrage, ici nous avons seulement considéré l’imprécision provenant des niveaux de gris. On remarquera qu’on peut encore gagner en robustesse en considérant les deux sources d’ambiguïté conjointement, notamment dans la détection d’objets qui sont à la limite des critères utilisés. Ainsi, si on reprend l’exemple de détection limite dans le cas flou et si on n’utilise pas des fonctions d’appartenance en tout ou rien pour les différents critères, on peut encore réduire l’amplitude du patron déployé autour des niveaux de gris. Dans ce cas, la détection sera certes dégradée comme on peut le voir à figure 2.16(h) mais toujours utilisable.

Ce type de configurations filtre/image d’intervalles flous pourrait donc lui aussi être utilisé pour la détection automatique de zones d’intérêt dans des volumes de tomosynthèse du sein comme une première étape dans une chaîne de détection/classification automatique de lésions. Néanmoins, dans ce cas, une lésion doit, en plus d’être représentée par une sur-densité, être de radiométrie stable.

On remarquera enfin que ce type d’approche semble être une alternative intéressante aux tentatives de relaxation de la notion de zone plate que l’on peut trouver dans la littérature (Soille, 2008; Meyer et Maragos,

1999, 2000). Néanmoins, comme on l'a vu jusqu'à présent, on se trouve ici dans un cadre bien plus large permettant de définir une plus grande gamme de filtres.

2.5 Conclusion

Dans ce chapitre nous avons proposé une extension possible dans le cadre du flou des filtres connexes. Cette formulation repose sur la représentation de l'imprécision contenue dans les niveaux de gris de l'image par des quantités floues. Selon le type de filtrage que l'on souhaite mettre en œuvre, ces quantités floues peuvent apparaître sous différentes formes. Ainsi pour des filtrages fonctionnant sur zones plates, on utilisera des images de nombres ou d'intervalles flous, alors que pour des filtres d'amincissement ou les ouvertures d'attribut on préférera des images d'ombres floues qui seront plus adaptées.

Nous avons aussi montré comment exprimer des filtres classiques comme les ouvertures d'aire ou les arasements volumiques. Ces extensions floues de filtres classiques nous permettent de nous conforter sur la cohérence de notre nouvel environnement théorique fondé sur les images floues. De plus en utilisant la flexibilité fournie par les ensembles flous, de nouveaux filtres peuvent être écrits. Nous avons proposé une étude générale de leur propriétés. De tels filtres ont aussi été illustrés par un exemple de détecteur d'objets circonscrits dans des *IIF* ou des *IOF* qui peut être une première étape de marquage de lésions circonscrites dans des volumes de tomosynthèse du sein.

Chapitre 3

Mise en œuvre des filtres connexes flous

Les filtres connexes (Vincent, 1993; Serra et Salembier, 1993; Salembier et Serra, 1995; Heijmans, 1997; Vachier, 2001) sont largement utilisés en traitement d'images. Certains, comme les filtres de nivellement (Meyer, 1998a,b; Vachier, 2001), permettent par exemple de simplifier des images et d'autres sont utilisés dans des tâches de segmentation. Dans tous les cas, ils reposent essentiellement sur la définition binaire de composantes connexes et leur comportement est sujet à la façon dont cette définition est utilisée pour traiter des images à niveaux de gris. Par exemple, des composantes connexes peuvent être extraites à partir des zones plates de l'image (Salembier et Serra, 1995; Crespo *et al.*, 1997) ou à partir de seuillages successifs de l'image (Vincent, 1993). Généralement dans le second cas, les filtres sont mis en œuvre en utilisant une représentation par arbre de coupes de l'image (Salembier *et al.*, 1998; Meijster et Wilkinson, 2002; Berger *et al.*, 2007). Les nœuds de tels arbres représentent les composantes connexes issues des différents seuillages de l'image tandis que les arcs correspondent à l'ordre d'inclusion, qui est fonction des niveaux de gris, de ces différentes composantes connexes.

Dans le chapitre précédent, nous avons proposé une extension générale des filtres connexes pour les images à niveaux de gris dans le cadre des ensembles flous. Cette extension repose sur la représentation des niveaux de gris d'une image par l'intermédiaire de quantités floues. L'extraction d'ensembles flous définis sur le domaine de l'image à partir de ces quantités floues associée à une définition floue de connexité permet de formuler de manière naturelle des extensions aux filtres connexes classiques. Selon la nature des quantités floues utilisées pour représenter la valeur des pixels, on définira des filtres dont le comportement se rapproche plutôt des filtres de nivellement (dans le cas de quantité floues décroissantes) ou des filtres par zones plates (dans le cas d'intervalles flous). Néanmoins, la question de la mise en œuvre n'a pas encore été abordée et peut poser problème dans la mesure où la quantité de données à traiter augmente grandement par rapport à des images à niveaux de gris classiques, de plus la définition mathématique des filtres ne semble pas adaptée à leur mise en œuvre de manière efficace.

La représentation sous forme d'arbres n'est pas limitée aux images à niveaux de gris. En fait, les ensembles flous peuvent eux aussi être représentés de la même manière. Cela est particulièrement valable lorsque l'on souhaite travailler avec les définitions classiques de composantes connexes floues dans la mesure où ces dernières peuvent être manipulées en considérant l'emboîtement de leurs α -coupes qui est adapté à une représentation arborescente. Le problème qui apparaît lorsque l'on veut utiliser ce formalisme pour représenter et filtrer les images floues est la nécessité d'avoir une représentation sous forme d'arbre pour chaque ensemble flou associé aux différents niveaux de gris. Même si la construction d'un tel arbre peut être faite en un temps quasi-linéaire (Najman et Couprie, 2006), ou même de façon parallèle (Wilkinson *et al.*, 2008), si le nombre de niveaux de gris est grand, la construction de manière indépendante de tous ces arbres n'est pas utilisable en pratique. En effet, cela résulterait en une complexité de l'ordre de $O(GN)$ (avec G le nombre de niveaux de gris et N le nombre de pixels dans l'image) dans le cas d'une mise en œuvre de l'étape de construction de complexité linéaire, ou même pire si on utilise des algorithmes plus lents. Néanmoins, deux ensembles flous extraits d'une image floue pour des niveaux de gris successifs sont généralement assez similaires l'un à l'autre dans la mesure où ils ne diffèrent que sur un nombre limité de pixels. Cela motive nos travaux sur le design d'un algorithme de mise à jour d'un arbre

représentant un ensemble flou en un autre arbre représentant un autre ensemble flou légèrement différent. Un tel algorithme pourrait aussi être utilisé dans d'autres champs applicatifs qui utilisent les représentations arborescentes.

Dans ce chapitre nous proposons un ensemble d'algorithmes visant à mettre à jour un arbre associé à un ensemble flou quand ce dernier voit ses degrés d'appartenance changer par endroit. Tout d'abord, nous introduisons les notations et définitions nécessaires à l'élaboration des algorithmes. Puis, nous proposerons des opérateurs permettant de faire croître/décroître un arbre associé à un ensemble flou en fonction des changements qu'il subit. Les preuves de validité de ces opérateurs seront aussi détaillées respectivement dans les sections 3.2 et 3.3. Enfin nous décrirons dans la section 3.4 les algorithmes résultant permettant le filtrage d'images floues.

3.1 Notations et définitions

Dans les développements théoriques qui suivent, nous réutiliserons les notations déjà introduites dans le chapitre précédent que nous rappelons dans un souci de clarté. Ainsi Ω correspondra au domaine de l'image, c'est-à-dire un domaine borné muni d'une connexité discrète. Le voisinage d'un point $p \in \Omega$ en utilisant la connexité définie sur Ω sera noté $\text{voisinage}(p)$. L'ensemble des niveaux de gris sera noté $\mathcal{G}$. L'ensemble des sous-ensembles flous définis sur Ω ($\Omega \rightarrow [0, 1]$) sera noté $\mathcal{S}$. L'ensemble 2^Ω représente l'ensemble des ensembles nets qui sont inclus dans Ω . De manière évidente, on a $2^\Omega \subset \mathcal{S}$. Enfin $\mathcal{F}$ représentera l'ensemble des ensembles flous définis sur $\Omega \times \mathcal{G}$ ($\Omega \times \mathcal{G} \rightarrow [0, 1]$).

Nous utiliserons aussi un ensemble de notations permettant de manipuler les structures arborescentes sous forme de graphes. Le choix d'une telle structure de données est motivée par des considérations de mise en pratique. En effet, tout comme c'est le cas pour la mise en œuvre de filtres connexes classiques, une telle représentation permet de faciliter, en terme de complexité, la manipulation des données pour effectuer l'étape de filtrage. Comme évoqué précédemment, une énumération des différentes composantes connexes associées aux différents niveaux de gris (ou α -coupes dans le cas d'ensembles flous) explose aussi bien en temps de calcul qu'en taille de mémoire. Enfin la structure arborescente proposée permet, comme on le verra tout au long de ce chapitre, de modifier efficacement un ensemble flou de manière locale, tout en étant capable de connaître les composantes connexes des différentes α -coupes de ce dernier, permettant ainsi la mise en œuvre d'une stratégie de filtrage rapide d'une image floue. Ainsi $\mathcal{V} = [0, 1] \times 2^\Omega$, représentera l'ensemble des nœuds possibles. Chaque nœud (α, P) contient un degré d'appartenance et un ensemble P de points inclus dans Ω . L'ensemble de tous les arcs possibles entre les différents nœuds est noté $\mathcal{E} = \mathcal{V} \times \mathcal{V}$. L'ensemble des arbres définis comme des graphes acycliques représentés par des couples (V, E) de nœuds $V \subseteq \mathcal{V}$ et d'arcs $(n_1, n_2) \in E$ sera noté $\mathcal{T}$. Un arc $(n_1, n_2) \in E$ s'interprète comme n_1 est le père n_2 . Remarquons que sans ajouter de contraintes supplémentaires, de tels arbres ne représentent pas forcément un sous-ensemble flou comme il le sera détaillé à la définition 3.1.10.

Ce formalisme, qui va être utilisé tout au long de ce chapitre, est nécessaire pour proposer une validation sans ambiguïté des aspects théoriques. Pour aider à la compréhension, un ensemble d'illustrations 1D des différentes définitions et théorèmes sera proposé. Remarquons cependant que ces exemples ont seulement pour but d'éclaircir les aspects formels et ne peuvent pas se substituer à ces derniers qui, eux seuls, permettent de s'assurer de la validité des méthodes proposées, notamment dans des dimensions supérieures.

Définition 3.1.1. Soit $\mathcal{A} : \mathcal{V} \rightarrow [0, 1]$ l'opérateur qui retourne le degré d'appartenance associé à un nœud :

$$\forall n = (\alpha, P) \in \mathcal{V} \quad \mathcal{A}(n) = \alpha \quad (3.1)$$

$\mathcal{A}$ est l'opérateur d'association d'un nœud vers l'ensemble des degrés d'appartenance. Comme on le verra plus loin, lorsque n est un nœud d'un arbre qui représente un ensemble flou, $\mathcal{A}(n)$ correspond au degré de l' α -coupe représentée par n .

Définition 3.1.2. Soit $\mathcal{P} : \mathcal{V} \rightarrow 2^\Omega$ l'opérateur qui associe à un nœud l'ensemble des points qu'il contient :

$$\forall n = (\alpha, P) \in \mathcal{V} \quad \mathcal{P}(n) = P \quad (3.2)$$

$\mathcal{P}$ est l'opérateur d'association d'un nœud vers l'ensemble de sous-ensembles de Ω .

Définition 3.1.3. Soit $\mathcal{N} : \mathcal{T} \rightarrow 2^{\mathcal{V}}$ l'opérateur qui associe à un arbre l'ensemble de ses nœuds :

$$\forall t = (V, E) \in \mathcal{T} \quad \mathcal{N}(t) = V \quad (3.3)$$

$\mathcal{N}$ est l'opérateur d'association d'un arbre vers l'ensemble des sous-ensembles de nœuds.

Définition 3.1.4. Soit $\mathcal{B} : \mathcal{T} \rightarrow 2^{\mathcal{E}}$ l'opérateur qui associe à un arbre l'ensemble de arcs entre ses nœuds :

$$\forall t = (V, E) \in \mathcal{T} \quad \mathcal{B}(t) = E \quad (3.4)$$

$\mathcal{B}$ est l'opérateur d'association d'un arbre vers l'ensemble des sous-ensembles des relations père/fils.

Définition 3.1.5. Soit $\mathcal{W} : \mathcal{T} \times \mathcal{V} \rightarrow 2^{\mathcal{V}}$ l'opérateur qui retourne l'ensemble des fils d'un nœud d'un arbre :

$$\forall t \in \mathcal{T}, \forall n \in \mathcal{N}(t) \quad \mathcal{W}(t, n) = \bigcup_{v \in \mathcal{N}(t)/(n, v) \in E} \{v\} \quad (3.5)$$

Définition 3.1.6. Soit $\mathcal{R} : \mathcal{T} \rightarrow \mathcal{V}$ l'opérateur qui retourne la racine d'un arbre :

$$\forall t \in \mathcal{T} \quad \mathcal{R}(t) = n \in \mathcal{N}(t) / \forall v \in \mathcal{N}(t) \quad (v, n) \notin \mathcal{B}(t) \quad (3.6)$$

Définition 3.1.7. Soit $\mathcal{D} : \mathcal{T} \times \mathcal{V} \rightarrow 2^{\mathcal{V}}$ l'opérateur qui retourne les descendants (fils, fils des fils, etc.) d'un nœud d'un arbre qui est défini de manière récursive comme :

$$\forall t \in \mathcal{T}, \forall n \in \mathcal{N}(t) \quad \mathcal{D}(t, n) = \bigcup_{v \in \mathcal{W}(t, n)} (\{v\} \cup \mathcal{D}(t, v)) \quad (3.7)$$

Remarquons que $n \notin \mathcal{D}(t, n)$ et que si n est une feuille alors $\mathcal{W}(t, n) = \emptyset$ et $\mathcal{D}(t, n) = \emptyset$.

Définition 3.1.8. $\mathcal{D}' : \mathcal{T} \times \mathcal{V} \rightarrow 2^{\mathcal{V}}$ est le même opérateur que celui précédemment défini, mais son résultat contient aussi n :

$$\forall t \in \mathcal{T}, \forall n \in \mathcal{N}(t) \quad \mathcal{D}'(t, n) = \mathcal{D}(t, n) \cup \{n\} \quad (3.8)$$

Définition 3.1.9. L'opérateur $sa : \mathcal{T} \times \mathcal{V} \rightarrow 2^{\mathcal{T}}$, qui associe à un nœud d'arbre l'ensemble des sous-arbres dont la racine est un fils de ce nœud, est défini comme suit :

$$\forall t = (V, E) \in \mathcal{T}, \forall n \in V \quad sa(t, n) = \bigcup_{v \in \mathcal{W}(t, n)} \left\{ \left(\mathcal{D}'(t, v), \bigcup_{(v_1, v_2) \in \mathcal{D}'(t, v)^2 / (v_1, v_2) \in E} \{(v_1, v_2)\} \right) \right\} \quad (3.9)$$

Remarquons que n n'appartient à aucun des sous-arbres résultants car leurs racines sont des fils de n .

Définition 3.1.10. Soit un arbre $t \in \mathcal{T}$, t est un emboîtement si :

$$\forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{A}(s) > \mathcal{A}(n) \quad (3.10)$$

$$\wedge \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n) \quad (3.11)$$

$$\wedge \forall n \in \mathcal{N}(t) \quad \bigcap_{s \in \mathcal{W}(t, n)} \mathcal{P}(s) = \emptyset \quad (3.12)$$

La notion d'emboîtement introduit un ordre sur les nœuds en fonction des α (équation 3.10). Elle implique aussi que suivant cet ordre, on doit avoir une certaine inclusion (équation 3.11) et elle interdit d'avoir un recouvrement entre les descendants d'un même nœud (équation 3.12).

Définition 3.1.11. $\forall t \in \mathcal{T}, \forall f \in \mathcal{S}$, t est une représentation de f si :

$$t \text{ est un emboîtement} \quad (3.13)$$

$$\forall p \in \Omega, \forall \alpha \in [0, 1] \quad p \in f_\alpha \Rightarrow \exists n \in \mathcal{N}(t) / \mathcal{P}(n) = \hat{\Gamma}_{f_\alpha}^p \quad (3.14)$$

$$\forall n \in \mathcal{N}(t), \forall p \in \Omega \quad p \in \mathcal{P}(n) \Rightarrow \hat{\Gamma}_{f_{\mathcal{A}(n)}}^p = \mathcal{P}(n) \quad (3.15)$$

avec $\hat{\Gamma}_N^p$ la composante connexe contenant le point p dans N , et f_α l' α -coupe d'un ensemble.

Dans cette définition, on ajoute une contrainte de correspondance des composantes connexes des α -coupes de f vers les nœuds de t (équation 3.14) et une association des nœuds de t vers les α -coupes de f (équation 3.15). La figure 3.1 illustre le fait qu'un arbre est une représentation compacte d'un ensemble flou, et notamment le fait qu'une α -coupe (α_0 dans la figure) peut être représentée de manière indirecte par un nœud n qui n'a pas le même degré d'appartenance (ici $\mathcal{A}(n) = \alpha_1$). En effet, tous les niveaux de quantification de α ne sont pas forcément représentés explicitement par un niveau de l'arbre.

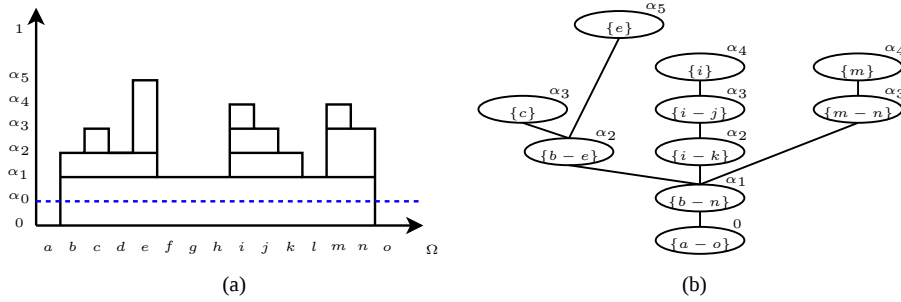


FIG. 3.1 – Représentation d'un ensemble flou en utilisant un arbre de coupes. La seule composante connexe contenue dans l' α -coupe α_0 est indirectement représentée par le nœud qui contient les points b à n et qui a un degré d'appartenance égal à α_1 .

L'ensemble des arbres $t \in \mathcal{T}$ qui sont des emboîtements est noté $\mathcal{T}^e$ ($\mathcal{T}^e \subset \mathcal{T}$).

Définition 3.1.12. L'ensemble $Q_p(T)$ des arbres inclus dans un ensemble d'emboîtements T et dont la racine contient p est défini de la façon suivante :

$$\forall p \in \Omega, \forall T \in 2^{\mathcal{T}^e} \quad Q_p(T) = \{t \in T / p \in \mathcal{P}(\mathcal{R}(t))\} \quad (3.16)$$

Théorème 3.1.1. Soit t un emboîtement. Si on considère un ensemble de nœuds frères dans cet arbre (ensemble des fils $sa(t, n)$ d'un nœud n donné), et un point $p \in \Omega$, alors p est inclus dans la racine d'au plus un des sous-arbres associés à ces frères. Cela signifie que l'information à propos d'un ensemble flou représenté par t à un point p est contenue dans une unique branche de t .

$$\forall t \in \mathcal{T}^e, \forall n \in \mathcal{N}(t), \forall p \in \Omega \quad (|Q_p(sa(t, n))| = 0) \vee (|Q_p(sa(t, n))| = 1) \quad (3.17)$$

avec $|\cdot|$ le cardinal d'un ensemble.

Démonstration. En utilisant la définition 3.1.10, on a : $\bigcap_{v \in \mathcal{W}(n)} \mathcal{P}(v) = \emptyset$. Ainsi pour un p donné, il existe au plus un seul $v \in \mathcal{W}(n)$ tel que $p \in \mathcal{P}(v)$. $\square$

3.2 Croissance d'un arbre représentant un ensemble flou

Nous allons maintenant introduire un nouvel opérateur destiné à fusionner deux arbres t_1 et t_2 autour des points p_1 et p_2 , qui sont contenus dans les racines respectives de t_1 et t_2 (c.f. figure 3.2). Concrètement, un des points va correspondre à l'endroit où l'ensemble flou croît (i.e. le degré d'appartenance associé à

ce point est remplacé par une valeur plus grande), et l'autre point va correspondre successivement aux différents points qui sont contenus dans son voisinage. Cela permet de reconnecter les sous-arbres, et ainsi les composantes connexes qui sont marquées par ces points. Cet opérateur va avoir un comportement similaire à celui récemment introduit par Wilkinson *et al.* (2008) pour reconnecter deux arbres représentant une image à niveaux de gris sur des sous-domaines voisins. Ici nous utilisons un formalisme de plus haut niveau dans le but d'utiliser notre opérateur dans un procédé destiné à mettre à jour/faire croître un arbre. Cela permet aussi d'avoir une expression plus générique par rapport à la structure de données utilisée lors de la mise en œuvre de notre algorithme comme il sera présenté à la section 3.4.

Remarquons que la capacité pour l'opérateur que nous introduisons de produire un arbre représentant un ensemble flou va dépendre des données passées en entrée (i.e. sous-arbres et marqueurs). Un processus pour sélectionner correctement ces données sera décrit plus tard dans cette section (c.f. définition 3.2.4).

L'opérateur de fusion que nous allons présenter se décompose en cinq cas distincts. Ces cas sont en fait obtenus en considérant toutes les possibilités en termes d'(in)égalité des degrés associés aux racines des arbres à fusionner ainsi que toutes les configurations au niveau de l'inclusion des marqueurs (p_1, p_2) dans les sous-arbres issus de ces racines. Une table de vérité montrant de manière plus formelle que toutes les possibilités sont bien couvertes par ces cinq cas est proposée à l'annexe C.1.

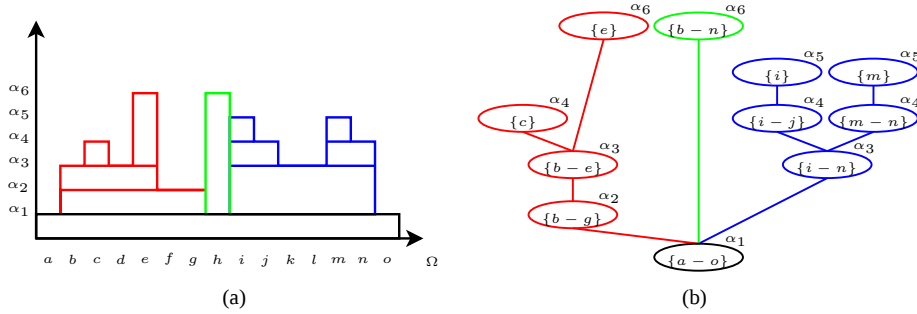


FIG. 3.2 – Illustration du processus de croissance. (a) Ensemble flou à représenter après fusion. La valeur au point h passe de α_1 à α_6 . (b) Arbre, qui n'est pas une représentation d'un ensemble flou, qui contient les sous-arbres (en rouge, vert et bleu) à fusionner. Pour restaurer les composantes connexes des différentes α -coupes, l'arbre en rouge et l'arbre en vert doivent être fusionnés en utilisant les marqueurs respectifs g et h , et l'arbre résultant doit être fusionné avec l'arbre bleu en utilisant les marqueurs respectifs h et i . Les marqueurs g, h et i sont utilisés pour sélectionner quels nœuds doivent être fusionnés, ils représentent les zones où la connexité peut avoir changé (voisinage du point qui a vu sa valeur croître).

Définition 3.2.1. Soit $\overline{M} : \mathcal{T}^e \times \mathcal{T}^e \times \Omega \times \Omega \rightarrow \mathcal{T}$ l'opérateur de fusion défini comme suit :

$$\forall t_1 = (V_1, E_1) \in \mathcal{T}^e, \forall t_2 = (V_2, E_2) \in \mathcal{T}^e, \forall (p_1, p_2) \in \Omega^2$$

cas 1 $(\mathcal{A}(\mathcal{R}(t_1)) = \mathcal{A}(\mathcal{R}(t_2))) \wedge ((|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 0) \vee (|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))| = 0))$

Dans ce cas, $\overline{M}(t_1, t_2, p_1, p_2) = (V_3, E_3)$ avec :

$$V_3 = \bigcup_{t \in sa(t_1, \mathcal{R}(t_1)) \cup sa(t_2, \mathcal{R}(t_2))} \{(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2)))\} \cup \mathcal{N}(t) \quad (3.18)$$

qui correspond à la nouvelle racine et à tous les nœuds de t_1 et t_2 excepté leurs racines.

$$E_3 = \mathcal{B}(t_1) \setminus \bigcup_{v \in \mathcal{W}(t_1, \mathcal{R}(t_1))} \{(\mathcal{R}(t_1), v)\} \cup \mathcal{B}(t_2) \setminus \bigcup_{v \in \mathcal{W}(t_2, \mathcal{R}(t_2))} \{(\mathcal{R}(t_2), v)\} \cup \bigcup_{v \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2))} \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), v)\} \quad (3.19)$$

où les deux premiers termes correspondent aux arcs des arbres t_1 et t_2 qui n'atteignent respectivement ni $\mathcal{R}(t_1)$ ni $\mathcal{R}(t_2)$. Le troisième terme permet d'introduire les arcs faisant référence à la nouvelle racine.

cas 2 $(\mathcal{A}(\mathcal{R}(t_1)) = \mathcal{A}(\mathcal{R}(t_2))) \wedge (|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 1) \wedge (|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))| = 1)$

Dans ce cas $\overline{M}(t_1, t_2, p_1, p_2) = (V_3, E_3)$ avec :

$$\begin{aligned} V_3 = & \{(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2)))\} \\ & \cup \bigcup_{t \in sa(t_1, \mathcal{R}(t_1))/p_1 \notin \mathcal{P}(\mathcal{R}(t))} \mathcal{N}(t) \\ & \cup \bigcup_{t \in sa(t_2, \mathcal{R}(t_2))/p_2 \notin \mathcal{P}(\mathcal{R}(t))} \mathcal{N}(t) \\ & \cup \mathcal{N}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2)) \end{aligned} \quad (3.20)$$

où la première partie correspondant à la création du nœud racine, les deuxième et troisième à l'insertion des sous-arbres dont la racine ne contient ni p_1 , ni p_2 , et la quatrième aux nœuds issus de la fusion entre les sous-arbres dont la racine contient l'un des deux précédents points.

$$\begin{aligned} E_3 = & \bigcup_{t \in sa(t_1, \mathcal{R}(t_1))/p_1 \notin \mathcal{P}(\mathcal{R}(t))} \mathcal{B}(t) \\ & \cup \bigcup_{t \in sa(t_2, \mathcal{R}(t_2))/p_2 \notin \mathcal{P}(\mathcal{R}(t))} \mathcal{B}(t) \\ & \cup \bigcup_{v \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2))/p_1 \notin \mathcal{P}(v) \wedge p_2 \notin \mathcal{P}(v)} \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), v)\} \\ & \cup \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), \mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2)))\} \\ & \cup \mathcal{B}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2)) \end{aligned} \quad (3.21)$$

où les deux premiers termes correspondent aux arcs des sous-arbres dont la racine ne contient ni p_1 , ni p_2 , le troisième terme aux arcs reliant la racine de ces sous-arbres à la nouvelle racine, le quatrième à l'arc reliant cette même racine à l'arbre résultant de la fusion des deux sous-arbres dont la racine contient p_1 ou p_2 , et le cinquième aux arcs issus de cette fusion.

cas 3 $(\mathcal{A}(\mathcal{R}(t_1)) < \mathcal{A}(\mathcal{R}(t_2))) \wedge (|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| \neq 0)$

Dans ce cas, $\overline{M}(t_1, t_2, p_1, p_2) = (V_3, E_3)$

$$\begin{aligned} V_3 = & \{(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2)))\} \\ & \cup \bigcup_{t \in sa(t_1, \mathcal{R}(t_1))/p_1 \notin \mathcal{P}(\mathcal{R}(t))} \mathcal{N}(t) \\ & \cup \mathcal{N}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2)) \end{aligned} \quad (3.22)$$

où la première partie correspond au nouveau nœud racine, la seconde partie aux nœuds des sous-arbres de t_1 dont la racine ne contient pas p_1 , la troisième à la fusion du sous-arbre restant avec t_2 .

$$\begin{aligned} E_3 = & \bigcup_{t \in sa(t_1, \mathcal{R}(t_1))/p_1 \notin \mathcal{P}(\mathcal{R}(t))} \mathcal{B}(t) \\ & \cup \bigcup_{v \in \mathcal{W}(t_1, \mathcal{R}(t_1))/p_1 \notin \mathcal{P}(v)} \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), v)\} \\ & \cup \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), \mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2)))\} \\ & \cup \mathcal{B}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2)) \end{aligned} \quad (3.23)$$

où le premier terme correspond aux arcs provenant des sous-arbres de t_1 dont la racine ne contient pas p_1 , le deuxième aux liens entre la racine de ces derniers et le nouveau nœud racine, le troisième au lien avec la racine de l'arbre résultant de la fusion, et le quatrième aux arcs résultant de cette même fusion.

cas 4 $(\mathcal{A}(\mathcal{R}(t_1)) < \mathcal{A}(\mathcal{R}(t_2))) \wedge (|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 0)$

Dans ce cas, $\overline{M}(t_1, t_2, p_1, p_2) = (V_3, E_3)$ avec :

$$\begin{aligned}
V_3 = & \{(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2)))\} \\
& \cup \mathcal{N}(sa(t_1, \mathcal{R}(t_1))) \\
& \cup \mathcal{N}(t_2)
\end{aligned} \tag{3.24}$$

avec dans l'ordre le nouveau nœud racine, l'ensemble des nœuds de t_1 privé de sa racine et l'ensemble de nœuds de t_2 .

$$\begin{aligned}
E_3 = & \bigcup_{t \in sa(t_1, \mathcal{R}(t_1)) \cup \{t_2\}} \mathcal{B}(t) \\
& \cup \bigcup_{v \in \mathcal{W}(t_1, \mathcal{R}(t_1))} \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), v)\} \\
& \cup \{((\mathcal{A}(\mathcal{R}(t_1)), \mathcal{P}(t_1) \cup \mathcal{P}(t_2)), \mathcal{R}(t_2))\}
\end{aligned} \tag{3.25}$$

où la première partie correspond à l'ensemble des arcs de t_2 et des sous-arbres de t_1 , en deuxième les arcs partant du nouveau nœud racine vers les sous-arbres de t_1 et en troisième le lien entre la nouvelle racine et la racine de t_2 .

cas 5 $\mathcal{A}(\mathcal{R}(t_1)) > \mathcal{A}(\mathcal{R}(t_2))$

Dans cette dernière configuration, on a $\overline{M}(t_1, t_2, p_1, p_2) = \overline{M}(t_2, t_1, p_2, p_1)$ et on se ramène aux cas 3 et 4.

Remarque 3.2.1. Notons qu'aucune autre configuration n'est possible pour l'opérateur $\overline{M}$ dans la mesure où la définition couvre toutes les combinaisons possibles de valeurs pour $|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))|$ et $|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))|$ et ce pour toutes les égalités/inégalités entre $\mathcal{A}(\mathcal{R}(t_1))$ et $\mathcal{A}(\mathcal{R}(t_2))$.

Illustrons maintenant la définition 3.2.1 sur des exemples correspondant aux quatre premiers cas.

Le cas 1 est illustré à la figure 3.3 pour $p_1 = h$ et $p_2 = i$ avec t_1 l'arbre de la figure 3.3(b) et t_2 l'arbre de la figure 3.3(c). Dans ce cas, on a $|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 0$ car seule la racine de t_1 contient le point h (p_1), et $|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))| = 1$ car il existe un sous-arbre de t_2 qui contient i (p_2) et dont la racine est un nœud fils de la racine de t_2 . Pour fusionner ces deux arbres dont les nœuds racine ont même degré d'appartenance, il suffit donc juste de fusionner les deux nœuds racine de t_1 et t_2 (nœud en noir sur la figure 3.3(d)) et d'y raccrocher tous les sous-arbres issus des fils des deux précédents nœuds racines (nœuds en rouge et bleu sur la figure 3.3(d)).

La figure 3.4 illustre le deuxième cas avec toujours $p_1 = h$ et $p_2 = i$. Les arbres t_1 et t_2 sont respectivement présentés aux figures 3.4(b) et 3.4(c). Pour les deux arbres on a p_1 et p_2 qui sont inclus respectivement dans les nœuds racines $\mathcal{R}(t_1)$ et $\mathcal{R}(t_2)$ mais aussi dans un fils de chacun de ces nœuds. Pour cette raison, on a $|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 1$ et $|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))| = 1$. Comme les nœuds racines ont le même degré d'appartenance, on peut les fusionner de la même manière que dans le cas précédent (nœud en noir sur la figure 3.4(d)). Néanmoins le traitement des sous-arbres restants doit se faire de manière différente : les sous-arbres dont la racine ne contient ni p_1 , ni p_2 peuvent être rattachés directement au nouveau nœud racine (par exemple l'arbre bleu hors de la zone grise sur la figure 3.4(d)), les deux sous-arbres dont les racines contiennent respectivement p_1 et p_2 ont quant à eux besoin d'être fusionnés (zone grise sur la figure 3.4(d)).

La troisième configuration est illustrée à la figure 3.5 avec toujours $p_1 = h$ et $p_2 = i$. Les arbres t_1 et t_2 sont représentés respectivement aux figures 3.5(b) et 3.5(c). Ce cas est assez similaire à la précédente configuration, la principale différence étant la non égalité des degrés d'appartenance des nœuds racines de t_1 et t_2 . Le nouveau nœud racine (en noir sur la figure 3.5(d)) est donc composé de l'union des points des racines de t_1 et t_2 et du degré d'appartenance de la racine de t_1 (qui est plus petit que celui de la racine de t_2). Les sous-arbres de t_1 dont la racine ne contient pas p_1 sont directement rattachés au nouveau nœud racine et le sous-arbre dont la racine contient p_1 est fusionné avec t_2 (en gris sur la figure 3.5(d)).

Le quatrième cas est illustré à la figure 3.6. Ce cas est très semblable au précédent sauf qu'il n'y a pas de sous-arbres de t_1 dont la racine contient le point p_1 . Ainsi on se retrouve dans un cas terminal où aucune fusion supplémentaire n'est nécessaire.

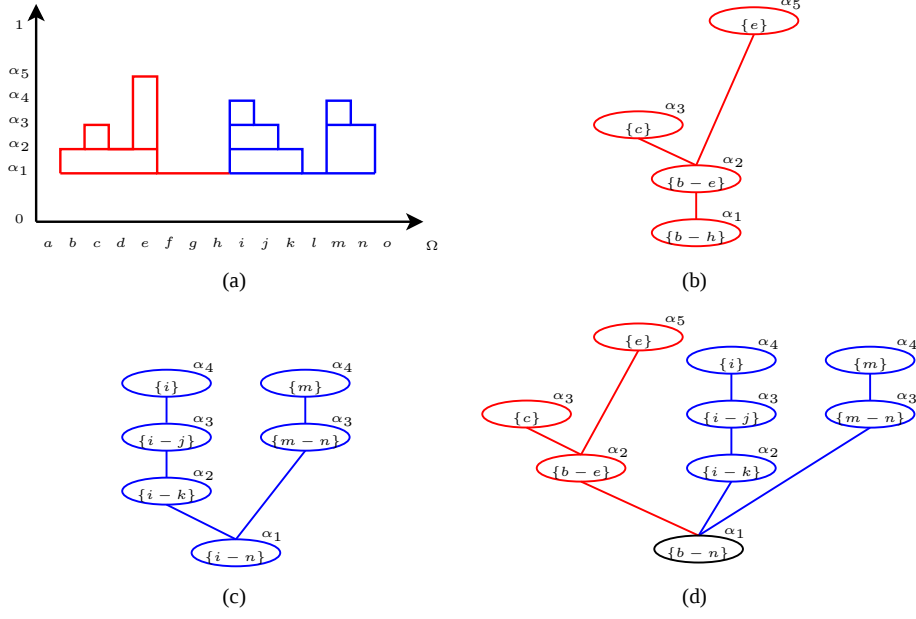


FIG. 3.3 – Opérateur $\bar{M}$ dans le cas 1. (a) Ensembles flous représentés par t_1 (en rouge) et par t_2 (en bleu). (b) t_1 . (c) t_2 . (d) $\bar{M}(t_1, t_2, h, i)$.

Théorème 3.2.1. *Sous certaines conditions, $\bar{M}$ donne un résultat dans $\mathcal{T}^e$:*

$$\begin{aligned} & \forall (t_1, t_2) \in \mathcal{T}^{e^2}, \forall (p_1, p_2) \in \Omega \\ & \forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \quad \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset \\ & \Rightarrow \begin{cases} \bar{M}(t_1, t_2, p_1, p_2) \in \mathcal{T}^e \\ \wedge (\mathcal{A}(\mathcal{R}(\bar{M}(t_1, t_2, p_1, p_2))) = \min(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{A}(\mathcal{R}(t_2)))) \\ \wedge (\mathcal{P}(\mathcal{R}(\bar{M}(t_1, t_2, p_1, p_2))) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2))) \end{cases} \end{aligned}$$

Cette hypothèse qui est assez forte correspond à l'idée que t_1 et t_2 représentent de manière partielle un ensemble flou sur des sous-domaines distincts de Ω . En d'autres termes, cela signifie que nous sommes dans des cas similaires à ceux illustrés aux figures 3.3, 3.4, 3.5 et 3.6.

Démonstration. Voir preuve en Annexe C.2 □

Définition 3.2.2. Deux nœuds $(n_1, n_2) \in \mathcal{V}$ sont dit compatibles dans $(t_1, t_2) \in \mathcal{T}^{e^2}$ si

$$\begin{aligned} & (n_1, n_2) \in \mathcal{N}(t_1) \times \mathcal{N}(t_2) \\ & \wedge \mathcal{P}(n_1) \cup \mathcal{P}(n_2) \text{ connexe} \\ & \wedge \mathcal{A}(n_1) \leq \mathcal{A}(n_2) \Rightarrow \nexists n'_2 \in \mathcal{N}(t_2) / ((\mathcal{A}(n_1) \leq \mathcal{A}(n'_2) < \mathcal{A}(n_2)) \wedge (\mathcal{P}(n_1) \cup \mathcal{P}(n'_2) \text{ connexe})) \\ & \wedge \mathcal{A}(n_2) \leq \mathcal{A}(n_1) \Rightarrow \nexists n'_1 \in \mathcal{N}(t_1) / ((\mathcal{A}(n_2) \leq \mathcal{A}(n'_1) < \mathcal{A}(n_1)) \wedge (\mathcal{P}(n_2) \cup \mathcal{P}(n'_1) \text{ connexe})) \end{aligned}$$

Cette définition permet de dire si deux nœuds peuvent être fusionnés dans le but de représenter partiellement ou complètement une composante connexe d'une α -coupe donnée comme illustré à la figure 3.7.

Théorème 3.2.2. Soit $(t_1, t_2) \in \mathcal{T}^{e^2}$ vérifiant les hypothèses du théorème 3.2.1. Alors $\forall (p_1, p_2) \in \Omega$:

$$\begin{aligned} & \forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \\ & (\mathcal{A}(n_1) \leq \mathcal{A}(n_2)) \wedge (p_1 \in \mathcal{P}(n_1)) \wedge (p_2 \in \mathcal{P}(n_2)) \wedge (n_1 \text{ et } n_2 \text{ compatible dans } t_1 \text{ et } t_2) \\ & \Rightarrow \exists n_3 \in \mathcal{N}(\bar{M}(t_1, t_2, p_1, p_2)) / ((\mathcal{A}(n_3) = \mathcal{A}(n_1)) \wedge (\mathcal{P}(n_3) = \mathcal{P}(n_1) \cup \mathcal{P}(n_2))) \end{aligned}$$

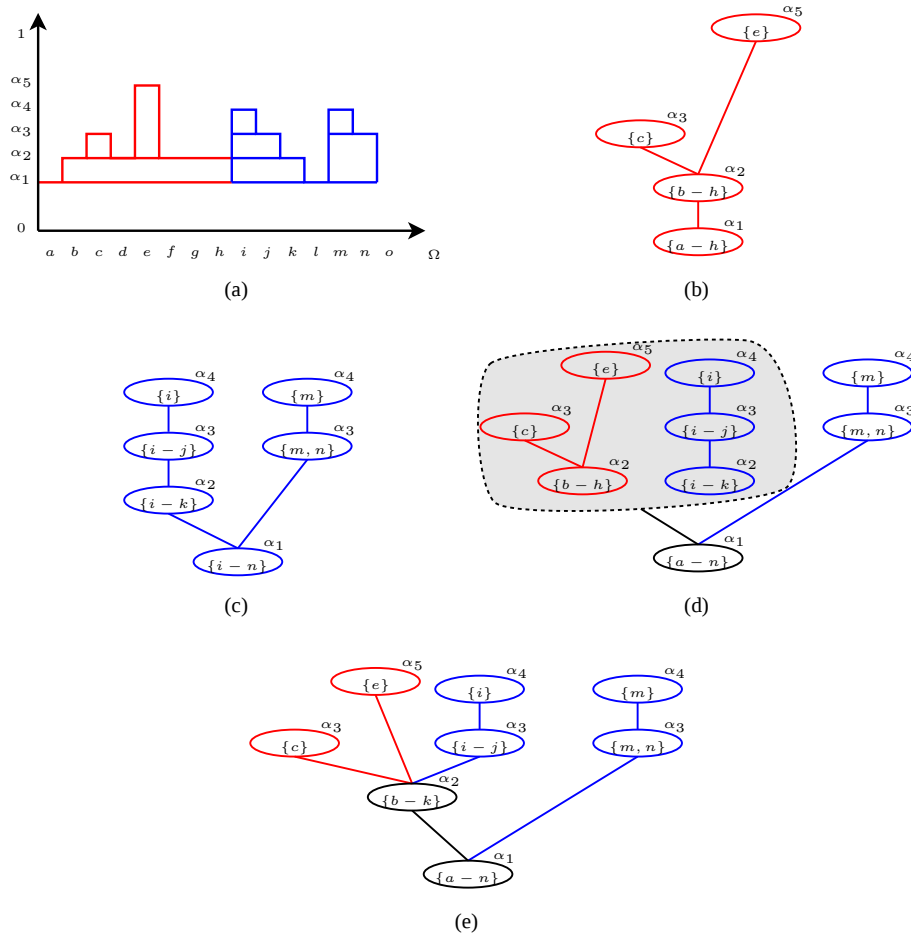


FIG. 3.4 – Opérateur $\overline{M}$ dans le cas 2. (a) Ensembles flous représentés par t_1 (en rouge) et t_2 (en bleu). (b) t_1 . (c) t_2 . (d) $\overline{M}(t_1, t_2, h, i)$ où les arbres dans la zone grise doivent être fusionnés par un nouvel appel à l'opérateur de fusion $\overline{M}$. (e) Arbre final après tous les appels récursifs à $\overline{M}$.

Remarque 3.2.2. Ce théorème peut s'interpréter comme le fait que la fusion de deux arbres qui sont des emboîtements sur des sous-domaines disjoints résulte en un arbre qui contient tous les nœuds qui représentent des composantes connexes correspondant à l'union des points contenus dans les paires de nœuds provenant des deux arbres, qui sont compatibles et qui contiennent respectivement p_1 et p_2 . En d'autres termes, l'opérateur $\overline{M}$ reconnecte de manière locale les composantes connexes des différentes α -coupes. Remarquons enfin que seuls des nœuds compatibles donnent naissance à de nouveaux nœuds après fusion, ainsi les nœuds d'origine sont soit conservés, soit agrandis.

Démonstration. Par définition de $\overline{M}$, et en utilisant les équations 3.11 et 3.12, on obtient directement le théorème 3.2.2. $\square$

Définition 3.2.3. Soit $\text{subst} : \mathcal{T} \times \mathcal{T} \times \mathcal{T} \rightarrow \mathcal{T}$, l'opérateur qui à $(t_1, t_2, t_3) \in \mathcal{T}^3$ associe l'arbre $\text{subst}(t_1, t_2, t_3)$ qui correspond à t_1 dans lequel on a remplacé t_2 par t_3 . Si t_2 n'est pas un sous-arbre maximal de t_1 , $\text{subst}(t_1, t_2, t_3) = t_1$.

La figure 3.8 illustre cette définition dans le cas où t_2 (c.f. figure 3.8(b)) est un sous-arbre maximal de t_1 (c.f. figure 3.8(a))

Introduisons maintenant une méthode pour mettre à jour un arbre t représentant un ensemble f , pour obtenir t^c représentation de f' (version modifiée en un unique point $\tilde{p}$ de f). L'idée est d'ajouter un nœud

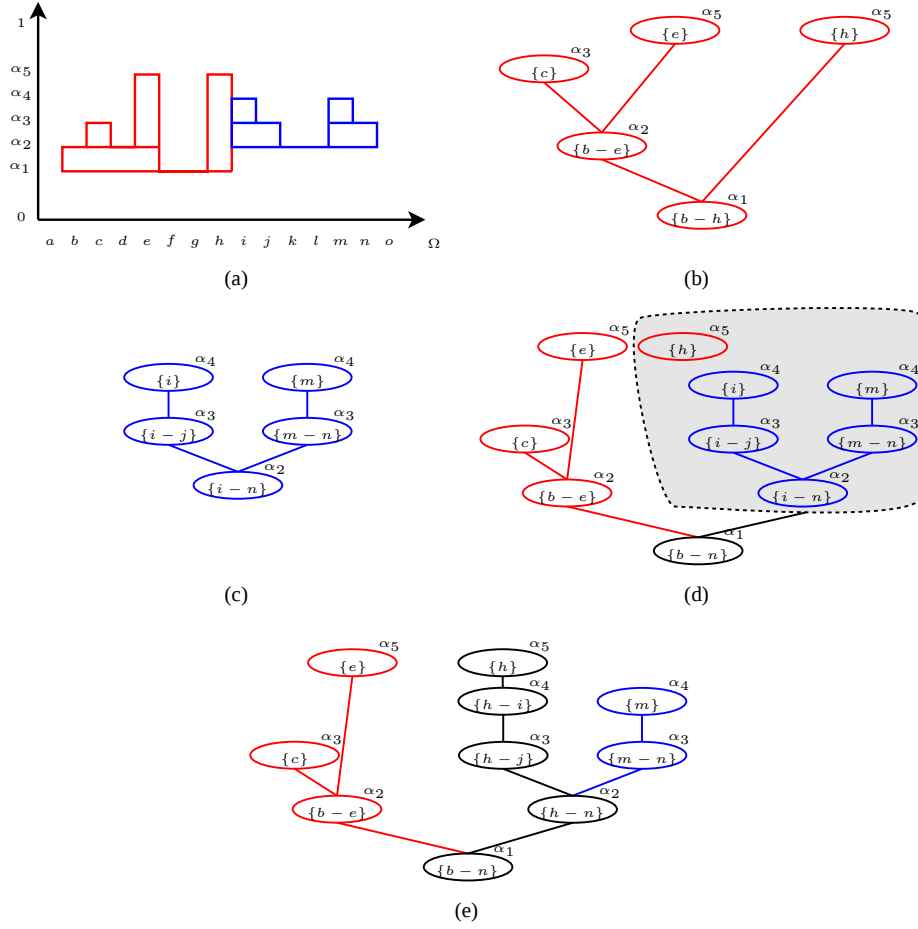


FIG. 3.5 – Opérateur $\overline{M}$ dans le cas 3. (a) Ensembles flous représentés par t_1 (en rouge) et t_2 (en bleu), (b) t_1 . (c) t_2 . (d) $\overline{M}(t_1, t_2, h, i)$ où les arbres dans la zone grise doivent être fusionnés par un nouvel appel à l'opérateur $\overline{M}$. (e) Arbre final après tous les appels récursifs à $\overline{M}$.

filis ne contenant que $\tilde{p}$ avec pour degré d'appartenance $f'(\tilde{p})$ au nœud contenant $\tilde{p}$ avec pour degré d'appartenance $f(\tilde{p})$ dans le but d'appeler un opérateur de fusion sur ce nœud et les nœuds qui contiennent un voisin de $\tilde{p}$. Ce processus est illustré à la figure 3.9.

Définition 3.2.4. Soit $M : \Omega \times 2^\Omega \times \mathcal{T}^e \rightarrow \mathcal{T}^e$ l'opérateur de fusion de nœuds d'un même arbre. M est défini de manière récursive par :

$$\text{cas 1} \quad \forall \tilde{p} \in \Omega, \forall \tilde{t} \in \mathcal{T}^e \quad M(\tilde{p}, \{\}, \tilde{t}) = \tilde{t}$$

$$\text{cas 2} \quad \forall \tilde{p} \in \Omega, \forall P = \{p_1..p_k\} \in 2^\Omega, \forall t \in \mathcal{T}^e \\ M(\tilde{p}, P, t) = \text{subst}(M(\tilde{p}, P \setminus \{p_k\}, t), t', t'') \text{ où } t' \text{ est un sous-arbre de } t^0 = M(\tilde{p}, P \setminus \{p_k\}, t) \text{ tel que :}$$

$$\mathcal{R}(t') = \underset{r \in \mathcal{N}(t^0) / (\tilde{p} \in \mathcal{P}(r)) \wedge (p_k \in \mathcal{P}(r))}{\text{argmax}} \mathcal{A}(r) \quad (3.26)$$

$$\wedge \mathcal{N}(t') = \mathcal{D}'(t^0, \mathcal{R}(t')) \quad (3.27)$$

avec, si on pose $T = \text{sa}(t', \mathcal{R}(t'))$, $T_Q = Q_{\tilde{p}}(T) \cup Q_{p_k}(T)$ et $T_{\overline{Q}} = T \setminus T_Q$, t'' défini de la manière suivante :
 – $t'' = t'$ si $|Q_{\tilde{p}}(T)| = 0 \vee |Q_{p_k}(T)| = 0$,

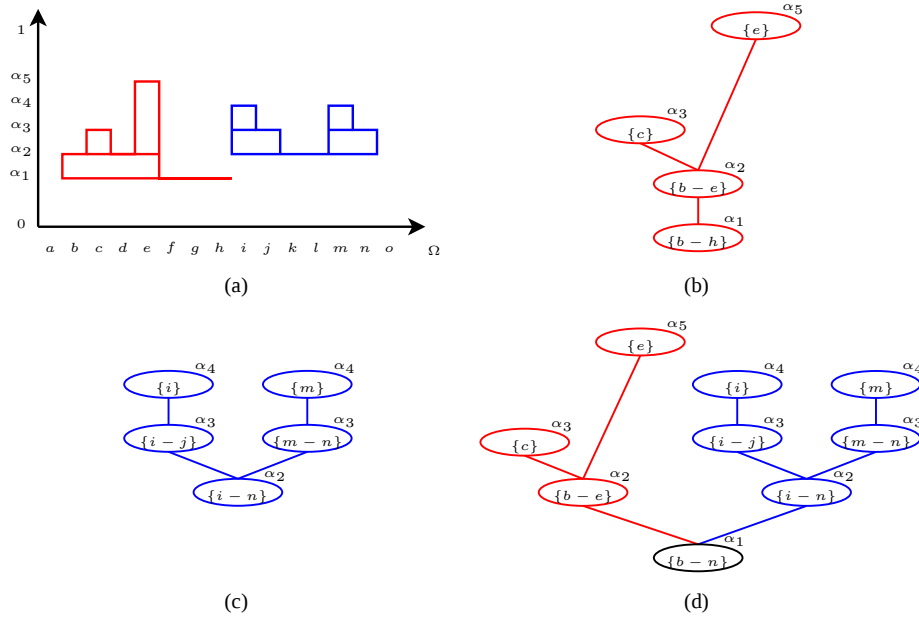


FIG. 3.6 – Opérateur $\bar{M}$ dans le cas 4. (a) Ensembles flous représentés par t_1 (en rouge) et par t_2 (en bleu). (b) t_1 . (c) t_2 . (d) $\bar{M}(t_1, t_2, h, i)$.

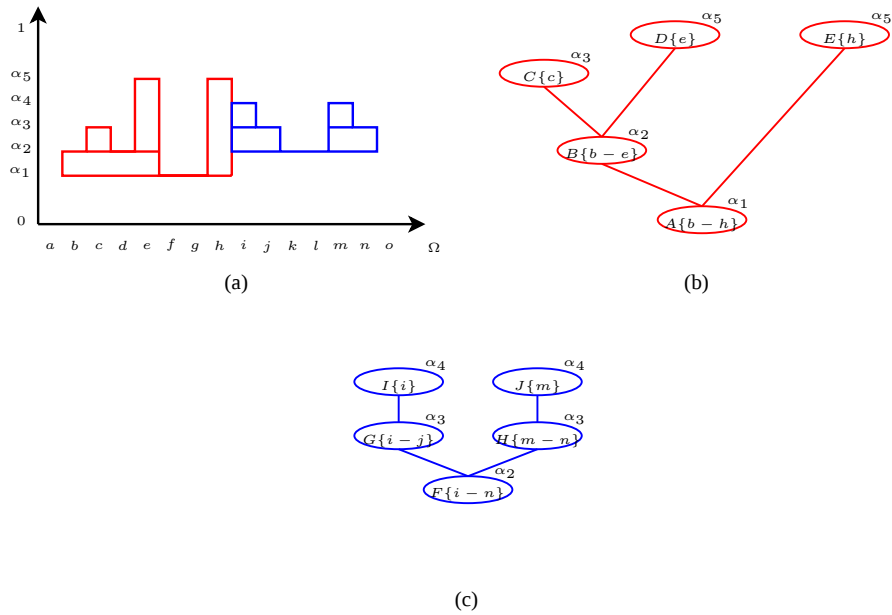


FIG. 3.7 – Exemple de nœuds compatibles dans les arbres t_1 (b) et t_2 (c) qui représentent un ensemble flou (a) sur des parties distinctes du domaine Ω . Seules les paires de nœuds (A, F) , (E, I) , (E, F) et (E, G) sont compatibles. La non-compatibilité est illustrée par les nœuds A et G puisque $A(A) \leq A(F) \leq A(G)$.

– sinon :

$$\mathcal{N}(t'') = \bigcup_{t''' \in T_{\bar{Q}}} \mathcal{N}(t''') \cup \mathcal{N}(\bar{M}(Q_{\bar{p}}(T), Q_{p_k}(T), \bar{p}, p_k)) \cup \{\mathcal{R}(t')\}$$

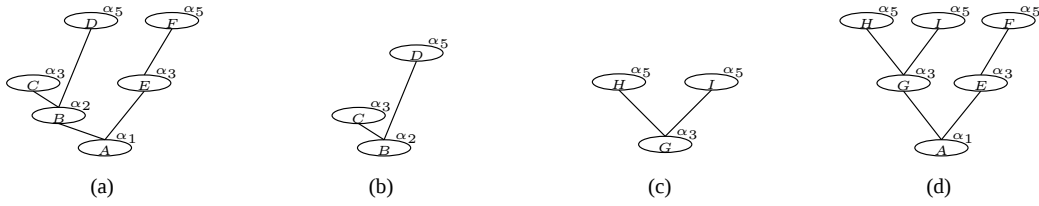


FIG. 3.8 – Exemple de résultat de l'opérateur *subst*. (a), (b), (c) et (d) représentent respectivement t_1 , t_2 , t_3 et $subst(t_1, t_2, t_3)$.

où le premier terme correspond aux nœuds des arbres non touchés, le second aux nœuds des arbres fusionnés et le troisième à la racine qui est inchangée ($\mathcal{R}(t'') = \mathcal{R}(t')$); et :

$$\begin{aligned} \mathcal{B}(t'') = & \bigcup_{t''' \in T_{\bar{Q}}} (\mathcal{R}(t'), \mathcal{R}(t''')) \cup \mathcal{B}(t''') \\ & \cup \mathcal{B}(\bar{M}(Q_{\bar{p}}(T), Q_{p_k}(T), \tilde{p}, p_k)) \cup (\mathcal{R}(t'), \mathcal{R}(\bar{M}(Q_{\bar{p}}(T), Q_{p_k}(T), \tilde{p}, p_k))) \end{aligned}$$

où le premier terme correspond aux arcs des sous-arbres inchangés (avec un arc entre leur racine et la racine de t'), et le second aux arcs des sous-arbres fusionnés.

L'arbre t' évoqué au cas 2 correspond au sous-arbre de t^0 résultant des premiers appels récursifs (équation 3.27) dont la racine contient les deux points servant de marqueurs pour la fusion et qui est tel que ses descendants ne contiennent jamais ces deux points en même temps. Cette dernière propriété peut se traduire par le fait que la racine de t' soit le nœud contenant p_k et $\tilde{p}$ qui a un degré d'appartenance maximal (c.f. équation 3.26) dans la mesure où les descendants d'un nœud d'un emboîtement ont un degré d'appartenance supérieur à ce dernier.

La figure 3.9 illustre les différents appels récursifs à M sur un exemple concret. Ici, l'arbre $\tilde{t}$ est présenté à la figure 3.9(b), avec $P = \{p_1, p_2, p_3, p_4\}$ les voisins de $\tilde{p}$. La fusion aux points $\tilde{p}$ et p_1 de $\tilde{t}$ (c.f. figure 3.9(c)) est fusionnée aux points $\tilde{p}$ et p_2 (c.f. figure 3.9(e)) pour être encore fusionnée aux points $\tilde{p}$ et p_3 (c.f. figure 3.9(g)) et finalement aux points $\tilde{p}$ et p_4 . Le résultat de ces fusions correspond à $M(\tilde{p}, P, \tilde{t})$ comme illustré à la figure 3.9(h). Dans le but d'identifier $t^0, t', t'', \dots$ à différentes étapes de récursion, on peut se référer par exemple aux figures 3.9(d) et 3.9(e) où t^0 est l'arbre de la figure 3.9(d) et t' l'arbre en pointillés rouges. Les sous-arbres $Q_{p_k}(T) = Q_{p_2}(T)$ et $Q_{\tilde{p}}(T)$ apparaissent aussi dans cette dernière figure. Finalement, la figure 3.9(e) correspond à $subst(t^0, t', t'')$ avec t'' l'arbre en pointillés bleus.

Théorème 3.2.3. $\forall t_1, t_2, t_3 \in \mathcal{T}^e \quad \mathcal{R}(t_2) = \mathcal{R}(t_3) \Rightarrow subst(t_1, t_2, t_3) \in \mathcal{T}^e$

Démonstration. Montrons les équations 3.10, 3.11 et 3.12.

Soit $t_1, t_2, t_3 \in \mathcal{T}^e$ tels que $\mathcal{R}(t_1) = \mathcal{R}(t_2)$

Equation 3.10 ($\forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$)

- $t_3 \in \mathcal{T}^e \Rightarrow \forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(t_3, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$
- or $\mathcal{A}(\mathcal{R}(t_2)) = \mathcal{A}(\mathcal{R}(t_3))$ car $\mathcal{R}(t_2) = \mathcal{R}(t_3)$
- ainsi $subst(t_1, t_2, t_3)$ vérifie l'équation 3.10.

Equation 3.11 ($\forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$)

- $t_3 \in \mathcal{T}^e \Rightarrow \forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(t_3, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$
- or $\mathcal{P}(\mathcal{R}(t_2)) = \mathcal{P}(\mathcal{R}(t_3))$ car $\mathcal{R}(t_2) = \mathcal{R}(t_3)$
- ainsi $subst(t_1, t_2, t_3)$ vérifie équation 3.11.

Equation 3.12 ($\forall n \in \mathcal{N}(t) \quad \bigcap_{s \in \mathcal{W}(t, n)} \mathcal{P}(s) = \emptyset$)

- $t_3 \in \mathcal{T}^e \Rightarrow \forall n \in \mathcal{N}(t_3) \quad \bigcup_{s \in \mathcal{W}(t_3, n)} \mathcal{P}(s) = \emptyset$

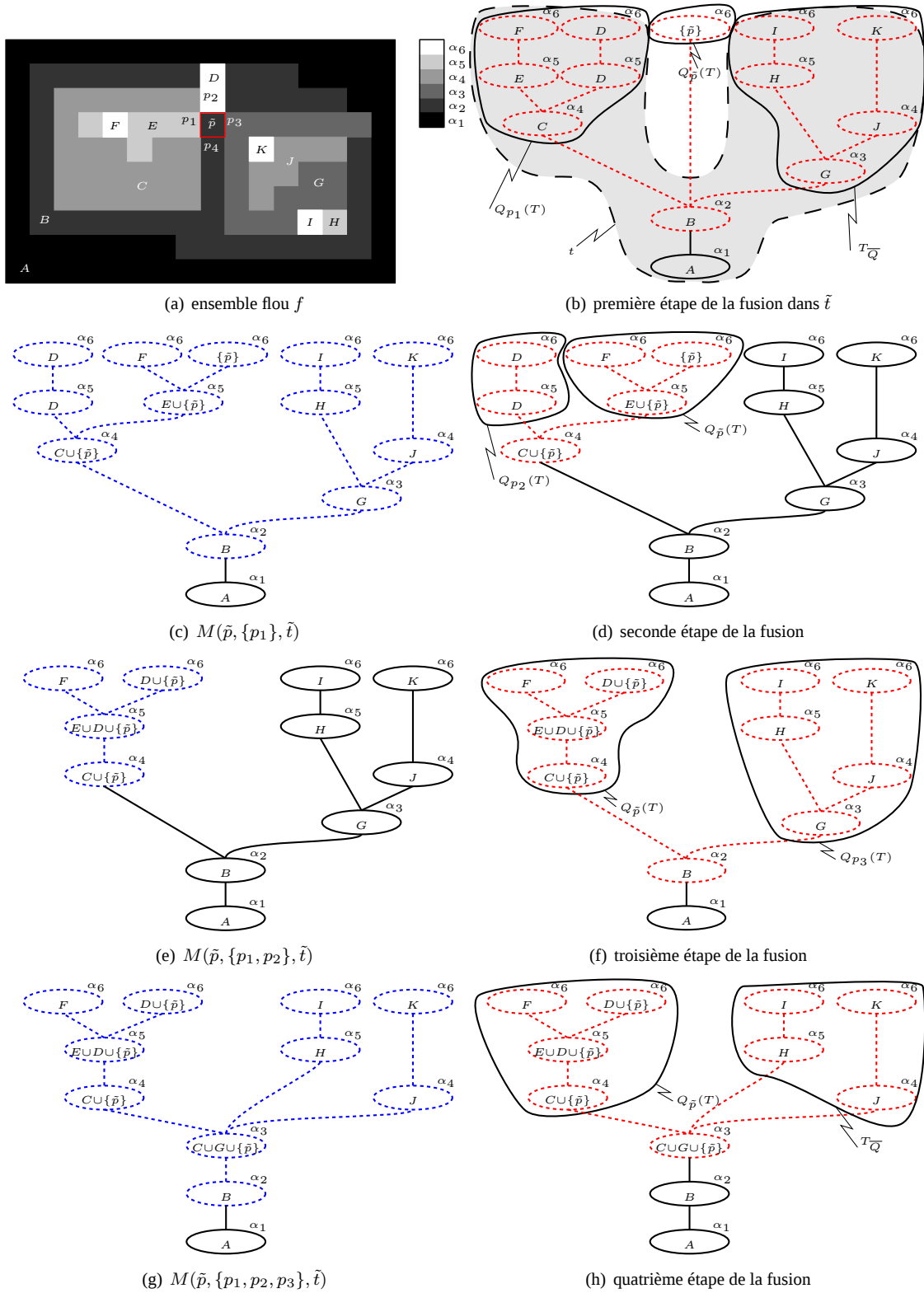


FIG. 3.9 – Utilisation de l'opérateur M sur l'ensemble flou f (a), qui voit sa valeur en $\tilde{p}$ devenir α_6 . (b) Arbre $\tilde{t}$, qui est une représentation de f avec un nouveau nœud contenant $\tilde{p}$, les deux sous-arbres $Q_{\tilde{p}}(T)$ et $Q_{p_1}(T)$ sont sélectionnés pour la première étape de fusion et t' apparaît en pointillés rouges. (c) Résultat de la première étape de fusion avec t'' en pointillés bleus. (d) Extraction de t' (en pointillés rouges), $Q_{\tilde{p}}(T)$ et $Q_{p_1}(T)$ pour la seconde étape de fusion. (e) Résultat de la seconde étape de fusion avec t'' en pointillés bleus. (f) Troisième étape de fusion avec t' en pointillés rouges. (g) Résultat de la troisième étape de fusion avec t'' en pointillés bleus. (h) Dernière étape de fusion avec t' et t'' en pointillés rouges car $|Q_{p_4}(T)| = 0$.

- or $\mathcal{P}(\mathcal{R}(t_2)) = \mathcal{P}(\mathcal{R}(t_3))$ car $\mathcal{R}(t_2) = \mathcal{R}(t_3)$
- ainsi $\text{subst}(t_1, t_2, t_3)$ vérifie l'équation 3.12.

Par conséquence, le théorème 3.2.3 est vérifié. $\square$

Théorème 3.2.4. $\forall \tilde{p} \in \Omega, \forall P \in 2^\Omega, \forall \tilde{t} \in \mathcal{T}^e \quad M(\tilde{p}, P, \tilde{t}) \in \mathcal{T}^e$

Démonstration. Voir annexe C.3. $\square$

Théorème 3.2.5. Soit $\tilde{p} \in \Omega$ et $(f, f') \in \mathcal{S}^2$ tels que :

$$\begin{cases} f(\tilde{p}) < f'(\tilde{p}) \\ \forall p' \in \Omega \quad f(p') = f'(p') \text{ si } p' \neq \tilde{p} \end{cases}$$

Soit $t \in \mathcal{T}^e$ un représentation de f et $\tilde{t}$ l'arbre tel que :

$$\mathcal{N}(\tilde{t}) = \mathcal{N}(t) \cup S_{\tilde{p}}^{f'(\tilde{p})} \quad (3.28)$$

$$\mathcal{B}(\tilde{t}) = \mathcal{B}(t) \cup \left(\underset{n \in \mathcal{N}(t)/\tilde{p} \in \mathcal{P}(n)}{\text{argmax}} \mathcal{A}(n), S_{\tilde{p}}^{f'(\tilde{p})} \right) \quad (3.29)$$

avec S_p^α le nœud qui a un degré d'appartenance égal à α et qui contient seulement p .

Sous ces conditions, $M(\tilde{p}, \text{voisinage}(\tilde{p}), \tilde{t})$ est une représentation de f' (c.f. figure 3.9).

Démonstration. Voir annexe C.4 $\square$

Ce dernier théorème permet d'utiliser l'opérateur M pour mettre à jour un arbre représentant un ensemble f pour obtenir un arbre représentant un ensemble f' qui contient f . Cette preuve repose sur les théorèmes 3.2.2 et 3.2.4.

3.3 Décroissance d'un arbre associé à un ensemble flou

Nous allons maintenant introduire un moyen de faire décroître un ensemble flou en un point. L'idée proposée dans cette section pour mettre à jour l'arbre représentant cet ensemble est de supprimer tous les nœuds qui peuvent ne plus représenter des composantes connexes dans les α -coupes de l'ensemble flou pour ensuite faire grossir l'arbre intermédiaire à l'aide des outils de la section précédente.

Définition 3.3.1. Pour un arbre donné représentant un ensemble flou, et pour un endroit où le degré d'appartenance de l'ensemble voit sa valeur décroître, l'opérateur $\mathcal{Y} : \mathcal{T}^e \times \Omega \times [0, 1] \rightarrow \mathcal{T}^e$ retourne une version élaguée de l'arbre d'entrée :

$$\forall t \in \mathcal{T}^e, \forall \tilde{p} \in \Omega, \forall \alpha \in [0, 1]$$

$$\mathcal{N}(\mathcal{Y}(t, \tilde{p}, \alpha)) = \mathcal{N}(t) \setminus Z_{f^t(\tilde{p}), \alpha}^{\tilde{p}}(\mathcal{N}(t)) \quad (3.30)$$

$$\mathcal{B}(\mathcal{Y}(t, \tilde{p}, \alpha)) = \mathcal{B}(t) \setminus \left\{ (n_1, n_2) \in \mathcal{B}(t) / (n_1 \in Z_{f^t(\tilde{p}), \alpha}^{\tilde{p}}(\mathcal{N}(t))) \vee (n_2 \in Z_{f^t(\tilde{p}), \alpha}^{\tilde{p}}(\mathcal{N}(t))) \right\} \\ \cup \left\{ (n_1, n_2) \in \mathcal{N}(t)^2 / \begin{aligned} & (n_2 \notin Z_{f^t(\tilde{p}), \alpha}^{\tilde{p}}(\mathcal{N}(t))) \\ & \wedge (\exists n_3 \in Z_{f^t(\tilde{p}), \alpha}^{\tilde{p}}(\mathcal{N}(t)) / n_2 \in \mathcal{W}(n_3)) \\ & \wedge \left(n_1 = \underset{n \in \mathcal{N}(t) / (n_2 \in \mathcal{D}(n)) \wedge (n \notin Z_{f^t(\tilde{p}), \alpha}^{\tilde{p}}(\mathcal{N}(t)))}{\text{argmax}} \mathcal{A}(n) \right) \end{aligned} \right\} \quad (3.31)$$

avec f^t l'ensemble flou représenté par t et $\forall p \in \Omega, \forall \alpha_1 \in [0, 1], \forall \alpha_2 \in [0, 1], \forall N \in \mathcal{V} \quad Z_{\alpha_1, \alpha_2}^p(N) = \{n \in N / (p \in \mathcal{P}(n)) \wedge (\mathcal{A}(n) > \alpha_2) \wedge (\mathcal{A}(n) \leq \alpha_1)\}$.

Cet opérateur retourne un arbre dont les nœuds (équation 3.30) sont ceux de l'arbre d'entrée privé de ceux qui contiennent $\tilde{p}$ et qui ont un degré d' α -coupe compris entre l'ancien et le nouveau degré d'appartenance de l'ensemble flou au point $\tilde{p}$ (en pointillés rouges dans la figure 3.10(c)). De manière similaire, les arcs de cet arbre sont ceux de l'arbre d'entrée privés de ceux qui font référence aux nœuds supprimés (deuxième partie de l'équation 3.31, les éléments supprimés apparaissent en pointillés rouge dans la figure 3.10(c)). Pour que le résultat de $\mathcal{Y}$ soit toujours un arbre, des arcs reliant les nœuds isolés par la suppression des précédents arcs et pointant vers leurs descendants directs qui sont préservés sont aussi ajoutés (troisième partie de l'équation 3.31 correspondant aux arcs en pointillés bleus dans la figure 3.10(d)).

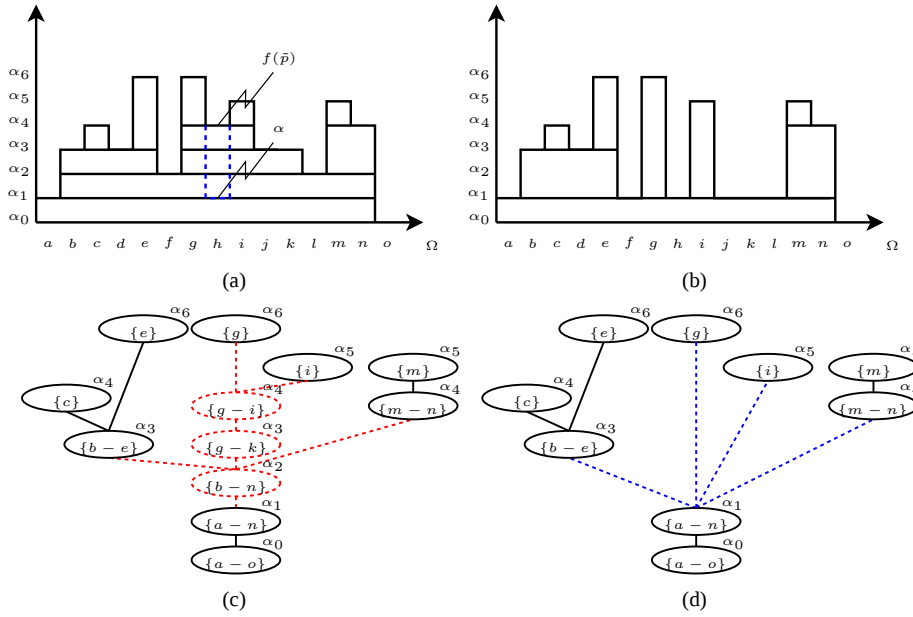


FIG. 3.10 – Fonctionnement de l'opérateur $\mathcal{Y}$ sur l'ensemble (a) au point h . L'arbre t présenté en (b) est la représentation de cet ensemble. L'ensemble (b) correspond au résultat de $\mathcal{Y}(t, h, \alpha)$ avec sa représentation sous forme d'arbre en (d). Les nœuds et les arcs en pointillés rouges dans (c) correspondent aux éléments supprimés respectivement dans l'équation 3.30 et dans l'équation 3.31. Les arcs en pointillés bleus dans (d) correspondent aux arcs ajoutés dans la dernière partie de l'équation 3.31.

Théorème 3.3.1. $\forall t \in \mathcal{T}^e, \forall n \in \mathcal{N}(t), \forall s \in \mathcal{D}(t, n) \quad (\mathcal{P}(s) \subseteq \mathcal{P}(n)) \wedge (\mathcal{A}(s) > \mathcal{A}(n))$

Démonstration. On obtient directement cette propriété en utilisant l'équation 3.10 sur t de manière récursive. $\square$

Théorème 3.3.2. *L'opérateur $\mathcal{Y}(t, \tilde{p}, \alpha)$ retourne un arbre représentant un ensemble flou qui est inclus dans l'ensemble flou (f^t) représenté par l'arbre (t) passé en entrée si le degré α est inférieur à $f^t(\tilde{p})$:*

$$\forall t \in \mathcal{T}^e, \forall \tilde{p} \in \Omega, \forall \alpha \in [0, 1], \quad f^t(\tilde{p}) > \alpha \Rightarrow f^{\mathcal{Y}(t, \tilde{p}, \alpha)} \subset f^t \quad (3.32)$$

Démonstration. Voir la preuve en annexe C.5 qui repose sur le théorème 3.3.1. $\square$

Introduisons maintenant un dernier théorème permettant la mise en œuvre de la suppression de certains nœuds et arcs par l'opérateur $\mathcal{Y}$ en considérant les voisins des points contenus dans les nœuds impactés par le changement de valeur de f au point $\tilde{p}$.

Théorème 3.3.3. $\forall t \in \mathcal{T}^e$ représentation d'un ensemble flou $f, \forall (n_1, n_2) \in \mathcal{B}(t)$

$$\mathcal{P}(n_1) \neq \mathcal{P}(n_2) \Rightarrow \exists (p_1 \in \mathcal{P}(n_1)) \wedge (p_2 \in \mathcal{P}(n_2)) / (f(p_1) = \mathcal{A}(n_1)) \wedge (p_1 \in \text{voisinage}(p_2))$$

Remarque 3.3.1. Ce théorème peut se traduire par le fait que pour tout couple de nœuds présentant un lien de parenté père/fils et qui ne représentent pas la même composante connexe, il existe un couple de points voisins tels que l'un soit contenu dans les deux nœuds et l'autre soit contenu uniquement dans le nœud père, et dans aucun autre nœud correspondant à une α -coupe de degré plus élevé. Ainsi si on regarde les voisins de tous les points p d'un nœud n qui vérifient $\mathcal{A}(n) = f(p)$, on peut lister tous les nœuds dont le père est n .

Démonstration. Soit $t \in \mathcal{T}^e$ une représentation d'un ensemble flou f . Soit $(n_1, n_2) \in \mathcal{B}(t)/\mathcal{P}(n_1) \neq \mathcal{P}(n_2)$

- l'équation 3.11 nous assure que $\mathcal{P}(n_2) \subset \mathcal{P}(n_1)$ car t est un emboîtement,
- l'équation 3.15 nous assure aussi que $\mathcal{P}(n_1)$ est connexe et que $\mathcal{P}(n_2)$ est connexe dans la mesure où t est une représentation,
- il existe donc au moins un couple de points voisins $(p_1, p_2) \in \mathcal{P}(n_1) \times \mathcal{P}(n_2)$ tel que $p_1 \notin \mathcal{P}(n_2)$ dans la mesure où on s'interdit l'inclusion totale de $\mathcal{P}(n_2)$ dans $\mathcal{P}(n_1)$,
- de plus, $\nexists n_3 \in \mathcal{W}(t, n_1)/p_1 \in \mathcal{P}(n_3)$ car sinon n_2 ne serait plus une composante connexe de f et par conséquent t ne serait plus une représentation. On en déduit donc que $f(p_1) = \mathcal{A}(n_1)$.

En conclusion, le théorème 3.3.3 est bien vérifié. $\square$

En utilisant l'opérateur $\mathcal{Y}$, on est capable de produire un arbre t' représentant un ensemble flou $f^{t'}$ qui est inclus dans celui représenté par l'arbre passé en entrée (t). En utilisant les développements de la section 3.2, on peut produire à partir de l'arbre t' un arbre t'' qui représente un sous-ensemble flou $f^{t''}$ qui ne diffère de f^t qu'en un seul point. On peut aussi envisager en pratique de diminuer f^t en plusieurs endroits en même temps, comme on le verra dans la section suivante.

3.4 Application au filtrage d'image de quantités floues

Dans cette section, nous allons proposer des algorithmes résultant des développements précédents. Même si ces algorithmes peuvent être utilisés dans de nombreuses applications, nous les illustrerons dans le contexte du filtrage d'images floues qui ont été définies dans le chapitre 2. Tout d'abord pour faciliter la mise en œuvre des filtres, nous introduirons un moyen de coder les quantités floues contenues dans ces images. Dans un deuxième temps, nous proposerons des algorithmes pour faire croître/décroître des arbres associés à des ensembles flous. Enfin un algorithme générique permettant de mettre en œuvre un grand nombre de filtres connexes flous sera proposé.

3.4.1 Images floues et représentation des niveaux de gris

Comme nous l'avons vu dans le chapitre 2, les images de quantités floues sont des images où la valeur des pixels est représentée par des quantités floues (ensembles flous définis sur le domaine des niveaux de gris $\mathcal{G}$). Un filtre connexe $\delta_\psi : \mathcal{F} \rightarrow \mathcal{F}$ défini pour une telle image floue F peut s'exprimer par un filtre $\psi : \mathcal{S} \rightarrow \mathcal{S}$ appliqué aux ensembles flous définis sur Ω extraits de F pour chaque niveau de gris $g : \forall p \in \Omega, \forall g \in \mathcal{G} \quad \delta_\psi(F)(p, g) = \psi(F^g)(p)$, avec $F^g = F(*, g)$. On rappellera que dans certains cas ψ peut dépendre du niveau de gris et ainsi que du contenu global de l'image F . Pour être considéré comme connexe ψ doit se comporter de manière connexe sur les différentes α -coupes des ensembles passés en entrée. D'un point de vue concret, une représentation sous forme d'arbre est généralement bien adaptée pour mettre en œuvre ψ .

Les valeurs des pixels des images floues peuvent être représentées de manière compacte en ne codant que leurs transitions. Dans la suite de ce document, nous prendrons cette représentation en partant de la plus grande valeur de niveau de gris comme proposé à la figure 3.11.

3.4.2 Croissance d'un ensemble flou

Nous décrivons dans cette sous-section l'algorithme de mise à jour d'un arbre t représentant un ensemble flou f en un arbre t' représentant un ensemble flou f' avec $f < f'$. Pour des raisons de simplicité, nous considérons que f et f' ne diffèrent qu'en un unique point $\tilde{p}$. La figure 3.12 illustre ce cas de figure, ainsi que la méthode proposée pour mettre à jour l'arbre. La première étape est de rajouter dans t un nœud

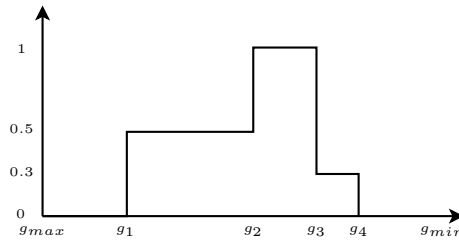


FIG. 3.11 – Nombre flou correspondant à l'intensité en un point représenté de droite à gauche (les plus grands niveaux de gris apparaissent à gauche). Ce nombre est codé par les transitions suivantes : $(0; g_{max})$, $(0.5; g_1)$, $(1; g_2)$, $(0.3; g_3)$, $(0; g_4)$.

ne contenant que le point à modifier dans f (en pointillés sur les figures 3.12(a) et 3.12(b)) avec son nouveau degré d'appartenance (c.f. équations 3.28 et 3.29). Une fois cela fait, l'idée est de fusionner ce nœud avec les nœuds voisins dans l'optique de rétablir la propriété pour chaque nœud de représenter une α -coupe du nouvel ensemble. Cette étape de fusion se fait en utilisant l'opérateur M (c.f. définition 3.2.4) de la manière suivante : pour chaque point p' voisin de $\tilde{p}$, nous cherchons tous les nœuds de t qui contiennent p' et qui ont une valeur d' α -coupe supérieure à l'ancien degré d'appartenance $f(\tilde{p})$. En faisant cela, nous sélectionnons deux sous-arbres ($Q_{\tilde{p}}(T)$ et $Q_{p'}(T)$ comme défini à la définition 3.2.4) associés respectivement aux points $\tilde{p}$ et p' (apparaissent en pointillés et en traits d'union dans les figures 3.12(c), 3.12(d), 3.12(e), 3.12(f)). Une fois les deux sous-arbres sélectionnés, il ne reste plus qu'à les fusionner nœud par nœud en partant de leur racine (c.f. définition 3.2.1 de $\overline{M}$). Cette fusion se fait juste en ajoutant les points provenant des deux sous-arbres pour un niveau α donné : nous rétablissons ainsi partiellement les composantes connexes correspondant aux α -coupes, comme nous pouvons le voir sur les figures 3.12(g) et 3.12(h) (c.f. remarque 3.2.2). Ici, partiellement, signifie que la connexité induite par la relation de voisinage entre $\tilde{p}$ et p' est restaurée. En fait, on est sûr de restaurer complètement la connexité grâce aux itérations du procédé pour tout $p' \in \text{voisinage}(\tilde{p})$.

L'algorithme 3.1 résume le processus que nous venons de décrire. Le théorème 3.2.5 nous assure que cet algorithme produit un arbre t' qui représente f' .

On peut aussi remarquer que le problème de fusion résolu par cet algorithme dans le cadre de l'augmentation de la valeur d'un ensemble flou en un point est similaire au problème de fusion de deux représentations arborescentes d'une image à niveaux de gris sur des sous-domaines différents dans le cadre d'un algorithme de construction parallèle d'arbre proposé par Wilkinson *et al.* (2008). La différence principale dans ce qui est proposé ici est que les nœuds sont stockés dans une pile avant d'être fusionnés pour respecter la définition de $\overline{M}$ alors que dans Wilkinson *et al.* (2008), les nœuds sont fusionnés de manière directe. Néanmoins, même si ces approches résolvent le même problème de fusion d'arbres, ce dernier ne correspond qu'à une sous-partie de ce que nous voulons faire ici : c'est-à-dire corriger la topologie d'un arbre pour refléter les changements correspondant à la mise à jour du degré d'appartenance en un point de l'ensemble flou représenté.

3.4.3 Réduction d'un ensemble flou

Nous présentons dans cette sous-section l'algorithme de mise à jour de l'arbre t (associé à un ensemble flou f) en un arbre t' (associé à un ensemble f') lorsque $f' < f$. Comme dans la section précédente pour des questions de simplicité nous supposons que f' et f ne diffèrent qu'en un point $\tilde{p}$.

Cette deuxième partie de réduction d'un ensemble flou est beaucoup plus complexe que la première, en effet il peut être assez délicat d'enlever un point d'un nœud d'un arbre, dans la mesure où un point peut déconnecter une composante connexe qui devrait ainsi être représentée par deux nœuds distincts comme nous pouvons le voir sur la figure 3.13(a).

Notre approche consiste à essayer de se rapprocher du cas où l'on augmente la valeur de f . Pour se faire il suffit de retirer tous les nœuds qui contiennent $\tilde{p}$ et qui ont une valeur d'appartenance comprise entre $f(\tilde{p})$ et $f'(\tilde{p})$. Les nœuds qui ne satisfont pas ce dernier critère ne pouvant pas être impactés par la variation de

Algorithme 3.1 : Algorithme de mise à jour d'un arbre correspondant à l'augmentation de la valeur d'appartenance d'un point d'un ensemble flou.

```

soit  $f$  l'ensemble flou à modifier;
soit  $t$  l'arbre qui représente l'ensemble flou  $f$ ;
soit  $R^+ = \{(p_i, \alpha_i)\}$  l'ensemble des points  $p_i$  où la fonction  $f$  croît;
tant que  $R^+ \neq \emptyset$  faire
    prendre une paire  $(p, \alpha) \in R^+$ ;
     $R^+ \leftarrow R^+ \setminus \{(p, \alpha)\}$ ;
    créer un nœud  $n_0 | (\mathcal{A}(n_0) = \alpha) \wedge (\mathcal{P}(n_0) = \{p\})$ ;
    ajouter  $n_0$  à la liste des fils des nœuds  $n'$  qui vérifient  $(\mathcal{A}(n') = f(p)) \wedge (p \in \mathcal{P}(n'))$ ;
    pour tous les  $p_i \in \text{voisinage}(p)$  faire
        soit  $n_i \in \mathcal{N}(t) | (p_i \in \mathcal{P}(n_i)) \wedge (\mathcal{A}(n_i) = f(p_i))$ ;
        soit  $S = \emptyset$  une pile de nœuds;
        tant que  $n_i \neq n_0$  faire
            si  $\mathcal{A}(n_i) > \mathcal{A}(n_0)$  alors
                empiler  $n_i$  dans  $S$ ;
                 $n_i \leftarrow \mathcal{D}(n_i)$ ;
            fin
            sinon si  $\mathcal{A}(n_i) < \mathcal{A}(n_0)$  alors
                empiler  $n_0$  dans  $S$ ;
                 $n_0 \leftarrow \mathcal{D}(n_0)$ ;
            fin
            sinon
                empiler  $n_i$  et  $n_0$  dans  $S$ ;
                 $n_i \leftarrow \mathcal{D}(n_i)$ ;
                 $n_0 \leftarrow \mathcal{D}(n_0)$ ;
            fin
        fin
    soit  $r \leftarrow n_0$  la racine commune;
    tant que  $S \neq \emptyset$  faire
         $n \leftarrow$  le haut de la pile  $S$ ;
         $\mathcal{P}(r) \leftarrow \mathcal{P}(r) \cup \mathcal{P}(n)$ ;
        si  $\mathcal{A}(r) < \mathcal{A}(n)$  alors
             $\mathcal{D}(n) \leftarrow r$ ;
             $r \leftarrow n$ ;
        fin
    fin
fin

```

$f(\tilde{p})$, il est inutile de les enlever.

Pour illustrer ce procédé, nous pouvons prendre l'exemple de la figure 3.13 où l'ensemble f voit sa valeur réduite au point g (figure 3.13(a)). La première étape consiste à enlever tous les points qui ont un degré d'appartenance égal à α_5 et qui appartiennent à un nœud n contenant g , c'est-à-dire les points d, g, i comme illustré à la figure 3.13(c). Cette suppression se fait de manière implicite par correction du lien de parenté des nœuds qui contiennent un voisin de ces points et qui pointent sur n (c.f. théorème 3.3.3). Une fois cela fait, le nœud n peut être supprimé en toute sécurité dans la mesure où il n'existe plus aucun nœud dont il est le père. En itérant ce procédé sur les autres nœuds contenant g et dont le degré d'appartenance est supérieur à $f'(g)$, nous obtenons un arbre qui représente un ensemble flou f'' qui est inclus dans f' . Il est alors possible de rajouter les points que l'on a supprimés, en utilisant l'algorithme de la figure 3.1.

Ces différentes étapes sont formalisées dans l'algorithme 3.2.

3.4.4 Filtrage et discussion

Les précédents algorithmes permettent de décrire un meta-algorithme (algorithme 3.3) pour le filtrage des images floues. Ainsi pour filtrer une image F , on peut partir avec un arbre ne contenant qu'un seul

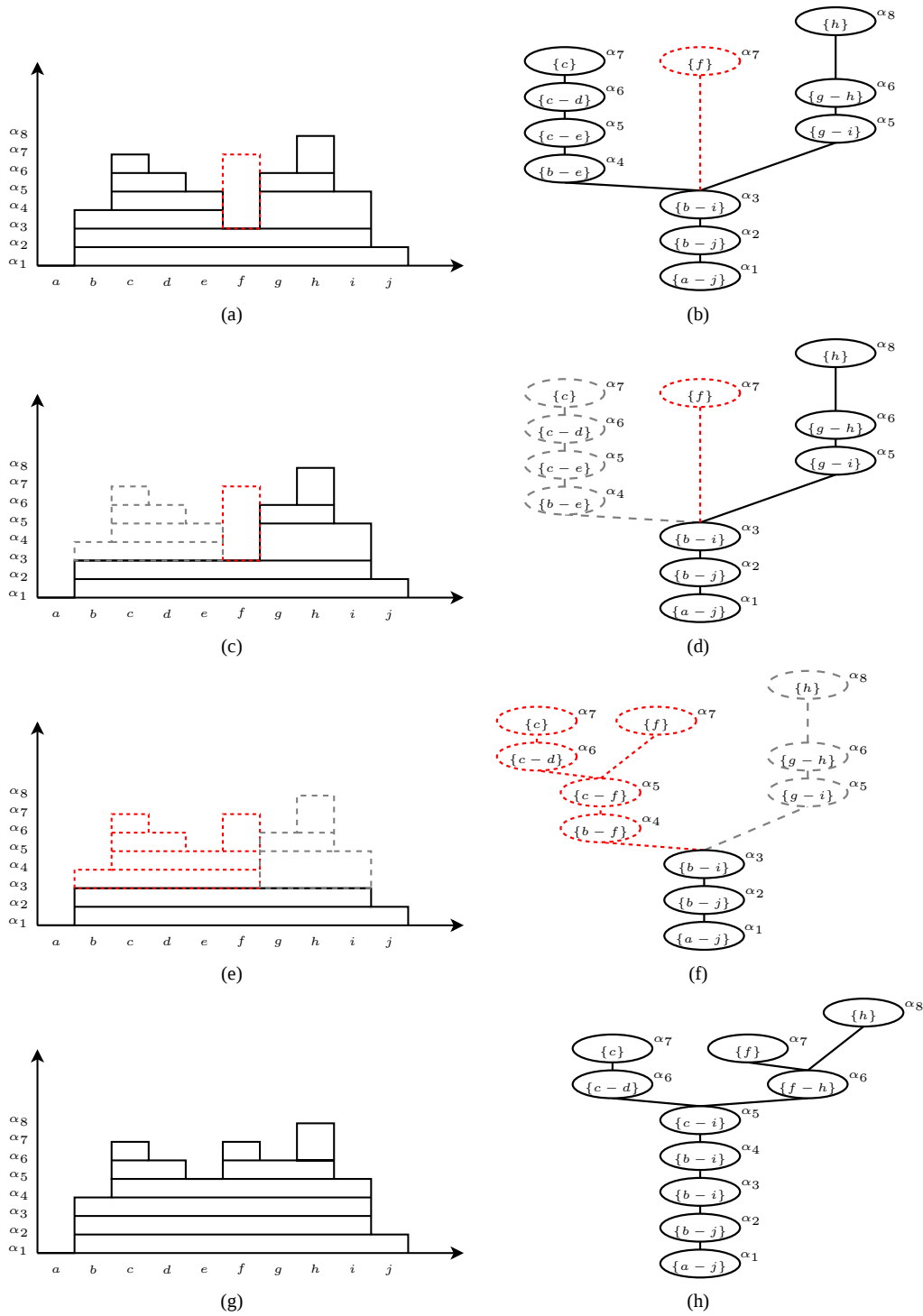


FIG. 3.12 – Exemple d'augmentation de la valeur d'un ensemble flou (a) en un point f . (c), (d) Sélection du sous-arbre contenant f (en pointillés) et du sous-arbre contenant son premier voisin (en traits d'union). (e), (f) Deuxième étape de fusion avec le second voisin de f . (g) Ensemble mis à jour. (h) Arbre mis à jour.

nœud dont le degré d'appartenance est 0 et qui contient tous les points de Ω (arbre en haut à gauche de la figure 3.14), dans le but de le mettre à jour par l'introduction des transitions entre les ensembles flous F^g

Algorithme 3.2 : Algorithme de mise à jour d'un arbre correspondant à la diminution de la valeur d'appartenance d'un point d'un ensemble flou.

```

soit  $f$  l'ensemble flou à modifier;
soit  $t$  l'arbre représentant l'ensemble flou  $f$ ;
soit  $R^- = \{(p_i, \alpha_i)\}$  l'ensemble de points  $p_i$  où  $f$  est censé décroître;
tant que  $R \neq \emptyset$  faire
    prendre un couple  $(p, \alpha) \in R$ ;
    soit  $n_0 = \underset{n \in \mathcal{N}(n) | p \in \mathcal{P}(n)}{\operatorname{argmax}} \mathcal{A}(n)$ ;
    si  $\mathcal{A}(n_0) > \alpha$  alors
        tant que  $\mathcal{A}(n_0) > \alpha, n_0 \leftarrow \mathcal{D}(n_0)$  faire
            pour tous les  $n | (p \in \mathcal{P}(n)) \wedge (\mathcal{A}(n) > \alpha)$ , faire
                pour tous les  $p' \in \mathcal{P}(n) | f(p') = \mathcal{A}(n)$  faire
                     $R^+ \leftarrow R^+ \cup \{(p', f(p'))\}$ ;
                    pour tous les  $p'' \in \text{voisinage}(p')$  faire
                        si  $\exists n' | (p'' \in \mathcal{P}(n)) \wedge (\mathcal{D}(n') = n)$  alors
                             $\mathcal{D}(n') \leftarrow n_0$ ;
                        fin
                    fin
                fin
                supprimer le nœud  $n$ ;
            fin
             $R^- = R^- \setminus (p, \alpha)$ ;
        fin
    fin
    utiliser l'algorithme 3.1 sur  $R^+$ ;
fin

```

et F^{g-1} associés respectivement aux niveaux de gris g et $g - 1$ (arbres à gauche dans la figure 3.14). En utilisant les arbres obtenus, un filtre connexe δ_ψ peut être utilisé (transformation des arbres à gauche vers les arbres à droite à la figure 3.14). Cela correspond généralement à extraire les branches des sous-arbres de l'arbre initial en fonction de la définition de composantes connexes floues que l'on utilise (Rosenfeld, 1979; Braga-Neto et Goutsias, 2003b; Nempont *et al.*, 2008) dans le but de les garder ou de les supprimer complètement ou partiellement.

Algorithme 3.3 : Algorithme de mise à jour d'un ensemble flou en plusieurs points.

```

soit  $f \leftarrow 0_F$ ;
soit  $t \leftarrow \langle 0, X \rangle$  l'arbre qui représente l'ensemble flou  $f$ ;
soit  $R = \{(p_i, \mu_i, g_i)\}$  l'ensemble des transitions de l'image floue (c.f. figure 3.11);
trier  $R$  de manière croissante par rapport à  $g_i$ ;
tant que  $R \neq \emptyset$  faire
    soit  $g$  le niveau de gris associé au premier élément de  $R$ ;
     $R^+ \leftarrow \{(p_i, \alpha_i) | (p_i, \alpha_i, g) \in R \wedge (\max\{\mathcal{A}(n) | p_i \in \mathcal{P}(n)\} > \alpha_i)\}$ ;
     $R^- \leftarrow \{(p_i, \alpha_i) | (p_i, \alpha_i, g) \in R \wedge (\max\{\mathcal{A}(n) | p_i \in \mathcal{P}(n)\} < \alpha_i)\}$ ;
     $R \leftarrow R \setminus R^+ \cup R^-$ ;
    utiliser l'algorithme 3.2 sur  $R^-$ ;
    utiliser l'algorithme 3.1 sur  $R^+$ ;
    filtrer l'arbre  $t$  mis à jour;
fin

```

Le gain de cette stratégie de mise à jour comparé au calcul direct d'un arbre pour chaque niveau de gris dépend du nombre de points à mettre à jour entre deux ensembles flous. En fait, pour un point p donné d'une image F , le nombre de mises à jour est lié à ce à quoi ressemble la quantité floue $F(p, *)$. Ainsi, un nombre flou avec un grand support introduira un grand nombre de transitions, alors qu'une faible quantification des degrés d'appartenance réduira le nombre de transitions. Dans le pire cas, on aura à mettre à jour l'arbre pour

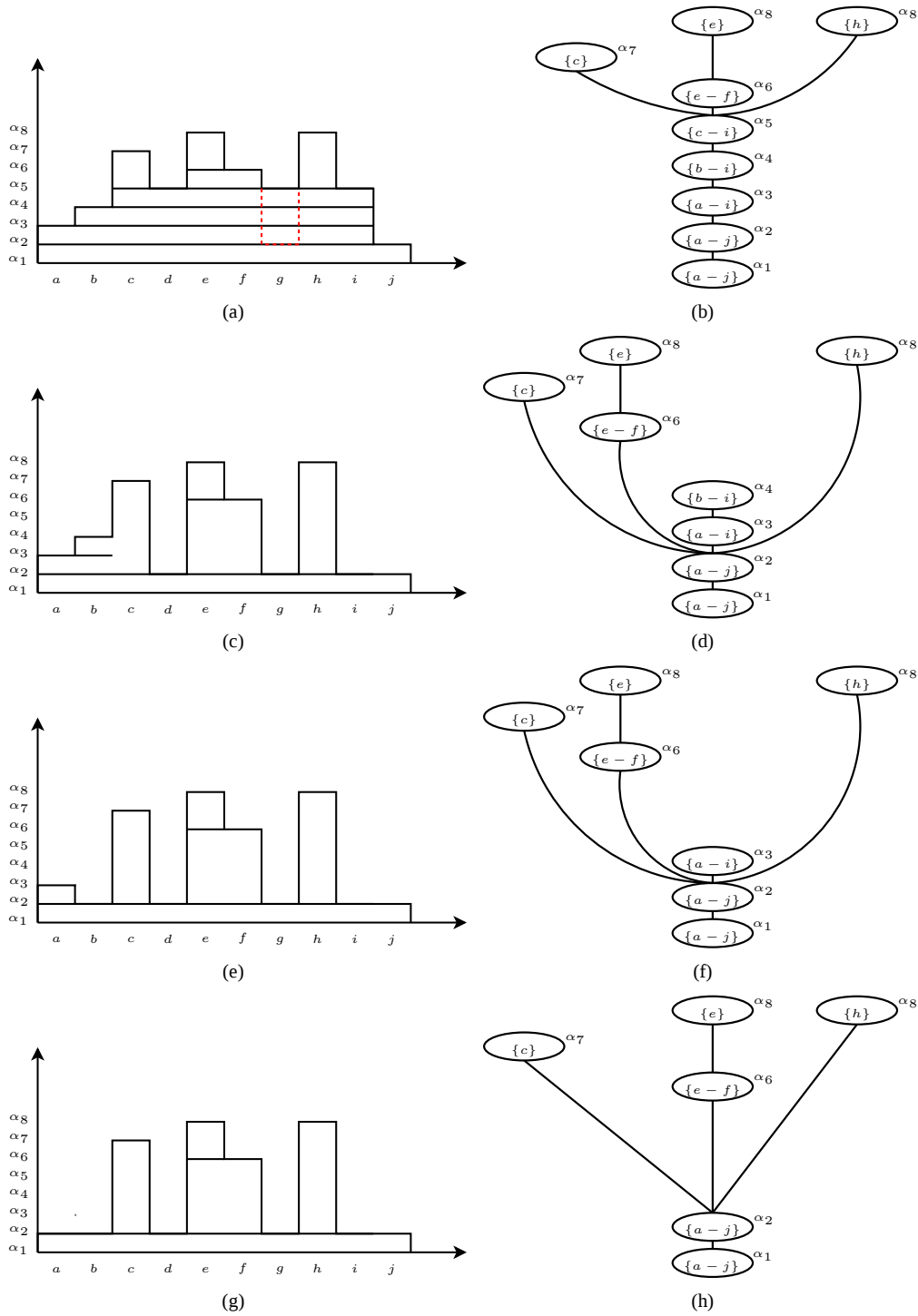


FIG. 3.13 – Exemple de diminution de la valeur d'un ensemble flou (a) en un point g . (c), (d) Première étape de suppression du nœud $\langle \alpha_5, c-i \rangle$. (e), (f) Deuxième étape de suppression du nœud $\langle \alpha_4, b-i \rangle$. (g), (h) Troisième étape de suppression du nœud $\langle \alpha_3, a-i \rangle$. Pour finaliser la procédure, il ne reste qu'à faire croître le dernier ensemble (g) aux points a, b, d, i .

tous les points du domaine de l'image. Néanmoins, un tel cas n'est possible que lorsque l'image contient une très grande quantité de flou, ce qui n'est généralement pas souhaitable dans la mesure où quand tout

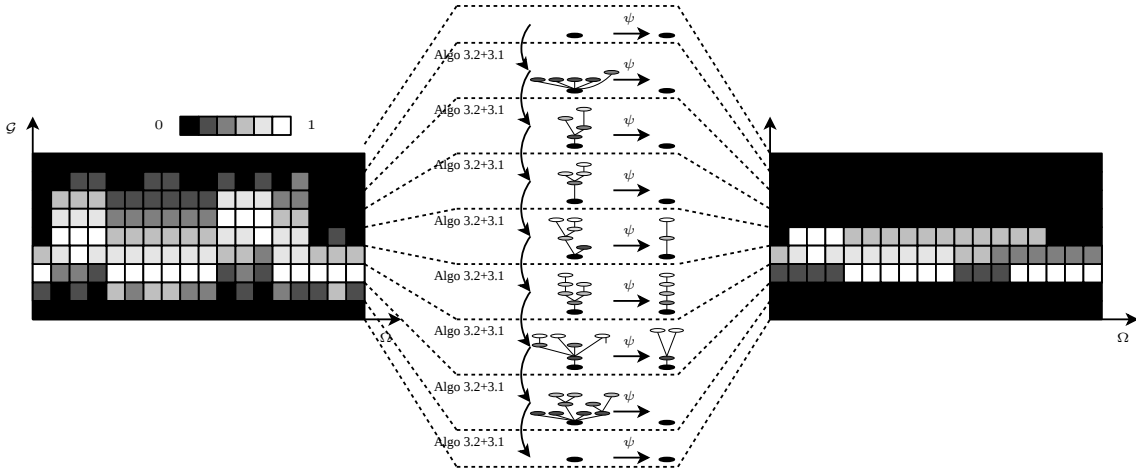


FIG. 3.14 – Filtrage d’une image de nombres flous (image originale sur la gauche et image filtrée sur la droite). Le filtrage de chaque niveau de gris en utilisant un opérateur connexe ψ est fait en utilisant une représentation arborescente des ensembles flous correspondants. Les arbres (à gauche) sont mis à jour de manière itérative en utilisant les algorithmes 3.2 et 3.1 en considérant les différences entre les ensembles flous successifs extraits de l’image originale en partant d’un arbre composé d’un seul nœud représentant Ω (arbre du haut).

devient possible, il est très difficile de décider quoi que ce soit.

Remarquons aussi que les appels à l’algorithme 3.2 dans l’algorithme 3.3 sont inutiles dans le cas où l’image à filtrer est une image d’ombres floues à cause de la propriété de décroissance par rapport aux niveaux de gris (c.f. définition 2.2.1). Grâce à cette propriété, si on encode les niveaux de gris en utilisant l’approche proposée (des plus grands niveaux vers les plus petits), on est assuré que les degrés d’appartenance des ensembles flous extraits de l’image d’ombre floue ne peuvent que croître ou rester constants quand le niveau de gris décroît.

D’un point de vue de la mise en œuvre, la façon dont sont écrits les algorithmes n’impose aucune contrainte sur la représentation informatique des données. Bien sûr, les représentations standard des arbres de coupes peuvent facilement être adaptées à ces algorithmes. Dans nos expérimentations, les nœuds sont étiquetés par des entiers, et la fusion de deux nœuds est faite par un processus de chaînage (un nœud peut pointer vers un autre nœud qui représente la même α -coupe). Les relations père/fils sont exprimées en utilisant des liens entre les représentants des nœuds. Cette structure de données est illustrée à la figure 3.15. Finalement, on peut remarquer que même si les nœuds peuvent être représentés par les points de l’image à la place d’entiers Berger *et al.* (2007), dans notre cas, on ne peut pas utiliser une telle astuce car après la mise à jour d’un arbre, deux nœuds peuvent contenir exactement les mêmes points. Pour éviter cela, on ne devrait utiliser que des représentations compactes qui interdisent de telles configurations.

Enfin, les algorithmes proposés suggèrent de filtrer un arbre après l’étape de mise à jour. Même si cela permet de proposer une approche générique pour mettre en œuvre un filtre, dans certains cas, cela peut ne pas être efficace. Bien sûr, pour des mises en œuvre réelles, les attributs utilisés pour décider quoi supprimer ou garder dans l’arbre à filtrer peuvent être calculés/mis à jour pendant la mise à jour de l’arbre.

3.5 Conclusion

Les arbres de coupes sont couramment utilisés pour représenter des images. Bien qu’ils soient principalement utilisés pour mettre en œuvre des filtres connexes sur des images à niveaux de gris, ils peuvent aussi représenter des ensembles flous permettant par la même occasion de manipuler des composantes connexes floues qui reposent sur une notion de connexité par α -coupes comme celle proposée par Rosenfeld (1979), Braga-Neto et Goutsias (2003b) ou Nempont *et al.* (2008). Comme on a pu le voir dans le chapitre précé-

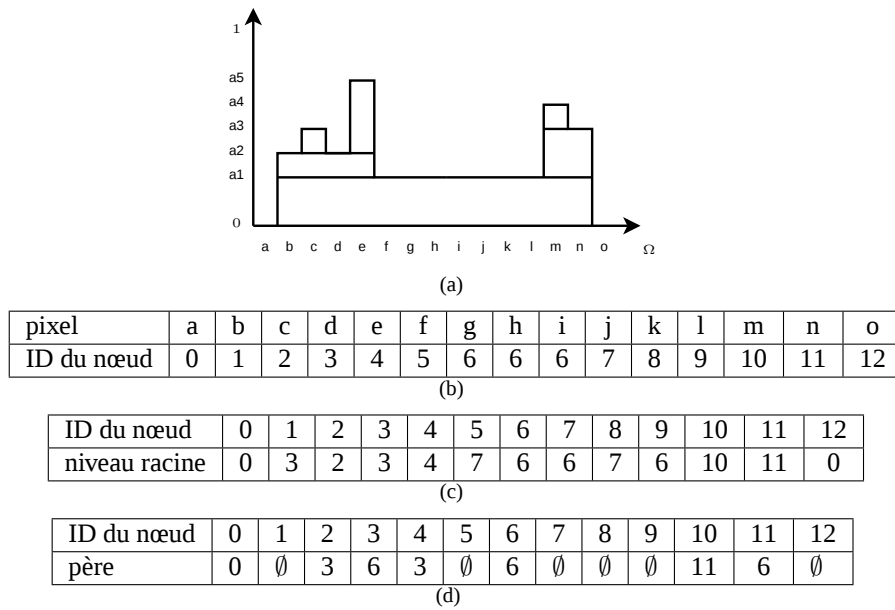


FIG. 3.15 – Structure de données utilisée pour représenter un arbre. (a) Ensemble flou à représenter. (b) Identifiants des nœuds associés aux points dans l'image. (c) Représentants des nœuds ($\text{niveau racine}[n] = n \Rightarrow n$ est un représentant). (d) Parents des nœuds ($\emptyset$ signifie que l'information est inutile puisque le nœud n'est pas un représentant).

dent, une extension des filtres connexes dans le cadre des images floues a été introduite, leur mise en œuvre nécessitant généralement une représentation arborescente pour chaque ensemble flou associé à chaque niveau de gris.

En utilisant l'idée que pour des niveaux de gris successifs, ces ensembles sont très similaires, nous avons introduit des algorithmes efficaces qui permettent de refléter les changements entre deux ensembles dans leurs représentations arborescentes. Ainsi, avec notre approche, si nous avons une représentation sous forme d'arbre d'un ensemble flou, on peut le mettre à jour dans le but de lui faire représenter une version modifiée de l'ensemble original, et ce sans recalculer complètement l'arbre. Cette étape de mise à jour repose principalement sur deux étapes distinctes. La première correspond à la mise à jour d'un arbre pour refléter une augmentation en un point du degré d'appartenance de l'ensemble flou représenté. Cela est fait en ajoutant un nouveau nœud contenant ce point et dont le degré d'appartenance correspond au nouveau degré d'appartenance. Ce nœud est ensuite fusionné avec les nœuds contenant un voisin du point considéré dans le but de restaurer la connexité des α -coupes du nouvel ensemble. La seconde étape consiste à mettre à jour un arbre dans le cas où l'ensemble représenté voit sa valeur réduire en un point. Pour cela, on supprime tous les nœuds qui contiennent ce point et qui ont un degré d'appartenance supérieur à la nouvelle valeur de l'ensemble au point considéré. Cela nous conduit à un arbre représentant un ensemble flous inclus dans l'ensemble flou cible. Ainsi on peut dès lors obtenir ce dernier en utilisant l'approche utilisée à la première étape. En combinant ces deux méthodes, il est possible de faire une mise à jour d'un arbre reflétant n'importe quel changement entre deux ensembles flous.

Même si ce travail a été développé pour filtrer des images floues, d'autres applications peuvent être considérées. Par exemple, des travaux récents Wilkinson *et al.* (2008) utilisent des résultats similaires à ce que nous proposons ici pour fusionner deux sous-arbres, dans un contexte de construction parallèle d'arbres de coupes d'une images. Néanmoins, on peut remarquer que cette problématique de fusion d'arbres ne correspond qu'à une partie de notre schéma de mise à jour d'arbres (sélection des sous-arbres, élagage, etc.). Finalement le formalisme introduit dans ce chapitre nous permet de prouver complètement la validité mathématique de notre approche.

En résumé, dans ce chapitre nous avons proposé une représentation formelle sous forme d'arbres des images floues. Cela nous a permis de développer des algorithmes efficaces permettant le filtrage de ces dernières. Ces algorithmes peuvent être utilisés de manière pratique pour des tâches de filtrage ou de seg-

mentation comme cela sera illustré au chapitre 4, tout en maîtrisant le surplus de complexité de traitement induit par le flou. Ainsi, le filtrage d'images floues devient un nouvel outil utilisable en traitement d'images.

Chapitre 4

Construction d'images floues

Dans les chapitres précédents nous avons vu que les filtres connexes sont fortement utilisés en traitement d'images. Ce type de filtres reposant sur la notion de filtrage de composantes connexes binaires extraites à partir de l'image à filtrer (soit par seuillage, soit par décomposition en zones plates), les dégradations que subit cette dernière peuvent avoir un impact négatif sur leur efficacité. Pour cette raison, les images floues ont été introduites. Ces dernières permettent de représenter l'imprécision que l'on peut avoir sur les niveaux de gris dans l'image. Grâce au cadre théorique dans lequel elles sont définies, on peut étendre la définition de filtres connexes usuels ou encore en définir de nouveaux dédiés à des applications spécifiques comme de la détection. Nous avons aussi abordé le problème de leur mise en œuvre de manière efficace dans le chapitre précédent, cependant, la question du conditionnement de l'imprécision pour construire de telles images n'a pas encore été abordée. Cette étape est bien évidemment essentielle dans le but de fournir un outil de travail complet qui permette de manipuler une image de A à Z. De plus de la qualité de la construction de ces images va dépendre la capacité à effectuer des traitements corrects par la suite.

Dans ce chapitre, nous allons discuter des approches plus ou moins évoluées que l'on peut utiliser pour construire des images floues. Dans un premier temps nous détaillerons les méthodes simplistes de construction évoquées jusqu'ici et détaillerons leurs faiblesses. Dans un second temps nous introduirons le concept clé de conversion de la composante de bruit contenue généralement dans les images en une imprécision de débruitage. Enfin, nous détaillerons l'extension d'une méthode de débruitage reposant sur la décomposition en ondelettes pour générer des images floues et nous étudierons son impact dans le cadre de la détection d'objets isodenses.

4.1 Méthodes générales de construction

Dans cette section nous allons discuter des méthodes de construction d'images floues utilisées jusqu'ici pour mettre en évidence leurs limitations. Nous commencerons par la construction d'images d'ombres floues par l'opérateur um , puis nous parlerons du déploiement d'un patron autour des niveaux de gris d'une image pour obtenir une image d'intervalles flous

4.1.1 Construction d'images d'ombres floues nettes

Comme on a pu le voir dans le chapitre 2, des images d'ombres floues peuvent être construites en utilisant juste la définition classique d'images d'ombres par l'intermédiaire de l'opérateur um . Néanmoins ce genre d'approche produit une image ne comportant aucune imprécision, et donc ne permet pas de profiter pleinement du cadre théorique des filtres connexes flous. Cependant, comme on l'a vu dans le filtrage de coupes de volumes de tomosynthèse du sein, les filtres connexes flous permettent de manipuler de l'imprécision sur les données en entrée, mais aussi sur la réponse du filtre. Ainsi une composante connexe, même non floue, extraite d'un niveau de gris d'une image floue, qui ne satisfait pas pleinement les critères d'un objet que l'on recherche peut survivre partiellement à l'étape de filtrage permettant ainsi de manipuler les cas limites.

4.1.2 Construction en utilisant un patron

Définition

Dans le cas des images d'intervalles flous et des images de nombres flous, la précédente approche reposant sur l'opérateur um n'est pas adaptée dans la mesure où elle ne permet pas de créer ce type d'images. Comme on l'a proposé dans le chapitre 2, l'approche la plus simple est de considérer un patron que l'on va déployer sur chaque niveau de gris de l'image nette $I \in \mathcal{I}$. Dans ce cas, l'imprécision est modélisée par une fonction d'appartenance $ng : \mathbb{R}^+ \rightarrow [0, 1]$, qui vérifie $ng(0) = 1$ et qui est décroissante. Cette fonction a pour but d'évaluer à partir d'une différence de niveaux de gris à quel point ils sont semblables. De manière plus formelle, une image floue $F \in \mathcal{F}$ sera construite à partir d'une image $I \in \mathcal{I}$ et d'une fonction ng en utilisant l'équation suivante :

$$\forall p \in \Omega, \forall g \in \mathcal{G} \quad F(p, g) = ng(|I(p) - g|)$$

Exemples de filtrage en utilisant un patron

Ce type de construction, bien que simpliste, permet déjà de tirer parti des avantages des filtres connexes flous comme on a pu le voir pour l'extraction de zones plates correspondant à des lésions dans le chapitre 2. Néanmoins, en étudiant de manière plus précise le comportement de cette méthode de construction d'images floues sur le résultat d'un filtrage, on peut s'apercevoir que cette méthode possède quelques lacunes. On se propose d'illustrer ces dernières sur une image synthétique. Dans cette expérience, la fonction ng sera définie de la manière suivante :

$$\forall x \in \mathbb{R}^+ \quad ng(x) = \max(0, \min(1, -ax + b))$$

avec a et b deux paramètres permettant de jouer sur la quantité de flou à introduire comme illustré dans la figure 4.1(c).

On remarquera qu'une telle fonction ng produit une image floue F où la valeur de chaque pixel correspond à un trapèze centré autour du niveau de gris associé au même point dans l'image nette (c.f. figure 4.1). De même, les paramètres a et b permettent de modifier la forme de ce trapèze.

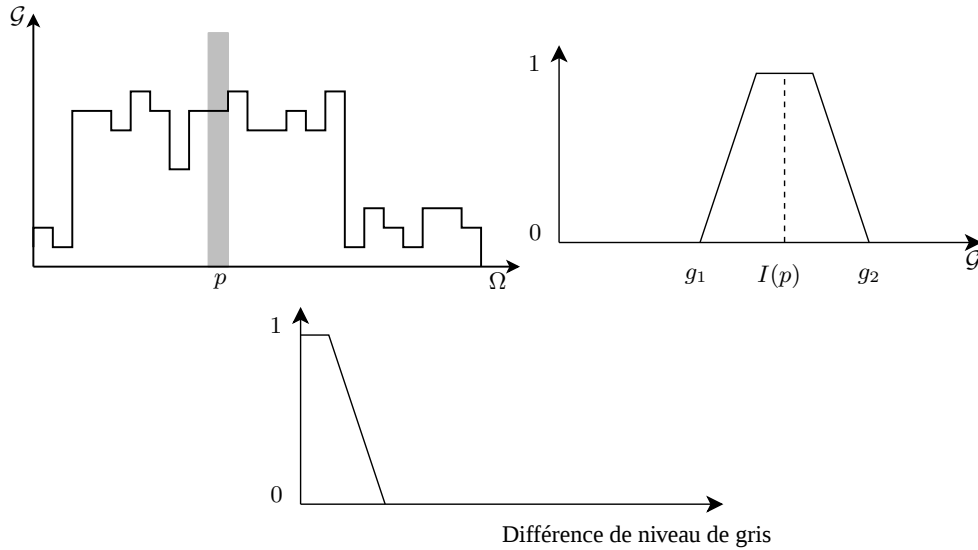


FIG. 4.1 – Construction d'une image d'intervalles flous en déployant un patron sur chaque niveau de gris (b) de l'image originale I (a). Le processus est détaillé en (b) pour le pixel en gris dans (a) avec $ng(g_1 - I(p)) = 0$ et $ng(I(p) - g_2) = 0$. La fonction ng utilisée est illustrée en (c).

L'image synthétique est une image constante par morceaux (c.f. figure 4.2(a)) composée de neuf disques de différents niveaux de gris. A titre d'exemple, cette image est corrompue par un bruit gaussien (c.f.

figure 4.2(b)) afin d'illustrer le comportement de notre approche par rapport au contraste et au bruit. Nous construisons à partir de cette image bruitée une image d'intervalles flous (les figures 4.2(d), 4.2(e) et 4.2(f) représentent différentes coupes de celle-ci) que nous filtrons par l'opérateur défini comme :

$$\delta(F)(*, g) = \bigvee \{h \in \mathcal{H}(F(*, g)) / (c_1 < \text{Card}(h) < c_2)\}. \quad (4.1)$$

avec Card l'aire d'un ensemble flou (Rosenfeld et Haber, 1985), c_1 et c_2 étant deux constantes qui entourent l'aire des disques dans l'image originale.

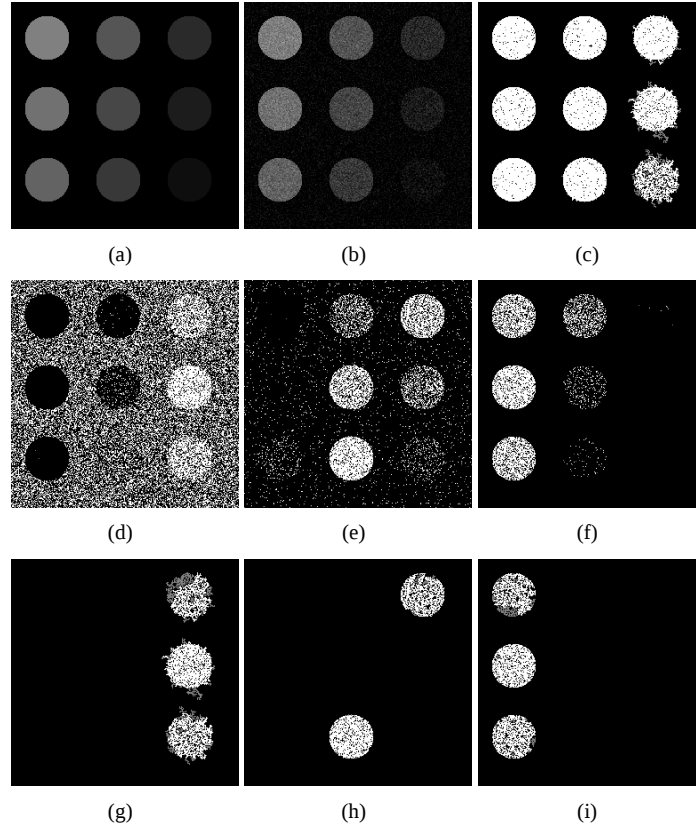


FIG. 4.2 – Exemple d'extraction de zones plates en fonction d'un critère d'aire sur l'image (b), version bruitée de (a). Les images (d), (e) et (f) représentent les coupes $F(*, 10)$, $F(*, 20)$ et $F(*, 40)$ de l'image floue avant filtrage. Les images (g), (h) et (i) représentent ces mêmes coupes après filtrage. L'image (c) représente l'agrégation du résultat du filtrage de tous les $F(*, g)$, $g \in \mathcal{G}$ en utilisant l'opérateur agg .

Le but est donc de savoir si les disques présents dans l'image non bruitée sont détectables dans l'image floue construite à partir de l'image bruitée. Les figures 4.2(g), 4.2(h) et 4.2(i) présentent les résultats du filtrage des coupes précédentes (figures 4.2(d), 4.2(e), 4.2(f)) de l'image d'intervalles flous. Nous voyons sur les différentes coupes que les composantes connexes floues ayant une aire trop grande ou trop petite sont supprimées. Pour un niveau de gris donné, ces structures supprimées proviennent principalement du bruit et de l'étalement du fond ou des disques situés à d'autres niveaux de gris. La figure 4.2(c) représente l'agrégation floue par l'opérateur agg le long des niveaux de gris en chaque pixel de l'image floue filtrée. Nous pouvons ainsi apprécier la qualité de la détection des zones plates satisfaisant le critère. Nous voyons que le disque le moins contrasté (en bas à droite dans la figure 4.2(a)) ressort quand même du fond, même s'il n'est pas très bien défini. Nous remarquons aussi que les zones plates obtenues contiennent des trous, et ce indépendamment du contraste. Cela s'explique par la modélisation de l'imprécision sur les niveaux de gris : le bruit gaussien a un support infini, alors que notre trapèze lui a un support fini. Nous pourrions pallier ce problème en faisant dépendre l'imprécision d'un pixel des valeurs de son voisinage.

Discussion

Bien que cette première façon de construire des images d'intervalles floue permette d'obtenir des résultats corrects en terme de détectabilité comme on a pu le voir sur une image synthétique corrompue avec un bruit Gaussien, le résultat obtenu n'est pas parfait. En effet les structures extraites contiennent des trous, la méthode ne semble donc pas adaptée à ce type de dégradation. Cette conclusion est tout aussi valable pour d'autres types de bruits ou d'autres modèles de dégradation. Dans la mesure où la chaîne de formation d'image n'est pas modélisée par cette méthode de construction, ce type de comportement certes assez intuitif ne permet que d'obtenir un résultat de qualité moyenne en réglant au cas par cas les paramètres de la fonction ng .

4.2 Construction d'image et débruitage

Le manque de modélisation n'est pas le seul problème de l'approche précédente lorsque l'on désire considérer des images qui contiennent du bruit statistique. En effet, il n'est pas possible d'utiliser un ensemble flou pour représenter ce type de variation. Cela a pour conséquence de rendre les images floues inaptes pour ce type d'applications. Une solution pour pallier ce problème serait de modifier la définition de ces images, en utilisant par exemple des variables aléatoires floues (Lopez-Diaz et Gil, 1997), pour faire en sorte qu'elles soient capables de représenter ce type de dégradation. Cependant, la structure des données serait encore plus complexe et donc très difficile à manipuler. Une autre solution peut être de tenter de supprimer la composante correspondant au bruit. Après discussion de ce principe, nous allons présenter une manière simple de le mettre en œuvre en utilisant des filtres de débruitage classiques et en illustrant les résultats que l'on peut obtenir sur la même image synthétique que précédemment. Enfin nous proposerons une approche générale pour construire une image floue à partir d'un filtre de débruitage quelconque.

4.2.1 Débruitage d'images et construction d'image floues

Dans le but de se débarrasser du bruit contenu dans les images nettes, de nombreuses méthodes de débruitage existent. Malheureusement, ces dernières dépendent souvent d'un paramètre et ne produisent pas un résultat parfait (on n'obtient pas l'image originale sans bruit). La plupart du temps, on obtient un compromis entre la suppression du bruit et la conservation de l'information présente dans l'image. Généralement, les paramètres sont fixés de telle manière à ce qu'ils optimisent un critère qui modélise ce compromis. Bien sûr, différents critères donnent généralement différents paramètres avec différentes propriétés. Par exemple on peut vouloir être sûr de conserver les structures faiblement contrastées en contrepartie d'une réduction sommaire du bruit dans l'image. Au contraire, on peut vouloir une image sans bruit et ce même si des structures disparaissent. En définissant de tels critères à cette étape, on prend de manière implicite une décision qui peut donc résulter en une perte d'information qui peut se révéler problématique sur des traitements qui doivent être faits en aval. Pour illustrer cela, on peut prendre l'exemple d'une image dans laquelle on recherche des zones iso-denses faiblement contrastées. Une étape de débruitage peut faire disparaître complètement ces zones et rendre leur détection par un filtre (connexe ou autre) impossible.

Pour éviter ce genre de problèmes, on peut essayer de garder le plus d'information possible dans le but de prendre une décision plus tard (au moment de la détection par exemple). De manière concrète, on peut essayer de reporter l'imprécision des paramètres du filtre de débruitage pour obtenir une imprécision sur les niveaux de gris débruités. Dans ce contexte, on remarquera qu'une image floue est tout à fait adaptée pour gérer ce genre d'information. De plus comme on l'a vu dans les chapitres précédents, un cadre théorique permet de définir des filtres permettant de tirer parti de l'imprécision de ces images pour prendre des décisions.

4.2.2 Construction d'image floues à partir de filtres de rang

Principe

Pour illustrer le principe de construction d'une image floue à partir d'une méthode de débruitage, on peut prendre l'exemple de filtres de rang. L'idée est de combiner plusieurs filtres de rang pour obtenir un

intervalle flou représentant les niveaux de gris possibles après débruitage. Pour l'exemple nous considérons les trois filtres de rang suivants : la dilatation, l'érosion et le filtre médian, en faisant l'hypothèse que les niveaux de gris obtenus après filtrage d'une image $I \in \mathcal{I}$ par les deux premiers filtres correspondent aux valeurs de niveau de gris qui ne sont pas possibles et que le résultat du filtre médian donne les valeurs de niveaux de gris les plus possibles. De manière intuitive, on peut justifier ce choix en considérant que la dilatation et l'érosion, qui sont généralement de mauvais filtres en terme de débruitage, fournissent des bornes pour les valeurs possibles de niveau de gris alors que le filtre médian a un meilleur comportement en terme de régularisation et donc de débruitage. On peut donc construire une *INF* en associant pour chaque pixel p de l'image, le nombre flou $F(p, *)$ défini comme un triangle qui vaut 1 pour la valeur de niveau de gris obtenu par le filtre médian, 0 pour les valeurs obtenues par érosion et dilatation et des degrés intermédiaires entre ces dernières comme illustré à la figure 4.3.

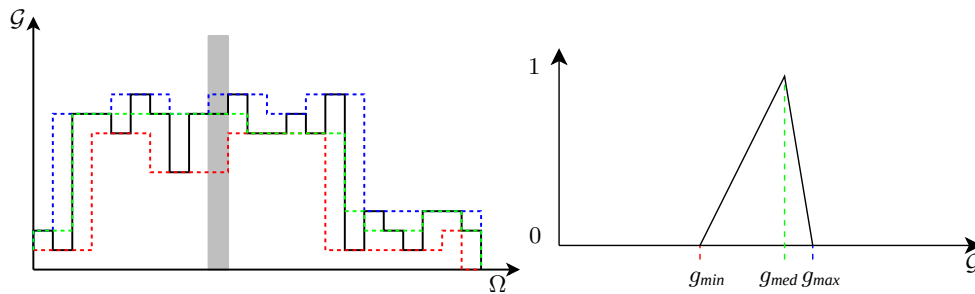


FIG. 4.3 – Etape de construction de nombres flous à partir de filtres de rang. (a) Image 1D originale en noir filtrée par une érosion (en rouge), une dilatation (en bleu) et un filtre médian (en vert). (b) Construction d'un niveau de gris flou pour le point marqué en gris dans (a) à partir des résultats des trois filtrages.

Dans ce type de procédé, la composante statistique est supprimée par l'étape de filtrage, et l'imprécision à propos de ce filtrage est introduite sur les niveaux de gris dans l'image floue.

Exemple de filtrage

Pour valider l'approche, on peut reprendre l'exemple de détection de la section précédente comme illustré à la figure 4.4. Dans cet exemple, l'image floue est construite à partir des trois filtres de rang utilisés précédemment (min, max et médian). Les figures 4.4(d), 4.4(e) et 4.4(f) présentent trois coupes associées aux niveaux de gris 10, 20 et 40 de l'image ainsi construite. On remarquera tout d'abord que pour chaque disque on retrouve la marque du contour pour tous les niveaux de gris compris entre celui du fond et celui du disque considéré. Cela s'explique par le fait que l'étape de filtrage par l'érosion et la dilatation vont introduire une grande imprécision au niveau des contours : le minimum dans le voisinage va venir du fond alors que le maximum va venir du disque. On remarquera aussi l'aspect légèrement morcelé de l'image qui est dû au filtrage médian. Enfin, on remarquera qu'il n'y a plus de trous dans les zones plates détectées comme on peut le voir à la figure 4.4(c) : ici la représentation par une fonction d'appartenance en forme de triangle des niveaux de gris est suffisante car il n'y a plus l'aspect du support infini du bruit gaussien (ce dernier est en partie éliminé par les filtrages).

4.2.3 Application du principe d'extension pour construire des images floues

On peut continuer sur l'idée précédente sur la transformation de la composante de bruit en imprécision de débruitage. En effet si on reprend le raisonnement de manière générale en considérant un filtre dépendant d'un paramètre quelconque, en supposant que l'on est capable de formuler l'imprécision sur notre critère de débruitage en définissant une fonction d'appartenance sur ce paramètre, on peut utiliser le principe d'extension (Zadeh, 1975) pour générer une image floue.

On peut définir une telle approche de manière générique. Soit $I \in \mathcal{I}$ une image nette, Υ_γ un filtre de débruitage avec γ un paramètre permettant d'influer sur le comportement de ce filtre. Si on est capable d'associer un degré d'appartenance μ_γ à chaque valeur de γ , on est capable d'associer à chaque niveau de

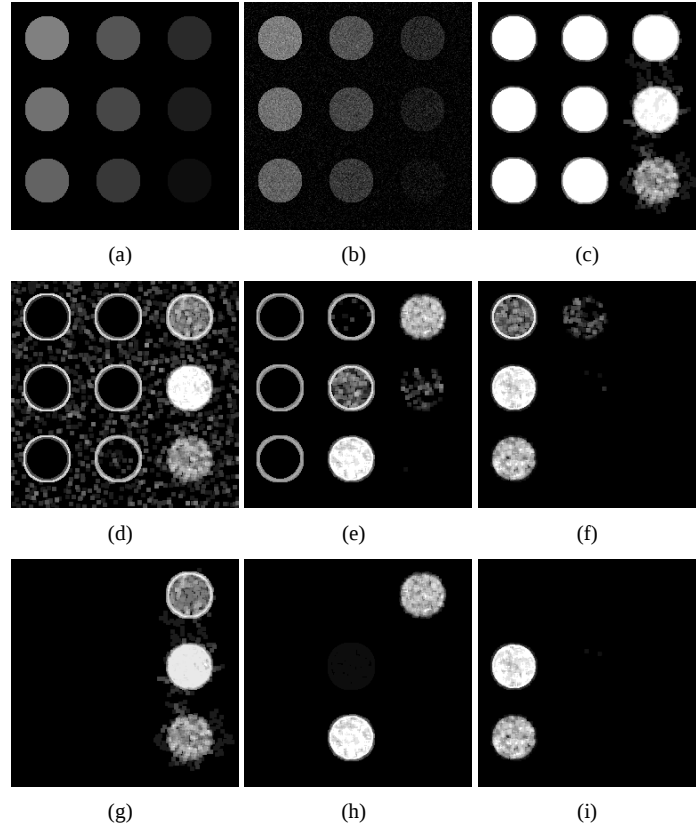


FIG. 4.4 – Exemple d'extraction de zones plates en fonction d'un critère d'aire sur l'image (b), version bruitée de (a). Les images (d), (e) et (f) représentent les coupes $F(*, 10)$, $F(*, 20)$ et $F(*, 40)$ de l'image floue avant filtrage construite à partir de filtres de rang. Les images (g), (h) et (i) représentent ces mêmes coupes après filtrage. L'image (c) représente l'agrégation du résultat du filtrage de tous les $F(*, g)$, $g \in \mathcal{G}$ en utilisant l'opérateur agg .

gris de chaque pixel un degré d'appartenance. Le principe d'extension nous donne :

$$\forall I : \Omega \rightarrow \mathcal{G}, \forall p \in \Omega, \forall g \in \mathcal{G} \quad F(p, g) = \sup_{\Upsilon_{\gamma}(I)(p)=g} \mu_{\gamma}$$

Cette méthode est illustrée dans la figure 4.5. Le problème d'une telle approche est le besoin de calculer un grand nombre d'images débruitées (une par valeur de γ). Pour cette raison, nous nous intéresserons dans le reste de ce chapitre à une méthode dédiée qui est une extension du débruitage par ondelettes dans le cadre des ensembles flous.

4.3 Débruitage flous par ondelettes

Dans cette section nous allons détailler une adaptation du filtrage par ondelettes dans le but de générer une image floue. Tout d'abord, nous rappellerons les principes d'une décomposition en ondelettes d'un signal 1D ainsi que son débruitage. Puis, nous détaillerons une méthode pour modéliser l'imprécision qui est introduite lors d'un tel débruitage ce qui nous permettra de construire des images floues. Nous verrons aussi comment moduler la quantité d'imprécision introduite dans l'étape de construction pour terminer enfin sur l'évaluation d'un tel procédé pour la détection de structures dans des images synthétiques et réelles.

4.3.1 Décomposition en ondelettes

L'analyse en ondelettes est une approche multi-résolution qui repose sur la définition suivante :

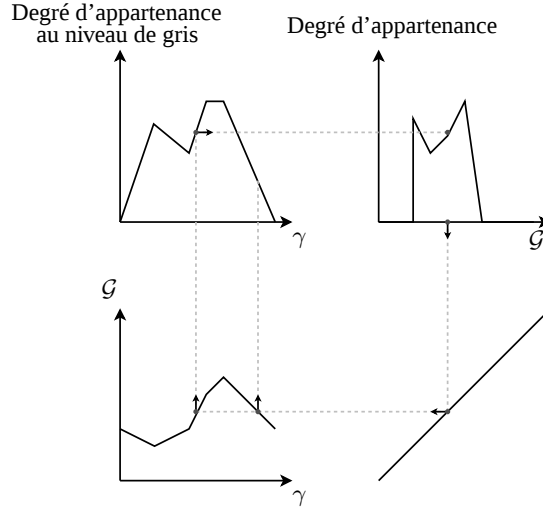


FIG. 4.5 – Construction d'un niveau de gris flou en utilisant le principe d'extension : pour chaque niveau de gris possible pour un pixel donné, les valeurs γ qui conduisent après débruitage à ce niveau de gris sont considérés et le degré d'appartenance maximal de ces derniers est associé au niveau de gris.

Définition 4.3.1. Une analyse multi-résolution est une séquence de sous-espaces $\{V_j\} \subset L_2(\mathbb{R})$ emboîtés et fermés vérifiant (Jansen, 2001) :

$$\begin{aligned} \forall j \in \mathbb{Z} \quad V_j &\subset V_{j+1} \\ \overline{\lim_{j \rightarrow \infty} V_j} &= \overline{\bigcup_{j \in \mathbb{Z}} V_j} = L_2(\mathbb{R}) \\ \lim_{j \rightarrow -\infty} \bigcap_{j \in \mathbb{Z}} V_j &= \{0\} \\ f(x) \in V_j &\Leftrightarrow f(2x) \in V_{j+1}, j \in \mathbb{Z} \text{ (invariance d'échelle)} \\ f(x) \in V_0 &\Leftrightarrow f(x+k) \in V_0, k \in \mathbb{Z} \text{ (invariance par décalage)} \\ \exists \phi(x) \in V_0 / \{\phi(x-k), k \in \mathbb{Z}\} &\text{ est une base stable pour } V_0. \end{aligned}$$

Le premier terme signifie que pour une échelle donnée j , V_j est un sous-espace de V_{j+1} , ce qui résulte en une perte d'information lorsque l'on projette un élément de V_{j+1} sur V_j . Le second terme signifie qu'avec une échelle assez grande, on peut représenter le signal original. Le troisième terme signifie que quand j tend vers $-\infty$, toute l'information à propos du signal est perdue. Les quatrième et cinquième termes impliquent que la même analyse peut être conduite indépendamment de l'échelle et de la localisation. Finalement, le dernier terme est nécessaire pour manipuler des espaces vectoriels de dimensions infinie. La fonction $\phi(x)$ est appelée fonction père ou fonction d'échelle : les fonctions représentant la base de V_0 en sont des versions décalées. En les étirant et en les normalisant, on peut produire les fonctions formant les bases associées aux différents V_j .

Comme on l'a dit, pour des échelles successives j et $j+1$ données, V_{j+1} permet de représenter des éléments que V_j ne contient pas. Pour cette raison, nous avons besoin d'introduire W_j le complément orthogonal de V_j dans V_{j+1} :

$$V_{j+1} = V_j \oplus W_j$$

De manière plus pratique, chaque W_j sera généré grâce à une base composée de translations d'une fonction mère ψ étirée. Cette dernière est aussi connue sous le nom d'ondelette.

Dans le but d'analyser une fonction donnée à différentes échelles, on doit être capable de l'exprimer dans la nouvelle base. Cela peut être fait en utilisant un banc de filtres comme proposé par Mallat (1989). Cette

décomposition repose sur quatre filtres $h, g, \tilde{h}$ et $\tilde{g}$. Ces derniers peuvent être déduits à partir de ϕ, ψ dans le cas d'ondelettes orthogonales et à partir de leur fonctions duales dans le cas d'ondelettes bi-orthogonales. Dans les deux cas on peut reconstruire de manière parfaite le signal original après décomposition.

Soit s_{j+1} une approximation du signal original f dans V_{j+1} . Ce banc de filtres peut être exprimé mathématiquement en utilisant les expressions suivantes :

$$s_{j,k} = \sum_{l \in \mathbb{Z}} \tilde{h}_{l-2k} s_{j+1,l}$$

où $s_{j,k}$ représente les coefficients d'approximation de s_{j+1} dans V_j ,

$$\omega_{j,k} = \sum_{l \in \mathbb{Z}} \tilde{g}_{l-2k} s_{j+1,l}$$

où $\omega_{j,k}$ représente les détails de s_{j+1} exprimés dans W_j .

La reconstruction s'exprime comme :

$$s_{j+1,l} = \sum_{k \in \mathbb{Z}} h_{l-2k} s_{j,k} + \sum_{k \in \mathbb{Z}} g_{l-2k} \omega_{j,k} \quad (4.2)$$

L'extension de ces outils théoriques en $2D$ se fait généralement en traitant une image le long des deux orientations successivement (Jansen, 2001). On commencera donc par filtrer l'image avec les filtres $\tilde{h}$ et $\tilde{g}$ orientés dans l'axe des x , puis le résultat sera à nouveau filtré avec les mêmes filtres orientés dans l'axe des y . Un exemple de résultat d'une telle décomposition est présenté à la figure 4.6. On voit que pour une échelle donnée, on a les détails horizontaux (coin supérieur droit), les détails verticaux (coin inférieur gauche) et diagonaux (coin inférieur droit) avec la décomposition à l'échelle suivante (coin supérieur gauche).

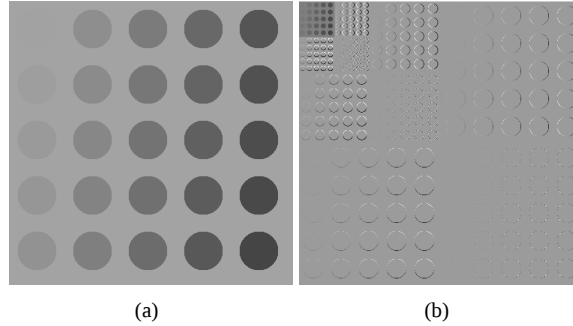


FIG. 4.6 – Exemple de décomposition sur 3 niveaux (b) en ondelettes (Haar) d'une image $2D$ (a).

4.3.2 Seuillage des coefficients d'ondelettes

Une fois que l'image est décomposée sur la base d'ondelettes, les coefficients $\{\omega_{j,k}\}$ peuvent être modifiés dans le but de changer l'aspect de l'image originale. On peut trouver dans la littérature (Jansen, 2001) plusieurs méthodes de débruitage s'appuyant sur les propriétés de décomposition en ondelettes et reposant sur la suppression des petits coefficients. Une première approche correspond à la suppression complète de certains coefficients : les plus grands sont conservés tandis que les plus petits sont mis à zéro. Cette approche est connue sous le nom de *seuillage dur* et s'exprime de manière mathématique en fonction d'un seuil λ comme :

$$tf(\omega) = \begin{cases} \omega & \text{si } \omega \geq \lambda \\ 0 & \text{sinon} \end{cases}$$

Une seconde approche, qui est plus continue, consiste en la soustraction, resp. l'addition, d'une constante donnée aux coefficients positifs, resp. négatifs, les coefficients qui changent de signe sont mis à zéro. Cette

approche est connue sous le nom de *seuillage doux* (Donoho, 1995). La définition mathématique pour un seuil λ s'écrit :

$$tf(\omega) = \begin{cases} \max(0, \omega - \lambda) & \text{si } \omega \geq 0 \\ \min(0, \omega + \lambda) & \text{si } \omega < 0 \end{cases} \quad (4.3)$$

Les deux approches reposent sur un seuil λ et produisent des images plus ou moins dégradées/bruitées en fonction de ce paramètre. Des techniques d'optimisation existent dans la littérature pour déduire de manière optimale ce dernier. Ainsi on peut citer la technique du seuillage universel (Donoho *et al.*, 1995), ou l'estimation non biaisée du risque de Stein aussi connue sous le nom de SURE (Donoho et Johnstone, 1995) ou encore l'approche de validation croisée généralisée souvent dénotée par GCV (Jansen *et al.*, 1997). Les deux dernières approches reposent sur la minimisation de l'erreur quadratique moyenne entre l'image débruitée et ce que serait l'image originale sans bruit. Dans le reste de ce chapitre, nous utiliserons essentiellement la GVC appliquée au *seuillage doux* pour illustrer notre approche. Bien entendu, ce choix est arbitraire et d'autres techniques pourraient tout aussi bien être utilisées.

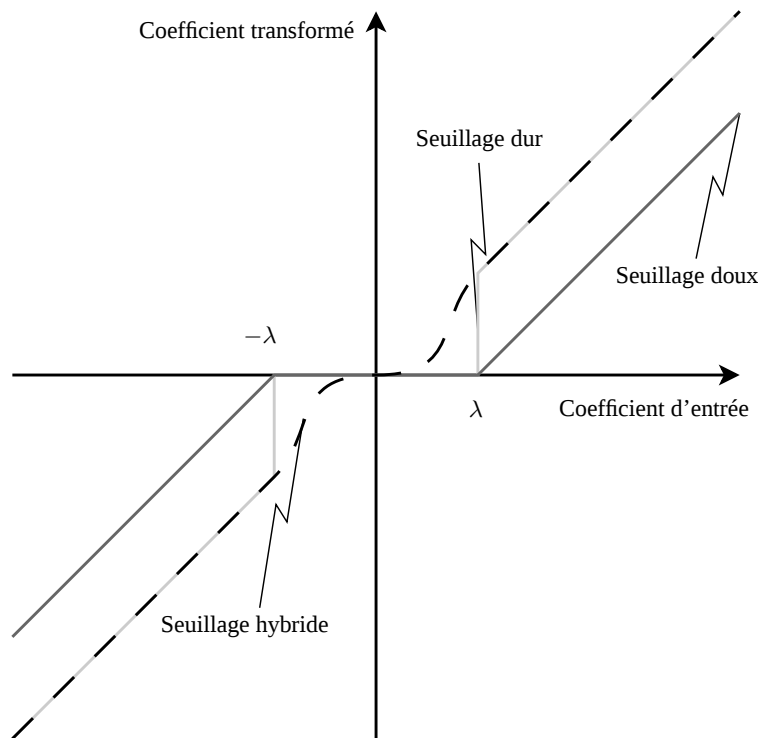


FIG. 4.7 – Fonctions de transfert couramment utilisées pour le débruitage par ondelettes.

La figure 4.7 illustre les fonctions de transfert associées à ces deux approches. On remarquera sur la même figure que des approches hybrides existent pour éviter la discontinuité de la fonction de transfert du *seuillage dur*.

4.3.3 Coefficients flous

Les seuils généralement utilisés avec les techniques de seuillage précédemment citées correspondent généralement à des compromis entre la quantité de bruit à enlever et la quantité de détails que l'on désire garder après traitement. Notre idée est que l'impact du choix que l'on fait à ce niveau peut être affaibli par l'introduction de l'imprécision sur ce paramètre. Cela est la raison fondamentale qui nous motive dans l'établissement d'un processus pour rendre flou les coefficients d'ondelettes ω_j .

Sans grande perte en termes de généralité, nous considérerons des nombres flous en utilisant une représentation LR qui est un cas particulier de la représentation LR d'intervalles flous. De telles quantités sont

représentées en utilisant deux fonctions $L : \mathbb{R}^+ \rightarrow [0, 1]$ et $R : \mathbb{R}^+ \rightarrow [0, 1]$ telles que $L(0) = R(0) = 1$, $L(1) = R(1) = 0$ et $\forall x > 1$ $L(x) = R(x) = 0$. Additionnellement, quatre paramètres (m_1, m_2, a, b) sont requis pour la définition de la fonction d'appartenance d'un intervalle $Q = (m_1, m_2, a, b)_{LR}$. Cette dernière est définie comme :

$$\begin{cases} \mu_Q(x) = L\left(\frac{m_1 - x}{a}\right) & \text{si } x \leq m_1 \\ \mu_Q(x) = 1 & \text{si } m_1 < x < m_2 \\ \mu_Q(x) = R\left(\frac{x - m_2}{b}\right) & \text{si } x \geq m_2 \end{cases}$$

Dans le cas d'un nombre flou, les paramètres m_1 et m_2 sont égaux. Cette représentation est illustrée à la figure 4.8 pour un intervalle flou.

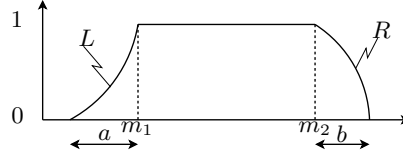


FIG. 4.8 – Représentation LR d'un intervalle flou.

De retour sur notre problème de génération de coefficients flous d'ondelettes, nous proposons une approche dans le cas du seuillage doux, fondée sur l'hypothèse que la valeur de coefficient seuillée la plus possible est obtenue par l'approche GCV, et que sa pire valeur est obtenue pour la valeur initiale du coefficient (c'est-à-dire $\lambda = 0$). Cela nous permet de transformer l'équation 4.3 pour obtenir le schéma de construction de coefficients flous suivant :

$$\mu_\omega(x) = \begin{cases} (\max(\omega - \lambda, 0), \max(\omega - \lambda, 0), 0, \beta(\omega - \max(\omega - \lambda, 0)))_{LR} & \text{si } \omega \geq 0 \\ (\min(\omega + \lambda, 0), \min(\omega + \lambda, 0), \beta(\min(\omega + \lambda, 0) - \omega), 0)_{LR} & \text{sinon} \end{cases}$$

avec β un paramètre qui nous permet de régler la quantité d'imprécision que l'on désire introduire. La figure 4.9 illustre ce procédé. Le noyau des coefficients flous, qui apparaissent en gris dans la figure, est obtenu à partir de la fonction de transfert optimale qui correspond à la méthode du GCV. La pente de ces nombres flous est plus ou moins douce en fonction du paramètre β . Quand ce dernier est égal à 1, la fonction d'appartenance devient nulle à partir de la valeur originale du coefficient, ce qui correspond au comportement intuitif que l'on a décrit précédemment.

4.3.4 Génération d'une image floue

Le problème principal restant se situe au niveau du processus de reconstruction du banc de filtres. En effet la nature de certains des éléments que l'on additionne et que l'on multiplie à l'équation 4.2 a changé : les coefficients d'ondelettes sont devenus des nombres flous. Néanmoins, le cadre des ensembles flous permet de redéfinir les opérateurs usuels pour intervalles flous. En utilisant une représentation LR, une grande partie des opérateurs peuvent facilement être mis en œuvre. Ainsi si l'on considère deux intervalles flous $Q_1 = (m, m', a, b)_{LR}$ et $Q_2 = (n, n', c, d)_{LR}$, on peut exprimer l'opérateur + flou noté $\oplus$ comme :

$$Q_1 \oplus Q_2 = (m + n, m' + n', a + c, b + d)_{LR} \quad (4.4)$$

Le produit entre un ensemble flous et un scalaire positif se note $\odot$ et s'écrit :

$$\alpha \odot Q_1 = (\alpha m, \alpha m', \alpha a, \alpha b)_{LR}$$

alors que dans le cas où le scalaire est négatif, il s'exprime comme :

$$\alpha \odot Q_1 = (\alpha m', \alpha m, -\alpha b, -\alpha a)_{RL}$$

Dans la mesure où l'équation 4.2 n'a besoin que de ces trois définitions dans le cas de coefficients $\{s_{j,k}\}$ et $\{\omega_{j,k}\}$ flous, cette expression peut se réécrire comme :

$$s_{j+1,l} = \sum_{k \in \mathbb{Z}}^{\oplus} h_{l-2k} \odot s_{j,k} + \sum_{k \in \mathbb{Z}}^{\oplus} g_{l-2k} \odot \omega_{j,k} \quad (4.5)$$

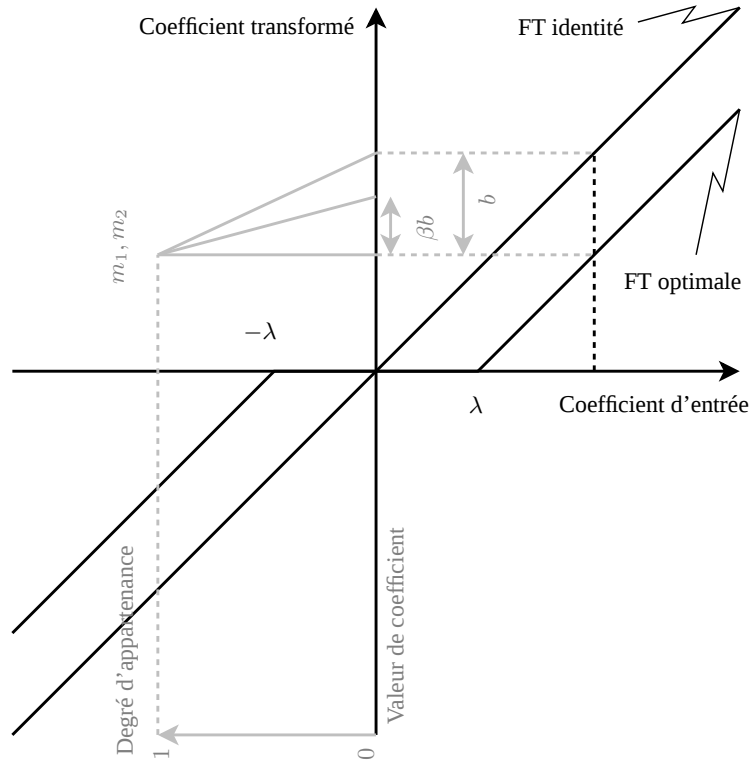


FIG. 4.9 – Construction d'un coefficient d'ondelettes flou. Un degré d'appartenance de 1 est associé à la valeur de coefficient obtenue pour un seuillage doux optimal. Le degré d'appartenance décroît quand le nouveau coefficient se rapproche de sa valeur originale. Les fonction de transfert *FT identité* et *FT optimale* correspondent respectivement aux fonctions obtenues pour $\lambda = 0$ et pour le λ obtenu par l'approche GCV.

avec $\sum^\oplus$ l'équivalent flou de $\sum$ selon l'équation 4.4.

La reconstruction complète d'une image floue grâce à cette expression donne une image de nombres flous. Dans la mesure où $\mathcal{G}$ est discret, l'échantillonnage des niveaux de gris flous sur cet ensemble peut déboucher sur des pixels où aucun niveau de gris n'a un degré d'appartenance égal à un. Clairement dans ce cas, nous ne manipulons plus une image de nombres flous. Pour pallier cette limitation, une dernière étape consistant à dilater le niveau de gris flou peut être faite : on soustrait, resp. ajoute, $\frac{s}{2}$ à m_1 , resp. m_2 , avec s le pas d'échantillonnage comme illustré dans la figure 4.10. Dans le cas où l'on désire une image d'ombres floues, on peut forcer de manière arbitraire la décroissance par rapport aux niveaux de gris à partir d'une image de nombres flous.

4.3.5 Quantité de flou

Comme expliqué précédemment, le paramètre β introduit dans le procédé de construction des coefficients flous permet de régler la quantité d'imprécision que l'on retrouvera dans l'image floue finale. Dans le but de quantifier cette dernière, on peut utiliser une mesure de degré de flou comme l'entropie floue proposée par De Luca et Termini (1972). Dans la mesure où les coefficients d'ondelettes sont des nombres réels, on utilisera dans notre contexte une version continue de la mesure :

$$E(\mu) = - \int_{-\infty}^{\infty} \mu(x) \log_2(\mu(x)) + (1 - \mu(x)) \log_2(1 - \mu(x)) dx$$

Théorème 4.3.1. $\forall \beta \in \mathbb{R}^+$, pour des intervalles flous $Q^\beta = (m_1, m_2, \beta a, \beta b)_{LR}$, on a $E(\mu_{Q^\beta}) = \beta E(\mu_{Q^1})$.

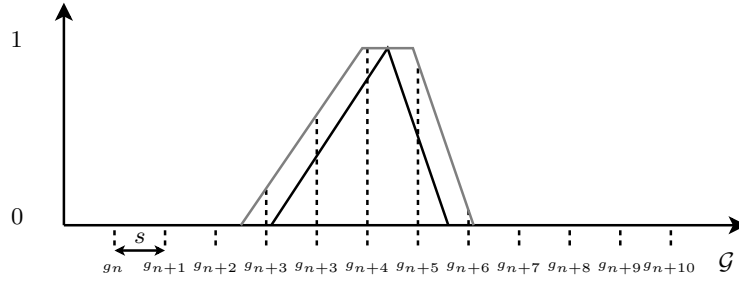


FIG. 4.10 – Echantillonnage sur $\mathcal{G}$ d'un niveau de gris flou. le nombre flou original (en noir) ne voit pas son noyau correspondre avec les niveaux de gris échantillonnés alors que sa version dilatée de $\frac{s}{2}$ (en gris) ne souffre pas de ce problème.

Démonstration. Par définition, on a :

$$\begin{aligned} E(\mu_{Q^\beta}) = & - \int_{-\infty}^{m_1} R\left(\frac{m_1-x}{\beta a}\right) \log_2 R\left(\frac{m_1-x}{\beta a}\right) dx \\ & - \int_{m_2}^{m_1} 1 \log(1) dx \\ & - \int_{m_2}^{\infty} L\left(\frac{x-m_2}{\beta b}\right) \log_2 L\left(\frac{x-m_2}{\beta b}\right) dx \\ & - \int_{-\infty}^{m_1} \left(1 - R\left(\frac{m_1-x}{\beta a}\right)\right) \log_2 \left(1 - R\left(\frac{m_1-x}{\beta a}\right)\right) dx \\ & - \int_{m_2}^{m_1} 0 \log(0) dx \\ & - \int_{m_2}^{\infty} \left(1 - R\left(\frac{x-m_2}{\beta b}\right)\right) \log_2 \left(1 - R\left(\frac{x-m_2}{\beta b}\right)\right) dx \end{aligned}$$

les second et cinquième termes sont nuls, et les autres termes sont des versions étirées de l'intégration de R et L sur $[0, 1]$, ce qui nous donne :

$$\begin{aligned} E(\mu_{Q^\beta}) = & -\beta a \int_0^1 R(x) \log_2(R(x)) dx \\ & -\beta b \int_0^1 L(x) \log_2(L(x)) dx \\ & -\beta a \int_0^1 (1 - R(x)) \log_2(1 - R(x)) dx \\ & -\beta b \int_0^1 (1 - L(x)) \log_2(1 - L(x)) dx \end{aligned}$$

Ainsi on a $E(\mu_{Q^\beta}) = \beta E(\mu_{Q^1})$. □

Introduisons maintenant la notion de degré de flou dans une image ou quantité de flou par pixel (*qfpp*)

Définition 4.3.2. Soit F une image floue, la quantité de flou par pixel s'exprime comme :

$$qfpp(F) = \frac{1}{|\Omega|} \sum_{p \in \Omega} E(F(p, *))$$

On peut finalement introduire un dernier théorème :

Théorème 4.3.2. Soit $I \in \mathcal{I}$ une image nette, $\beta \in \mathbb{R}^+$ et F^β la version floue de I calculée à partir de β , on a :

$$qfpp(F^\beta) = \beta qfpp(F^1)$$

Démonstration. Soit $I \in \mathcal{I}$ une image nette. En utilisant l'équation 4.5 et les définitions de $\oplus$ et $\odot$, il existe $\{Q_i^\gamma = (m_1^i, m_2^i, \gamma a^i, \gamma b^i)\}$ définis pour $\gamma \in \mathbb{R}^+$ tels que $\forall \beta \in \mathbb{R}^+$ l'image floue F^β vérifie : $\forall p \in \Omega$ $F^\beta(p, *) = Q_i^\beta$. Ainsi, en utilisant le théorème 4.3.1, on peut déduire que $E(F^\beta(p, *)) = E(Q_i^\beta) = \beta E(Q_i^1) = \beta E(F^1(p, *))$.

Finalement, en utilisant la définition 4.3.2 on peut conclure que le théorème 4.3.2 est vérifié. □

Ce théorème est utile puisque si l'on veut construire plusieurs images floues $\{F^\beta\}$ qui ont une valeur de *qfpp* donnée, on a seulement besoin de produire une image floue F^1 , et déduire l'ensemble des β nécessaires. De plus, pour chaque β , F^β peut directement être calculé à partir de F^1 dans la mesure où β est un facteur des paramètres a^i et b^i des nombres flous contenus dans l'image floue reconstruite.

4.3.6 Expérimentations sur images synthétiques

Dans le but d'évaluer la robustesse de la méthode, nous l'avons testée sur deux images synthétiques respectivement dégradées avec un bruit Gaussien et un bruit de Poisson. L'expérimentation reprend le principe, tout en étant plus complète, de celles utilisées pour l'évaluation des méthodes de construction d'image floues en déployant un patron sur les niveaux de gris d'une image nette et en utilisant les filtres de rang. L'image originale est composée de 25 disques disposés dans un fond uniforme comme illustré à la figure 4.11. Les niveaux de gris des disques en partant du coin en haut à gauche, vers le coin en bas à droite sont décroissants à partir de l'intensité du fond (égale à 275) moins un. Cela permet de considérer des disques avec un rapport contraste à bruit (RCB) croissant. Sans bruit, une simple notion de zones plates classique en plus d'un critère sur la taille des disques permet de tous les récupérer. Dans notre expérimentation, nous évaluons la capacité à extraire les mêmes zones à partir d'une *IIF* en utilisant la notion floue de zones plates.

Expérimentations

Dans chaque cas, l'image dégradée est décomposée sur une base d'ondelettes introduite par Daubechies (1988) en utilisant quatre niveaux de décomposition. Les coefficients d'ondelettes sont alors convertis en nombre flous dans le but de générer une image d'intervalles flous en utilisant la méthode décrite dans les sous-sections précédentes. Pour extraire les disques, les images floues sont filtrées en utilisant le filtre présenté à l'équation 4.1.

Tout comme dans les expérimentations précédentes, le résultat de ce filtre sera interprété en utilisant l'opérateur *agg*. Un exemple de détection dans l'image de la figure 4.11(b) est donné à la figure 4.11(e). Idéalement, l'ensemble flou résultant de la détection est égal à un dans les disques et à zéro dans le fond. Ainsi, pour chaque disque, on calcule une mesure de similitude (Rifqi, 1996) avec l'ensemble net égal à un dans le disque considéré et à zéro dans le reste de l'image dans le but d'évaluer comment se comporte la méthode pour différents rapports contraste à bruit. En pratique la mesure suivante a été utilisée :

$$\forall (f_1, f_2) \in \mathcal{S}^2 \quad S(f_1, f_2) = \frac{\sum_{p \in \Omega} \min(f_1(p), f_2(p))}{\sum_{p \in \Omega} f_2(p)}$$

Evidemment, la quantité de flou introduite dans l'image va jouer un rôle important sur la capacité à extraire les disques de l'image. En utilisant le théorème 4.3.2, la mesure de similitude précédente pour les différents disques peut facilement être évaluée pour divers valeurs de *qfpp* sans avoir besoin de reconstruire une image floue à chaque fois.

Bruit gaussien

La figure 4.12 illustre les capacités de détection de l'approche sur l'image de la figure 4.11(b), qui contient du bruit gaussien d'écart type de 16,5. On observe que pour chaque rapport contraste à bruit, il y a une quantité de flou par pixel optimale pour la détection du disque correspondant. De trop petites valeurs de *qfpp* ne permettent pas de reconnecter les disques tandis que de trop grandes valeurs tendent à fusionner le disque et le fond ne permettant ainsi aucune détection. Un second aspect clé est que comme ces valeurs optimales dépendent du rapport contraste à bruit du disque cible, il n'est pas possible de construire une unique image floue qui permette de détecter de manière optimale chacun des disques. Les valeurs optimales de *qfpp* pour les disques de faible contraste sont proches de 0 alors qu'elles sont bien plus grandes pour des disques mieux contrastés. Néanmoins, dans un cadre de détection, des compromis peuvent être trouvés. En effet, si on choisit une *qfpp* optimisant la détection d'une structure assez contrastée, la détection de structures plus contrastées sera au moins d'aussi bonne qualité : dans la plupart des applications, un masque imparfait représentant les structures recherchées peut suffire pour initialiser d'autres traitements comme par exemple de la segmentation.

On peut se référer aux images des figures 4.11(d), 4.11(e) et 4.11(f) pour avoir un aperçu des détections dans l'image bruitée pour différentes quantités de flou. On remarque que pour de faibles *qfpp*, les premiers disques qui commencent à apparaître sont les moins contrastés (c.f. figure 4.11(d)). Néanmoins ils ne sont que faiblement détectés. Si on augmente la quantité de flou dans l'image, on peut obtenir une détection correcte de quasiment tous les disques (c.f. figure 4.11(e)). Seuls les deux moins bien contrastés posent

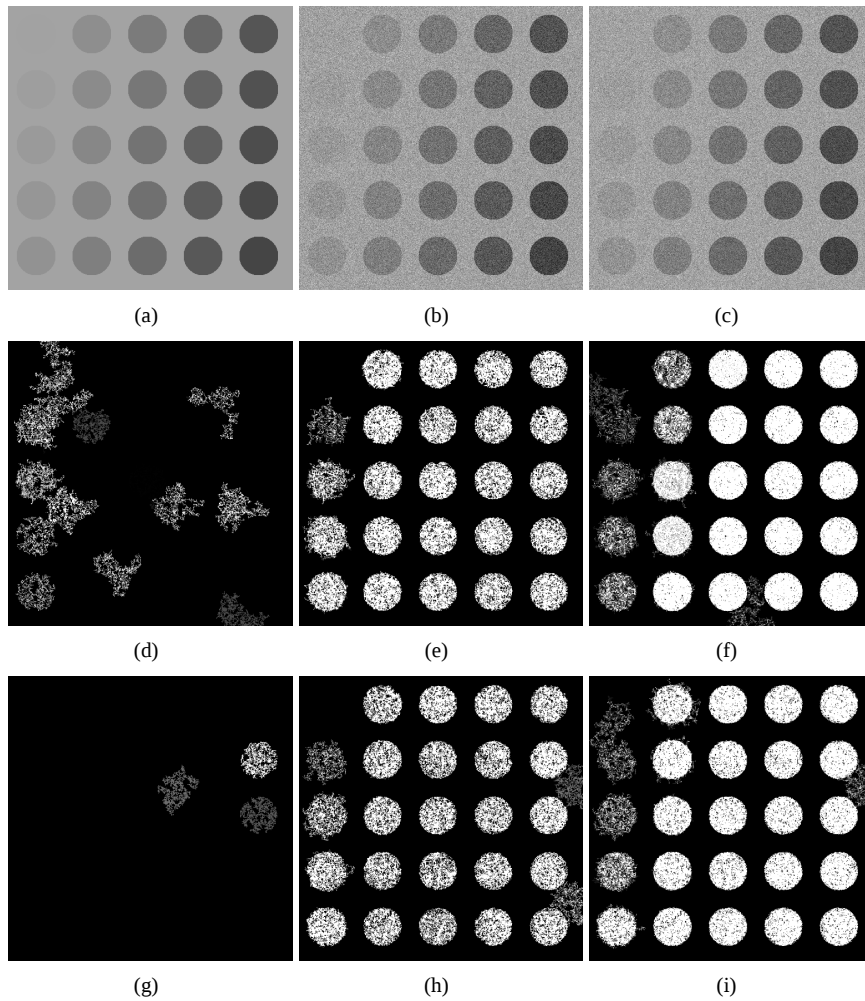


FIG. 4.11 – (a) Image originale composée de 25 disques de différents contrastes. (b) La même image corrompue avec un bruit gaussien (écart type de 16,5). (c) La même image avec un bruit de Poisson. (d), (e) et (f) Exemples de détection à partir de l'image (b) pour des $qfpp$ égales à 3,5, 6, 12. (g), (h) et (i) Exemples de détection à partir de l'image (c) pour des $qfpp$ égales à 3,5, 6, 9.

problème : l'un n'étant pas détecté et l'autre faiblement. Enfin le dernier exemple (c.f. figure 4.11(f)) illustre le critère d'optimalité de la quantité de flou à introduire. En effet, on peut remarquer que les disques les mieux contrastés sont mieux détectés que dans l'exemple précédent alors que les disques les plus faiblement contrastés sont moins bien détectés (colonne de gauche). On remarquera enfin que dans certains cas comme par exemple dans la figure 4.11(f), des structures autres que les disques sont détectées. Cela est dû au critère utilisé pour extraire les disques. En effet, on n'utilise que l'aire des composantes connexes, ainsi des composantes reconnectées de manière artificielle peuvent aussi correspondre à ce critère. De plus les images présentées étant des images d'agrégation, le fait que certaines de ces zones soient connexes aux disques ne signifie pas qu'elles soient réellement connexes dans l'image floue puisque les détections peuvent provenir de niveaux de gris différents. De manière pratique, ce fait se vérifie, de plus lorsque les disques sont vraiment connexes dans l'images floue à des zones reconnectées de manière abusive, les définitions de connexité d'un ensemble flou nous assurent que les degrés d'appartenance décroissent rapidement comme on peut le voir par exemple sur le second disque de la seconde colonne dans la figure 4.11(f). Enfin, on peut aussi ajouter qu'à cause de l'imperfection de l'image floue, un disque est représenté par plusieurs zones plates floues dans l'image floue, ainsi en ne considérant pas celles précédemment décrites, en utilisant par exemple un critère plus évolué que simplement l'aire, on peut tout de même détecter les disques. Finalement, dans la

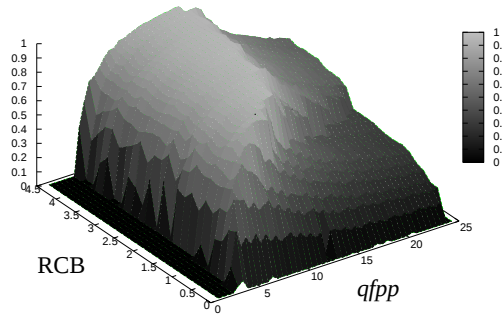


FIG. 4.12 – Evolution de la similitude entre les disques présents dans l'image de la figure 4.11(a) et leur détection en utilisant l'équation 4.1 couplée à l'opérateur *agg*, dans l'image corrompue avec un bruit gaussien de la figure 4.11(b). Les similitudes sont données en fonction du rapport contraste à bruit et de la quantité de flou par pixel introduite dans l'image floue.

mesure où le but de l'expérience n'est pas ici de faire un détecteur de disque, mais plutôt de vérifier si ces disques sont présents dans l'image floue et s'ils peuvent être récupérés, cela n'a que peu d'impact sur les conclusions que l'on peut tirer de cette expérimentation.

Bruit de Poisson

Dans le cas d'une image contenant un bruit de Poisson, on ne peut pas utiliser directement l'approche de débruitage par ondelettes dont les seuils optimaux sont fournis par GCV. En effet cette technique fait l'hypothèse que le bruit présent dans l'image est gaussien (Jansen, 2001). Néanmoins, on peut appliquer un pré-traitement sur l'image nette en utilisant les travaux de Anscombe (1948) visant à stabiliser la variance du bruit. De cette manière, on obtient des données avec une distribution plus proche du cas gaussien. Ce pré-traitement consiste à appliquer la fonction de transfert suivante :

$$tf(x) = 2\sqrt{x + \frac{3}{8}}$$

L'évolution de la détectabilité pour cette configuration de bruit est différent du cas gaussien comme on peut le voir dans la figure 4.13. Ici, de petites valeurs de *qfpp* permettent de détecter en premier les disques de fort contraste à bruit. Ce résultat est assez intuitif dans la mesure où le fond a une intensité plus grande que celle des disques, ce qui a pour conséquence que le bruit à l'intérieur de ces derniers devient de plus en plus faible à mesure que le contraste augmente (l'intensité du disque, et par conséquent la variance du bruit diminuent). La variation autour de ces zones plates étant plus faible, une quantité de flou moins importante est requise pour les extraire. Enfin, les conclusions à propos des valeurs optimales de *qfpp* et de leur dépendance au contraste sont encore valides.

Tout comme dans le cas gaussien, on peut se référer aux figures 4.11(g), 4.11(h) et 4.11(i) pour apprécier le résultat de la détection dans les images floues construites à partir de l'image bruitée de la figure 4.11(c). La première image illustre le fait que les disques les plus contrastés sont les premiers à être détectés lorsque l'on augmente progressivement la quantité de flou dans l'image. La deuxième figure illustre un compromis sur la quantité de flou à utiliser pour extraire de manière correcte un grand nombre de disques. Enfin en comparant cette dernière détection à la troisième image, on peut s'apercevoir que les remarques sur l'optimalité de la quantité de flou à introduire dans l'image sont toujours valables : les disques les mieux contrastés sont mieux détectés alors que les moins bien contrastés se dégradent.

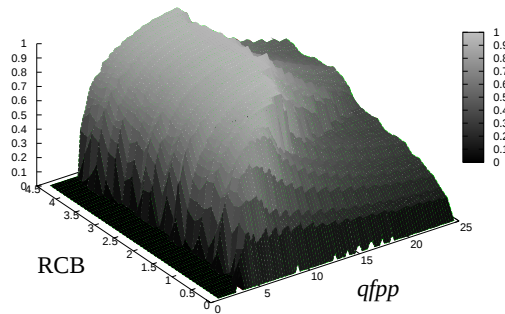


FIG. 4.13 – Evolution de la similitude entre les disques présents dans l'image de la figure 4.11(a) et leur détection en utilisant l'équation 4.1 couplée à l'opérateur *agg*, dans l'image corrompue avec un bruit de Poisson de la figure 4.11(c). Les similitudes sont données en fonction du rapport contraste à bruit et de la quantité de flou par pixel introduite dans l'image floue.

4.3.7 Expérimentation sur une image réelle

Pour valider le comportement de la détection de zones plates floues, une expérimentation a été effectuée sur une image réelle. L'image considérée est une région d'intérêt extraite d'une coupe d'un volume de tomosynthèse du sein contenant une lésion de radiométrie à peu près constante. Cette image, qui est présentée à la figure 4.14(a), présente aussi une superposition entre la lésion et une autre structure invalidant partiellement l'hypothèse que l'opacité est d'intensité constante. Pour voir si des zones plates floues correspondent à cette dernière dans l'image floue créée à partir de la même approche par ondelettes que précédemment, celle-ci a été filtrée par le même opérateur que pour les images synthétiques, c'est-à-dire celui défini à l'équation 4.1, ses paramètres étant adaptés à la taille de la lésion. On remarquera encore que ce filtre assez simpliste n'a pas pour but de détecter uniquement les lésions dans la mesure où seul un test sur la surface des objets est effectué. Au contraire, il sert juste à vérifier la présence de composantes connexes floues compatibles avec la lésion présente.

La figure 4.14 expose le résultat du filtrage pour différentes quantités de flou par pixel. Dans la mesure où la méthode de construction est la même que pour les exemples synthétiques, les images floues associées à ces différentes valeurs sont produites grâce au théorème 4.3.2. On remarquera qu'une partie interne de la lésion n'est pas complètement détectée alors que le reste l'est pour les valeurs de quantité de flou les plus faibles comme on peut le voir sur les figures 4.14(d) à 4.14(h). Cela est dû au fait qu'une structure se superpose à la lésion à cet endroit, rendant ainsi l'hypothèse de stabilité radiométrique invalidée aux frontières de cette zone. Lorsque l'on augmente encore la quantité de flou, on reconnecte ces deux parties comme illustré aux figures 4.14(i) et 4.14(j). Néanmoins, on s'apercevra qu'en faisant cela, la quantité de flou n'est plus optimale pour le reste de la lésion, dégradant ainsi légèrement le résultat comme on peut le voir sur la partie supérieure droite de l'opacité à la figure 4.14(j). Ce phénomène d'amélioration et de dégradation simultanées est visible sur un intervalle de *qfpp* illustré par les figures 4.14(g) à 4.14(k). Au-delà de cet intervalle, la qualité de la détection se dégrade de manière globale validant ainsi les conclusions tirées précédemment sur les images synthétiques.

4.4 Conclusion

Dans ce chapitre, nous avons abordé le problème de la construction des images floues. Nous avons proposé des méthodes générales comme le déploiement d'un patron sur les niveaux de gris originaux et l'utilisation de filtres de rang, mais aussi un cadre général de construction d'images à partir de méthodes de débruitage. En effet, dans la mesure où les images floues ne peuvent représenter correctement un bruit statistique, le débruitage des images est nécessaire. L'idée principale développée dans ce chapitre est de

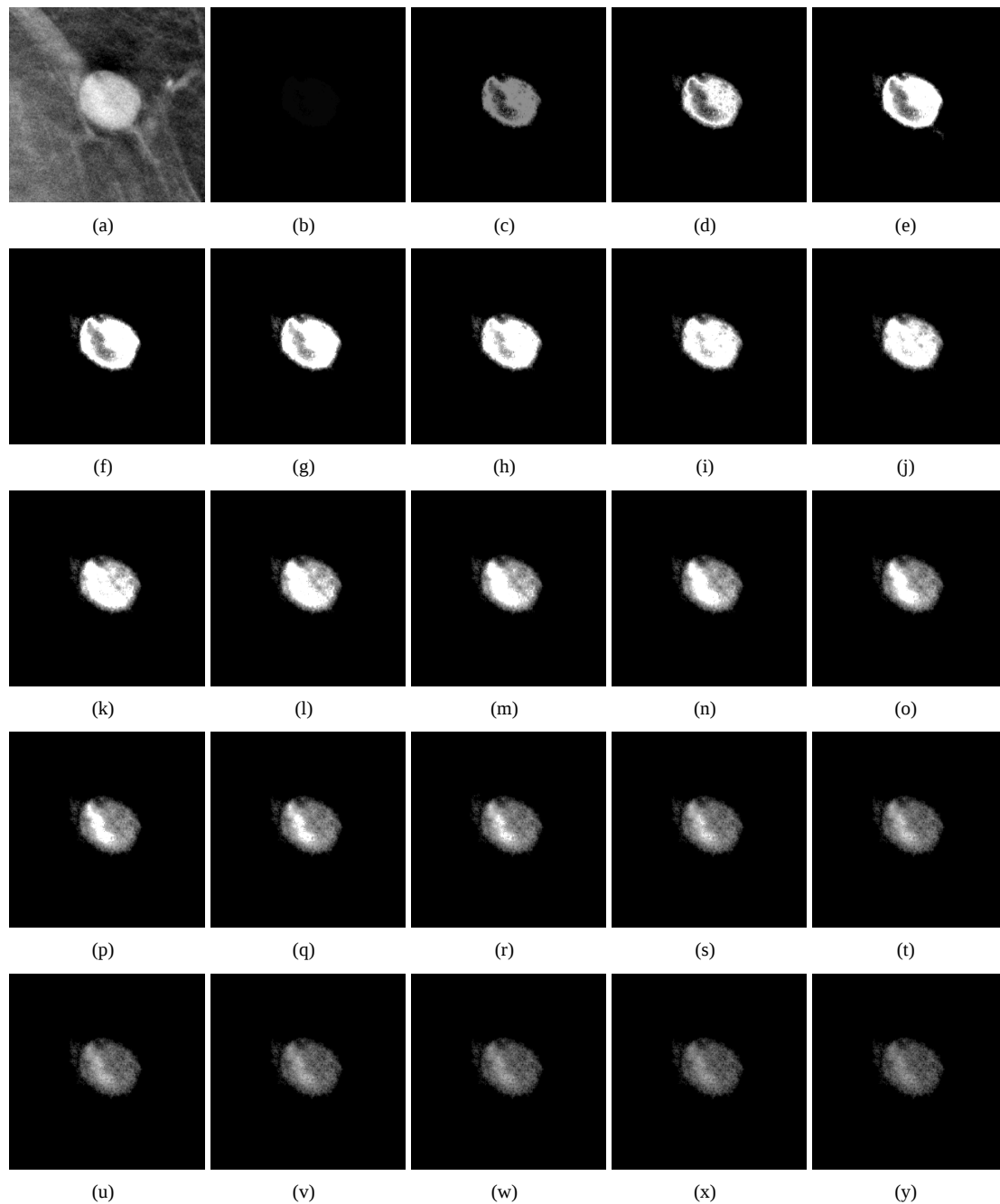


FIG. 4.14 – Extraction des zones plates floues correspondant à une lésion dans une coupe d'un volume de tomosynthèse du sein. (a) Région d'intérêt contenant la lésion. Les figures (b) à (y) correspondent aux détections pour des $qfpp$ croissantes.

modéliser l'imprécision produite dans cette étape de débruitage de manière à la transmettre dans les images floues. De cette manière on transforme le bruit statistique présent dans les images en imprécision de débruitage. Dans ce contexte une extension des méthodes de débruitage par ondelettes a été proposée pour construire des images floues.

Une étude sur des images synthétiques a aussi été conduite permettant de mettre en évidence que la façon de construire les images floues, et notamment la quantité d'imprécision introduite dans l'image,

influe directement sur la capacité qu'a un filtre connexe à détecter diverses structures. Ainsi, on a montré que dans le cas d'images contenant un bruit gaussien ou de Poisson, la quantité optimale de floue à introduire dépendait du rapport contraste à bruit des structures à détecter. Ces conclusions ont aussi été validées sur une image réelle extraire d'une coupe de volume de tomosynthèse de sein.

Chapitre 5

Segmentation floue

La caractérisation des lésions est une étape importante dans une chaîne de détection automatique de lésions. Dans le cas des opacités, il est assez naturel de délimiter les structures suspectes à l'aide d'un contour. Malheureusement, certaines lésions peuvent parfois être assez mal discernables, posant ainsi le problème de leur localisation, de leur délimitation voire de leur présence. Dans ce chapitre, nous allons proposer des méthodes de segmentation reposant sur la logique floue et permettant de prendre en compte ces problèmes.

Dans un premier temps nous détaillerons le contexte de la segmentation floue. Nous parlerons ensuite d'une méthode de segmentation reposant sur des seuillages multiples. Puis, nous présenterons une approche d'extraction de contours flous à partir d'une méthode de segmentation proposée par Peters *et al.* (2007). Nous finirons enfin par l'extension d'une méthode de segmentation reposant sur des principes de programmation dynamique (Timp et Karssemeijer, 2004).

5.1 Intérêt des contours flous

Certains problèmes comme la difficulté de définir un unique contour ou le doute sur la présence d'une lésion peuvent être pris en compte en utilisant le formalisme de la logique floue. Dans un premier temps nous rappellerons quels sont les défauts d'une image et leur implication en terme de segmentation puis nous parlerons du formalisme des contours flous.

5.1.1 Interprétation d'une image : imprécision contre incertitude

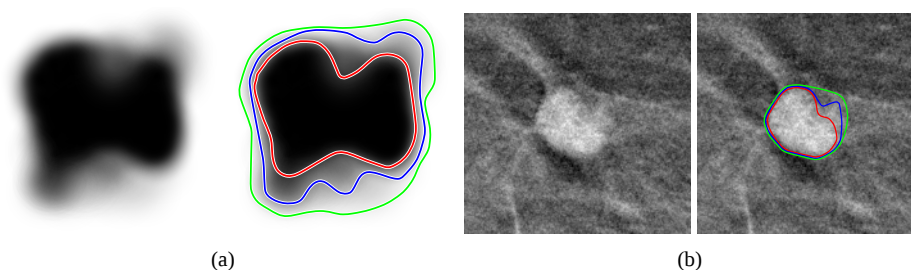


FIG. 5.1 – Illustration de l'imprécision pour la définition d'un contour sur une image synthétique (a) et sur une image contenant une lésion circonscrite (b).

Comme beaucoup de problèmes de traitement d'images, la segmentation de lésions est une tâche difficile. En effet, certaines de ces structures peuvent parfois être assez mal discernables. En pratique on peut remarquer deux types de défaut responsables de ces difficultés. Tout d'abord, une lésion peut avoir une

partie de son contour difficilement localisable, même si sa présence ne porte pas à discussion. Ce type de défaut correspond à de l'imprécision. La figure 5.1 présente un exemple d'imprécision sur une lésion et sur un exemple synthétique. Comme on le voit à la figure 5.1(a), ce genre de défaut se traduit par des contours faiblement définis et mal localisés. Le deuxième type de défaut correspond à l'incertitude, c'est-à-dire que l'on n'est pas sûr de l'endroit où est la lésion ni même de sa présence. La figure 5.2 illustre l'incertitude sur une lésion et une image synthétique.

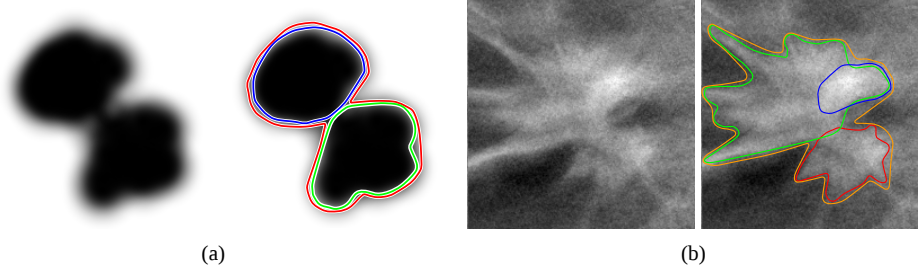


FIG. 5.2 – Illustration de l'incertitude pour la définition d'un contour sur une image synthétique (a) et sur une image contenant une lésion spiculée (b).

5.1.2 Formalisme des contours flous

Des travaux ont déjà été proposés pour prendre en compte les deux types de défauts présentés précédemment. L'idée est de ne pas segmenter de manière unique une lésion mais plutôt d'extraire différents contours possibles. Ainsi pour une segmentation, on associe un ensemble de contours inclus les uns dans les autres, eux mêmes associés à un degré d'appartenance à la classe contour. Ce formalisme appelé contour flou permet de modéliser l'imprécision d'un contour. L'incertitude est quant à elle prise en compte par l'extraction possible de plusieurs contours flous pour une même lésion. Originellement, ce formalisme a été proposé avec une méthode de segmentation par multi-seuillage. Cette approche sera discutée à la section 5.2. Briellement, l'idée est de considérer que chacun des maxima de l'image est une lésion potentielle (modélisation de l'incertitude), et que chaque composante connexe issue de tous les seuillages possibles de l'image est un contour possible (modélisation de l'imprécision).

5.2 Segmentation par seuillages multiples

Les images que nous cherchons à segmenter sont des images de tomosynthèse contenant des surdensités. De manière idéale, ce type d'image devrait être composé de voxels dont l'intensité est proportionnelle au matériau représenté. Ainsi, on peut envisager qu'un seuillage de l'image nous donne un contour correct de la lésion.

5.2.1 Principe

Comme dans la méthode proposée par Peters (2007), chaque contour est extrait par les différents seuillages de l'image. On lui associe un degré d'appartenance qui est une conjonction de différents degrés de satisfaction de propriétés qu'un contour possible doit vérifier.

De manière pratique, pour réaliser une segmentation on a besoin d'un marqueur représenté par un pixel q . De manière formelle, on définit l'ensemble des contours possibles dans une image I pour un maxima m donné :

$$C_m^q = \bigcup_g \left\{ \varphi(C) / (C \in H(X_g^+(I))) \wedge \left(\hat{\Gamma}_{X_{I(q)}^+(I)}^q \cap C \neq \emptyset \right) \wedge (m \in C) \right\}$$

avec φ un opérateur croissant, $H(A)$ l'ensemble des composantes connexes incluses dans un ensemble A , $X_g^+(I)$ l'opérateur de seuillage d'une image I au niveau de gris g et $\hat{\Gamma}_A^q$ la composante connexe de A

contenant le point q . Plus concrètement le rôle de l'opérateur φ est de supprimer les trous contenus dans les ensembles C ainsi que de régulariser les contours par une fermeture morphologique.

En considérant tous les maxima m_i de l'image tels que $C_i^p \neq \emptyset$, on obtient différents ensembles de contours possibles pour un même marqueur p . Le fait de considérer plusieurs ensembles de contours possibles permet de manipuler l'incertitude sur la structure à segmenter. Pour obtenir des contours flous à partir des différents ensembles C_m^q , il ne reste qu'à calculer un degré d'appartenance pour chacun des contours. Cela se fait en définissant des fonctions d'appartenance à la classe contour en fonction d'attributs différents. Dans notre contexte de segmentation de lésion nous avons utilisé un critère de compacité combiné à un critère de contraste. Ces deux critères se traduisent par deux fonctions d'appartenance ν_{comp} et ν_{trast} définies sur les valeurs possibles d'attribut. Ainsi le degré d'appartenance d'un contour $C \in C_m^q$ se calcule de la manière suivante :

$$\mu_C = \nu_{\text{comp}}(\text{compacité}(C)) \top \nu_{\text{trast}}(\text{contraste}(I, C)) \quad (5.1)$$

où $\top$ est une t-norme.

Remarquons que par la propriété de φ et par l'emboîtement des différentes composantes connexes issues des différents seuillages de l'image, les contours inclus dans les contours flous obtenus sont emboîtés.

5.2.2 Lien avec les images d'ombres floues

On remarquera que ce formalisme pour calculer et extraire des contours flous à partir d'une image est fortement lié au formalisme introduit au chapitre 2 concernant le filtrage d'images d'ombres floues. En fait dans le cas d'un φ identité, si l'on adapte l'équation 5.1 pour définir un filtre flou destiné à traiter une image d'ombre non floue, l'image filtrée contient, entre autres, tous les contours des différents $C_{m_i}^p$. Dans le formalisme du chapitre 2, un tel filtre s'exprimerait grâce à une famille d'opérateurs $\Psi = \{\psi^F : \mathcal{S} \rightarrow \mathcal{S}/F \in \mathcal{F}\}$ tels que $\forall f \in \mathcal{S}, \forall p \in \Omega$:

$$\psi^F(f)(p) = \max_{h \in \mathcal{H}(f)} (\min(h(p), \nu_{\text{comp}}(fcomp(h)))) \top \max_{h \in \mathcal{H}(f)} (\min(h(p), \nu_{\text{trast}}(fctrast(F, h))))$$

avec $\mathcal{F}$ l'ensemble des images floues, $\mathcal{S}$ l'ensemble des ensembles flous définis sur le domaine Ω de l'image, $\mathcal{H}(f)$ l'ensemble des composantes connexes floues contenues dans l'ensemble f , $fcomp$ et $fctrast$ des versions floues des opérateurs de compacité et de contraste.

Toujours dans le même formalisme, le traitement d'une image d'ombres floues F se fait pour tout point $p \in \Omega$ et pour tout niveau de gris g à l'aide de l'opérateur $\delta_\Psi(F)(p, g) = \psi^{g, F}(F(*, g))(p)$. Si maintenant, on considère que l'image F est une image d'ombres non floues issue de I , c'est-à-dire que les degrés d'appartenance qu'elle contient sont soit 0 soit 1, on a pour tout niveau de gris g , $F(*, g) = X_g^+(I)$. Donc l'opérateur δ_Ψ peut dans ce cas se reformuler comme :

$$\delta_\Psi(F)(p, g) = \psi^{g, F}(X_g^+(I))(p)$$

Ainsi dans cette même configuration, comme chacun des opérateurs ψ ne travaille que sur des ensembles f nets. Leur expression peut donc se reformuler comme :

$$\psi^F(f)(p) = \begin{cases} 0 & \text{si } p \notin f \\ \nu_{\text{comp}}(\text{compacité}(\hat{\Gamma}_f^p)) \top \nu_{\text{trast}}(\text{contraste}(I, \hat{\Gamma}_f^p)) & \text{sinon} \end{cases}$$

Concrètement, cela se traduit par le fait que les composantes connexes nettes se transforment après filtrage en composantes connexes floues dont tous les points ont un même degré d'appartenance calculé de la même manière que dans l'équation 5.1. Si on remarque enfin que l'ensemble des composantes connexes considérées correspond à toutes les composantes connexes issues de tous les seuillage possible de l'image I (c.f. opérateur δ_Ψ), on s'aperçoit que l'on récupère ainsi l'union des ensembles C_m^q pour tous les couples (q, m) possibles. Ainsi le filtrage par δ_Ψ est un moyen d'extraire tous les contours des différents contours flous que l'on pourrait extraire de l'image en les associant de manière implicite à un degré d'appartenance. Pour obtenir les différents contours flous, il ne reste qu'à réintroduire les contraintes d'inclusion de q et m , ce qui peut être fait de manière directe si on considère une représentation arborescente de l'image.

5.2.3 Illustration et limitations

La figure 5.3 présente un exemple de segmentation floue par multi-seuillages d'une lésion. On remarque que tous les contours collent de manière similaire à une même partie de la lésion. En fait cette zone ne contient que peu d'imprécision. Au contraire, dans la partie en haut à droite de la lésion, les contours ont tendance à différer. Ce comportement est souhaitable dans la mesure où c'est dans cette zone que la lésion est la plus mal définie.

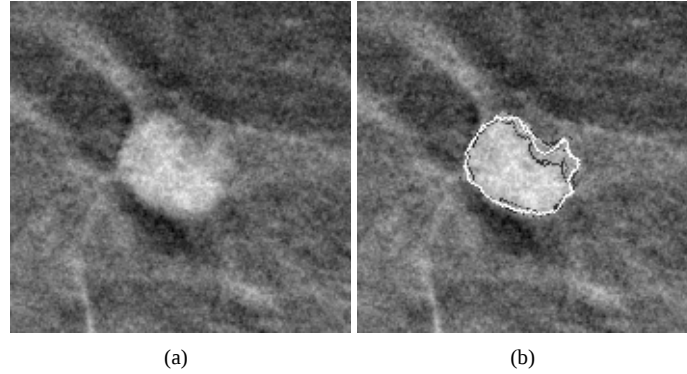


FIG. 5.3 – Segmentation par multi-seuillage d'une lésion circonscrite présentant une zone mal définie.

Le cas précédemment étudié se rapproche assez du cas idéal. En effet, la reconstruction 3D a permis de supprimer les superpositions entre cette lésion et les autres structures du sein. Malheureusement, dans certains cas la superposition n'est pas complètement supprimée, mettant ainsi à mal cette méthode de segmentation.

5.3 Formulation du problème dans le cadre des ensembles de niveaux

Les ensembles de niveaux sont communément utilisés pour la segmentation dans différents domaines d'application liés à l'image.

Dans cette section nous allons étendre une méthode de segmentation proposée par Peters *et al.* (2007), originellement destinée à la segmentation de lésions dans des images de projection. Dans un premier temps nous introduirons brièvement le formalisme de la segmentation par ensembles de niveaux, puis nous continuerons sur l'approche proposée par Peters *et al.* (2007) pour enfin finir sur l'extension de la méthode pour l'extraction de plusieurs contours.

5.3.1 Ensemble de niveaux

Le cadre des ensembles de niveaux permet d'exprimer et de résoudre le problème de segmentation de manière générale. Ce genre d'approche repose sur la représentation d'un contour par une fonction continue de Lipschitz (Osher et Sethian, 1988; Osher et Fedkiw, 2002) $\phi : \Omega \rightarrow \mathbb{R}$. Cette fonction s'interprète de la manière suivante :

$$\forall p \in \Omega \quad \begin{cases} \phi(p) > 0 & \text{si } p \text{ est à l'intérieur du contour} \\ \phi(p) = 0 & \text{si } p \text{ est sur le contour} \\ \phi(p) < 0 & \text{si } p \text{ est à l'extérieur du contour} \end{cases} \quad (5.2)$$

avec p un point appartenant au domaine de l'image Ω .

Originellement, cette fonction est une fonction de distance signée représentant de manière implicite (niveau 0) un contour. Le problème de segmentation peut se reformuler de manière générale en posant une énergie sur cette fonction, et par conséquent sur le contour représenté. Cette énergie peut contenir différents éléments comme une régularité a priori, ainsi que des termes d'attache aux données. Un exemple d'énergie sera détaillé plus loin. Cette formulation présente l'avantage d'offrir un cadre théorique rigoureux pour l'extraction du contour d'énergie minimale (Chan et Vese, 2001; Vese et Chan, 2002; Gout *et al.*, 2005).

5.3.2 Méthode originale

Dans un système de détection, Peters *et al.* (2007) a proposé de segmenter les lésions contenues dans les images de projection servant à la reconstruction. L'énergie proposée dans cette approche est un modèle hybride qui prend en compte la régularité du contour, le gradient sous ce dernier, ainsi que l'homogénéité des régions segmentées :

$$E(\phi) = \mu \int_{\Omega} \delta(\phi(p)) |\nabla \phi(p)| dp + \nu \int_{\Omega} H(\phi(p)) dp + \alpha E_{\text{région}}(\phi) + \beta E_{\text{bord}}(\phi) + \gamma E_{\text{pression}}(\phi) \quad (5.3)$$

avec H la fonction de Heavyside, δ sa dérivée (distribution de Dirac) et :

$$E_{\text{région}}(\phi) = \lambda_1 \int_{\Omega} |I(p) - \bar{c}_1|^2 H(\phi(p)) dp + \lambda_2 \int_{\Omega} |I(p) - \bar{c}_2|^2 (1 - H(\phi(p))) dp \quad (5.4)$$

avec I l'image à segmenter, $\bar{c}_1$ et $\bar{c}_2$ les valeurs de gris moyenne à l'intérieur et à l'extérieur du contour.

$$E_{\text{bord}}(\phi) = - \int_{\Omega} \delta(\phi(p)) g(|I(p)|) dp \quad (5.5)$$

où $g(|I(p)|)$ est une fonction d'arrêt destinée à ralentir le contour en proximité des zones de fort gradient. Cette fonction est donnée par :

$$g(|I(p)|) = |\nabla(G * I)(p)|$$

avec $G * I$, une version filtrée passe bas de I qui est le résultat de la convolution de I avec un noyau Gaussien. La fonction $g(|I(x, y)|)$ vaut zéro dans les régions homogènes et devient grande dans les zones de fort gradient.

Finalement, le terme de pression est donné par :

$$E_{\text{pression}}(\phi) = \int_{\Omega} \delta(\phi(p)) dp \quad (5.6)$$

5.3.3 Segmentation par ensembles de niveaux flous

En utilisant le formalisme précédemment introduit, nous allons proposer une façon d'obtenir différents contours possibles pour une même lésion. L'idée principale est d'utiliser un schéma d'évolution du contour similaire à celui proposé par Chan et Vese (2001). En effet ces derniers proposent d'utiliser des approximations de la fonction de Heavyside et du Dirac assez souples, c'est-à-dire des fonctions qui ont des valeurs non négligeables lorsque l'on s'éloigne de 0. En fait, ce genre d'approche permet de faire converger le contour aussi bien au niveau de l'interface qu'au niveau du reste de l'image. Avec un tel schéma d'évolution on peut remarquer que les contours éloignés de l'interface peuvent avoir un intérêt pour caractériser l'objet segmenté. L'idée proposée est donc de seuiller la fonction ϕ à plusieurs niveaux pour obtenir plusieurs contours candidats (c.f. figure 5.4).

Pour chaque contour, une valeur d'énergie est calculée et un degré d'appartenance à la classe *contour* en est déduit. La valeur d'énergie est obtenue en utilisant l'équation 5.3 et en ajoutant une constante à la fonction ϕ pour que le contour sélectionné se retrouve au niveau zéro. Pour calculer l'énergie, des approximations plus strictes de δ et H sont utilisées. Le candidat avec la plus faible énergie est alors considéré comme vérifiant complètement la propriété *est un contour*, et ainsi un degré d'appartenance lui est associé. Les degrés d'appartenance des contours restants sont calculés grâce à l'équation suivante :

$$\mu(\mathcal{C}) = \max(0, 1 - c * (e_{\mathcal{C}} - e_{\min})) \quad (5.7)$$

où c est une constante positive, $\mathcal{C}$ le contour considéré et $e_{\mathcal{C}}$ son énergie.

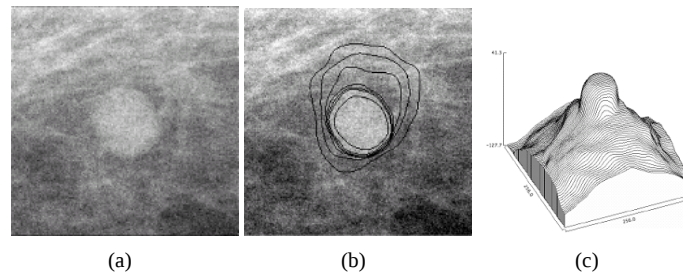


FIG. 5.4 – Ensembles de niveaux flous. L'image originale (a) est segmentée en utilisant l'énergie de l'équation 5.3. Des coupes d'iso-niveau sont extraites à partir de la fonction ϕ (c) pour obtenir un ensemble de contours possibles (b).

5.3.4 Limitations

La méthode précédemment présentée a plusieurs limites. Tout d'abord, la minimisation d'énergie se fait par des méthodes de descente de gradient qui n'interdisent pas de se retrouver dans des minima locaux. Remarquons que ce comportement peut dans certains cas être souhaité car rien ne nous assure que l'énergie soit optimale pour tous les contours. De plus cette même énergie fait référence à plusieurs constantes qui peuvent influencer fortement le résultat de la segmentation. Ces dernières sont assez délicates à optimiser si on souhaite segmenter des structures de formes assez variées. On verra au chapitre 6 une astuce pour pallier en partie cette limitation lorsque l'on souhaite différencier les lésions circonscrites des lésions spiculées. Enfin on remarquera que l'utilisation de l'astuce d'extraire plusieurs contours à partir de la fonction ϕ après convergence peut être discutable. En effet, on utilise juste un effet de bord de l'étape de minimisation pour extraire plusieurs contours.

5.4 Approche par programmation dynamique

L'approche que nous allons proposer maintenant repose sur la formulation du problème de segmentation par un chemin dans un graphe. Dans un premier temps, nous introduirons le formalisme originellement proposé pour la segmentation de lésions en mammographie standard par Timp et Karssemeijer (2004). Nous détaillerons ensuite notre proposition d'extension de cette approche pour obtenir des contours flous. Nous finirons par une évaluation des performances des méthodes nette et floue sur des lésions provenant de coupes de tomosynthèse.

5.4.1 Travaux existants

Représentation de l'image

La première étape dans ce que proposent Timp et Karssemeijer (2004) est de changer de repère de représentation de l'image. L'idée est de passer en coordonnées polaires en se servant d'un centre estimé de la structure à segmenter. L'idée derrière cette représentation est de considérer que la structure à segmenter a une forte compacité. Ainsi dans le cadre idéal le contour dans la représentation polaire délimitant un disque dans le domaine cartésien est une droite. La figure 5.5 illustre la conversion d'une telle image.

La correspondance n'est pas exacte entre les pixels dans les deux images, il est donc nécessaire de procéder à des interpolations pour pouvoir convertir une image dans le domaine polaire. Ce point constitue une faiblesse de l'approche puisque il y a perte potentielle d'information à ce niveau. Néanmoins, avec un pas d'échantillonnage fin pour les angles, on peut limiter la perte d'information pour les points qui sont éloignés du centre de la lésion au détriment d'une redondance d'information pour les zones proches de ce même centre.

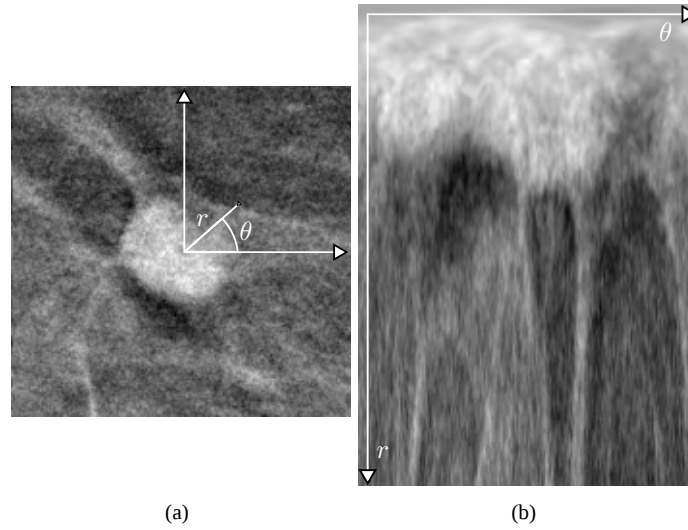


FIG. 5.5 – Représentation en domaine polaire d'une image contenant une lésion circonscrite.

Notion de chemin et de coût associé

Une fois dans le domaine polaire, l'idée est représenter un contour possible par un chemin. En associant à tous les contours possibles, c'est-à-dire à tous les chemins reliant la première colonne à la dernière colonne de l'image à segmenter, un coût modélisant la qualité du chemin, on pose un cadre permettant la résolution de notre problème de segmentation. Un tel coût s'écrit de manière générique de la façon suivante :

$$\mathcal{C}(c) = \sum_{\theta \in [0; \theta_{\max}]} M(\theta, c(\theta)) \quad (5.8)$$

avec $c(\theta)$ le rayon associé à l'angle θ pour le chemin c considéré et M une matrice de coûts. Les indices de la matrice correspondent aux coordonnées dans l'image polaire. Pour pouvoir faire le tour de la lésion nous considérons un intervalle d'angle allant de 0 à 2π . Ces angles étant échantillonnés, nous notons $\theta_{\max}$ l'indice d'angle dans M correspondant à 2π .

Timp et Karssemeijer (2004) proposent une matrice de coûts M composée de trois termes : un terme d'attache aux contours de l'image lié au gradient de cette dernière, un terme de niveau de gris préférentiel pour le contour et un terme de taille a priori. Dans la mesure où ce troisième terme est discutable puisque les lésions peuvent avoir des tailles très variables, nous l'avons omis dans nos expérimentation. Notre formule de coût s'écrit donc comme :

$$M(\theta, r) = \beta g(\theta, r) + (1 - \beta)n(\theta, r)$$

avec $\beta \in [0; 1]$ un paramètre de pondération entre le terme lié au gradient (g) et le terme lié au niveau de gris du contour (n).

Le terme lié au gradient a été obtenu par une simple convolution de l'image en coordonnées polaires avec une dérivée de gaussienne dans la direction des rayons. Cela revient à considérer que la lésion est assez compacte et que par conséquent le gradient dans cette direction équivaut au gradient dans la direction normale au contour. Un exemple d'une carte de coût obtenue par ce procédé est illustré à la figure 5.6(a).

Pour le terme lié à l'intensité des pixels par lesquels passe le contour, nous avons utilisé un terme similaire à ce que proposent Timp et Karssemeijer (2004) :

$$n(\theta, r) = \sqrt{|I(\theta, r) - \mu|}$$

où I est l'image en coordonnée polaire et $\mu = \alpha\mu_{int} + (1 - \alpha)\mu_{ext}$ avec respectivement μ_{int} et μ_{ext} les moyennes à l'intérieur et à l'extérieur de la lésion et $\alpha \in [0; 1]$ un facteur de pondération. Originellement, Timp et Karssemeijer (2004) proposent deux méthodes pour calculer μ_{int} et μ_{ext} , l'une reposant

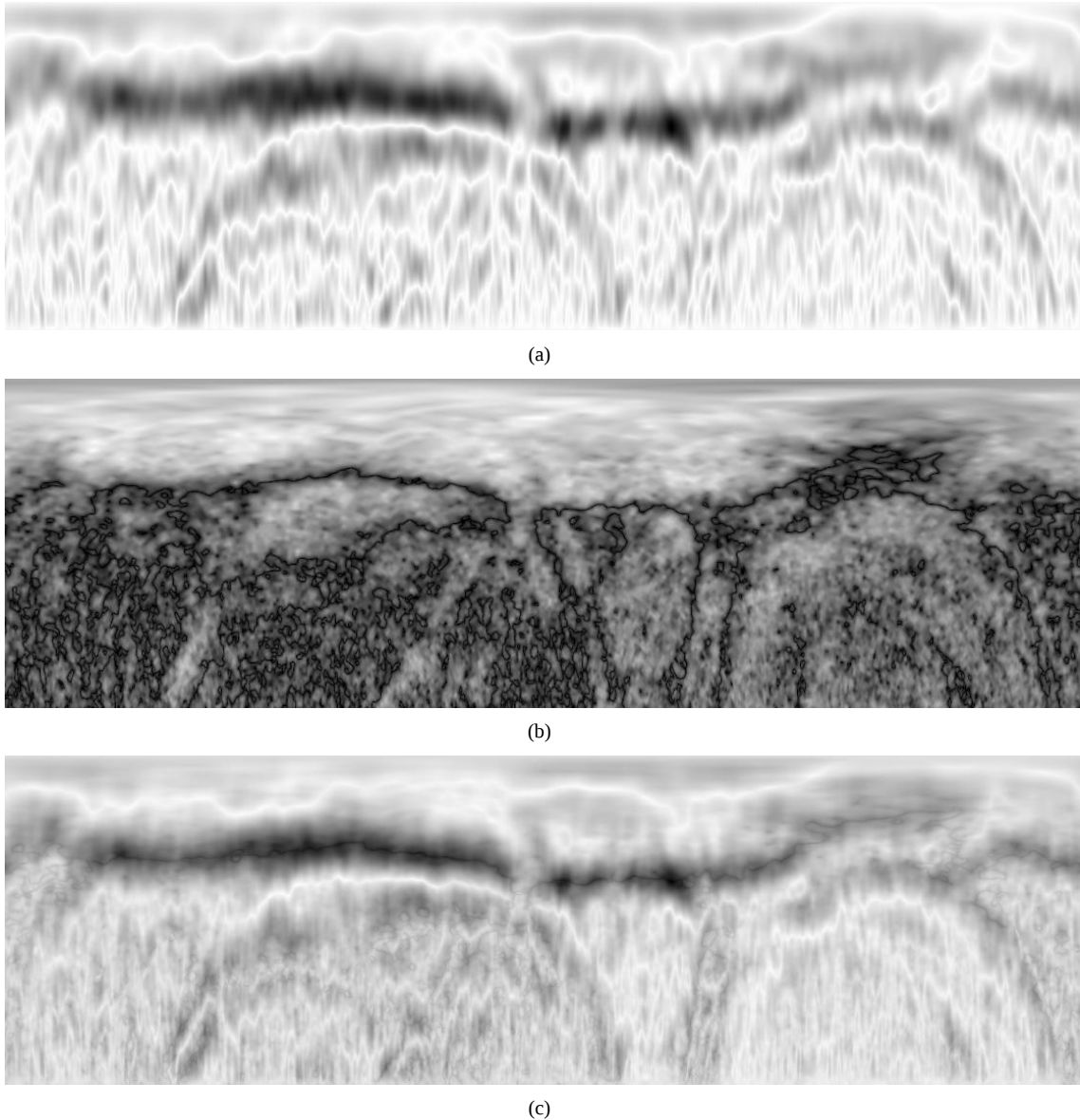


FIG. 5.6 – Exemple de calcul d’une matrice de coûts M . (a) Coût correspondant au gradient. (b) Coût correspondant au niveau de gris du contour. (c) Matrice de coûts M .

sur la taille des lésions et l’autre sur l’analyse d’histogramme de l’image. La contrainte introduite par la première approche n’étant pas acceptable nous ne l’avons pas utilisée. La deuxième approche modélise l’histogramme de l’image par un mélange de gaussiennes. Nous avons testé cette hypothèse sur des régions d’intérêt issues de coupes de tomosynthèse pour voir si elle était invalide sur ces données. Malheureusement, cela ne semblait pas être le cas. Nous avons donc opté pour une méthode à deux itérations : dans la première étape nous considérons que μ_{int} et μ_{ext} sont tous les deux égaux à la moyenne dans l’image, puis les deux moyennes issues de cette première segmentation sont enfin utilisées pour la deuxième étape. Une autre solution est d’utiliser l’information issue de l’étape de marquage. En effet comme on l’a vu au chapitre 2, les marqueurs que l’on peut obtenir avec des filtre connexes flous ont la forme des structures détectées, ainsi on peut utiliser cette information pour calculer μ_{int} et μ_{ext} . Un exemple de matrice de coûts lié à l’intensité du pixel est présentée à la figure 5.6(b)

Domínguez et Nandi (2007) proposent des raffinements par rapport à la méthode introduite par

Timp et Karssemeijer (2004). Par exemple ils suggèrent le remplacement du terme lié au gradient par un terme reposant sur la cohérence intrinsèque (Rao et Schunck, 1989; Mudigonda *et al.*, 2001). Cela a été testé en pratique sur nos données de tomosynthèse sans grande amélioration par rapport à la méthode simple.

Contrairement aux contours actifs, l'équation 5.8 ne contient pas de terme de régularité. Pour réintroduire cette contrainte et ainsi avoir des contours réguliers, on peut s'interdire de considérer les chemins trop irréguliers. Ainsi on définit l'ensemble des contours réguliers comme :

$$\mathcal{P} = \{c/\forall \theta \in [0; 2\pi] \quad |c(\theta) - c(\theta - 1)| \leq f\}$$

avec $f \in \mathbb{N}$ le paramètre permettant de jouer sur la régularité.

On remarquera que le choix de f est lié à la quantification des angles et des rayons. En effet si, pour un f constant, on augmente la résolution des angles on s'autorisera à considérer des contours plus irréguliers. Inversement si on augmente la quantification des rayons, on aura des contours plus réguliers.

Le problème de segmentation se formule donc comme suit :

$$\hat{c} = \arg \min_{c \in \mathcal{P}} \mathcal{C}(c)$$

L'approche proposée pour introduire de la régularité n'est pas la seule possible. On pourrait tout aussi bien modifier l'équation 5.8 pour y introduire un terme additif correspondant à la longueur du contour. Dans ce cas, la régularité du contour serait introduite de manière plus globale, c'est-à-dire que si le contour est très régulier sur une grande partie, il peut varier fortement à un endroit localisé. Dans l'approche proposée par Timp et Karssemeijer (2004), on force une certaine régularité de manière plus locale.

Résolution du problème par programmation dynamique

Pour résoudre le problème précédent, on peut considérer un graphe et ainsi utiliser les algorithmes de recherche de coûts minimaux (Timp et Karssemeijer, 2004). Le coût du chemin étant une addition de coûts locaux, on peut de manière efficace trouver le chemin optimal en utilisant les concepts de programmation dynamique. Pour cela il faut calculer une matrice de coûts cumulés :

$$\begin{cases} \overline{M}(0, r) = c(0, r) \\ \overline{M}(\theta + 1, r) = \min_{-f \leq l \leq f} \{ \overline{M}(\theta, r + l) + M(\theta + 1, r) \} \end{cases} \quad (5.9)$$

où f est toujours le paramètre utilisé pour l'introduction de régularité dans les contours.

La dernière colonne de la matrice $\overline{M}$ correspond aux coûts minimaux des chemins se terminant à différent rayons. Ainsi en prenant celui de coût minimal, on peut refaire le parcours inverse et trouver les points du chemin optimal.

Un dernier point abordé par Timp et Karssemeijer (2004) est le fait de s'assurer d'avoir un contour fermé. En effet, le formalisme précédent n'assure en rien que les points de départ et d'arrivée du chemin coïncident. L'idée que les auteurs proposent est de répliquer la matrice de coûts M pour contraindre les extrémités des chemins optimaux à se correspondre. Une fois la matrice de coûts répliquée, on peut utiliser l'équation 5.9 pour obtenir le chemin de coût minimal dans la matrice répliquée dont seule la partie centrale sera considérée pour l'extraction du contour. Formellement l'équation 5.9 se réécrit de la manière suivante :

$$\begin{cases} \overline{M}(0, r) = c(0, r) \\ \overline{M}(\theta + 1, r) = \min_{-f \leq l \leq f} \{ \overline{M}(\theta, r + l) + M((\theta + 1) \bmod N, r) \} \end{cases} \quad (5.10)$$

où N représente la largeur de la matrice M .

Le principe de matrice de coûts cumulés étendue est présenté à la figure 5.7. Dans cet exemple, on voit dans la matrice de coûts (c.f. figure 5.7(a)) que le chemin optimal (en noir) ne voit pas ses points de départ et de fin correspondre. Cela se vérifie sur la matrice de coûts cumulés (c.f. figure 5.7(b)). Si on utilise l'astuce de la matrice étendue (c.f. figure 5.7(c)), on peut limiter cet effet. En effet comme on peut le voir, les extrémités de la partie centrale de cette matrice sont influencées par les répliquations qui la précèdent et qui la suivent rendant ainsi le contour plus vraisemblable.

Des résultats de segmentations seront présentés à la section 5.4.3.

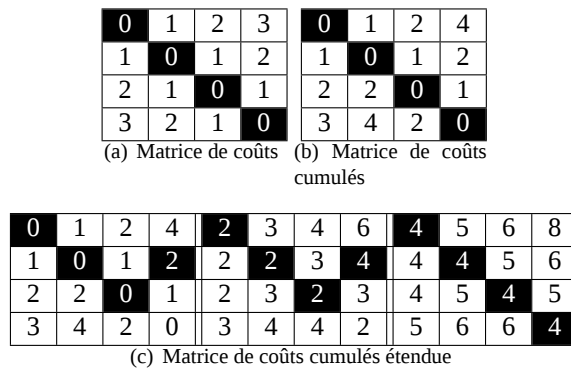


FIG. 5.7 – Illustration du principe de matrice de coûts cumulés. Les chemins optimaux apparaissent en noir.

5.4.2 Obtentions de contours flous

Nous allons maintenant étudier différentes façons d'étendre la méthode précédente dans l'optique d'obtenir des contours flous.

Pénalisation

Une première idée pour extraire plusieurs contours à partir de la matrice de coûts M est de la modifier en mettant les coûts des points par lesquels passent les contours déjà trouvés à $+\infty$. On peut ainsi obtenir plusieurs contours en itérant le processus tant qu'il existe un coût de chemin non infini reliant la première et la dernière colonne de la matrice de coûts. Ce genre d'approche n'empêche néanmoins pas l'obtention de contours inclus les uns dans les autres comme il est requis par Peters (2007). En effet, comme on peut le voir à la figure 5.8, un paramètre f supérieur à 0 permet de telles configurations.

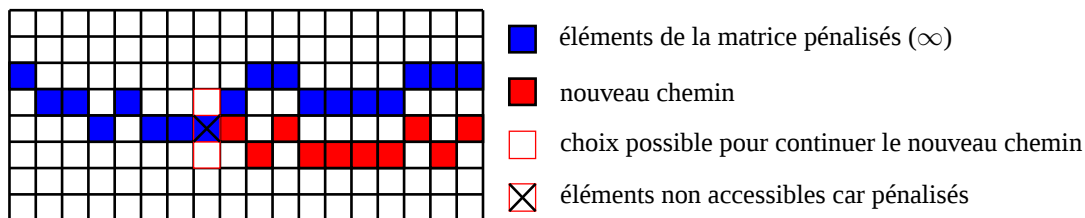


FIG. 5.8 – Exemple de pénalisation n'empêchant pas le croisement de deux contours.

Une solution pour forcer l'inclusion est de pénaliser une bande assez large autour du meilleur contour (c.f. figure 5.9). La largeur de cette bande est bien évidemment liée à la grandeur du paramètre f . Cette deuxième approche a aussi l'avantage de fournir des contours qui sont plus différents les uns des autres. Néanmoins dans le cas où le contour optimal passe dans une vallée plus large que la bande de pénalisation, on peut observer un effet de dilatation ou d'érosion : les différents contours obtenus sont semblables au meilleur contour à une échelle près. L'algorithme 5.1 résume ce processus d'extraction de contours.

Distance entre contours

Dans le but d'obtenir des contours qui portent une information différente, on peut procéder par sélection de contours significatifs par différence. L'idée est de ne considérer un nouveau contour que si il diffère assez des contours déjà retenus. Toute la difficulté repose donc sur la définition de différence entre contours. Nous

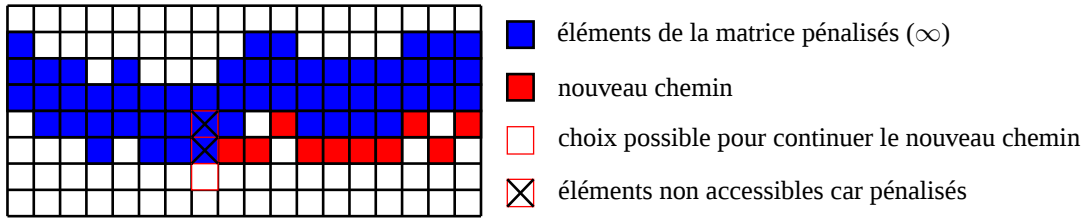


FIG. 5.9 – Exemple de pénalisation par bande empêchant le croisement de deux contours.

Algorithme 5.1 : Extraction de différents contours par la méthode de pénalisation par bande.

```

soit  $E = \emptyset$  un ensemble de contours ;
calculer la matrice  $M$  ;
tant que cirière d'arrêt non atteint faire
    calculer la matrice cumulée  $\bar{M}$  en utilisant l'équation 5.10 ;
    extraire le meilleur contour  $c$  à partir de  $\bar{M}$  ;
    pour chaque  $\theta \in [[0; \theta_{\max}]]$  faire
        pour chaque  $b \in [[-f; f]]$  faire
             $M(\theta, c(\theta) + b) = \infty$  ;
        fin
    fin
     $E = E \cup \{c\}$  ;
fin
retourner  $E$ 

```

proposons de définir celle-ci de la manière suivante :

$$d(c_1, c_2) = \max_{\theta \in [[0; \theta_{\max}]]} (|c_1(\theta) - c_2(\theta)|)$$

En utilisant un seuil ($d_{\min}$) sur cette mesure, on définit donc une procédure pour sélectionner différents contours possibles pour une même structure. L'approche est résumée à l'algorithme 5.2.

Degré d'appartenance

Pour obtenir des contours flous, il est nécessaire d'associer un degré d'appartenance aux différents contours (c.f. algorithme 5.2). Une première approche pour faire cela est d'utiliser le coût des différents contours extraits. Le degré d'appartenance d'un contour c devient donc proportionnel à $\mathcal{C}(c)$:

$$\mu_c = \frac{\mathcal{C}(c) - C_{\max}}{C_{\max} - C_{\min}} \quad (5.11)$$

avec $C_{\max}$ et $C_{\min}$ le maximum et le minimum des coûts des contours extraits.

En pratique on remarquera que ces coûts sont fortement liés à l'a priori introduit dans la méthode de segmentation. Une autre méthode possible est d'associer de forts degrés d'appartenance aux zones où beaucoup de contours similaires peuvent être extraits. Pour ce faire nous proposons de construire une carte de densité D de contours. Cette dernière est obtenue en faisant l'union des K premiers contours que l'on peut obtenir à partir de la matrice de coûts M . On notera l'ensemble de ces contours E . Cette carte de densité est ensuite filtrée à l'aide d'un filtre passe bas et le degré d'appartenance d'un contour c s'obtient grâce à l'équation suivante :

$$\mu_c = \frac{\sum_{\theta=0}^{\theta_{\max}} D(\theta, c(\theta)) - \min_{c' \in E} \sum_{\theta=0}^{\theta_{\max}} D(\theta, c'(\theta))}{\max_{c' \in E} \sum_{\theta=0}^{\theta_{\max}} D(\theta, c'(\theta)) - \min_{c' \in E} \sum_{\theta=0}^{\theta_{\max}} D(\theta, c'(\theta))} \quad (5.12)$$

Algorithme 5.2 : Extraction de différents contours par la méthode des distances.

```

soit  $E = \emptyset$  un ensemble de contours ;
calculer la matrice  $M$  ;
calculer la matrice cumulée  $\overline{M}$  en utilisant l'équation 5.10 ;
extraire le meilleur contour  $c$  à partir de  $\overline{M}$  ;
pour chaque  $\theta \in [[0; \theta_{\max}]]$  faire
     $M(\theta, c(\theta)) = \infty$  ;
fin
 $E = \{c\}$  ;
tant que critère d'arrêt non atteint faire
    répéter
        calculer la matrice cumulée  $\overline{M}$  en utilisant l'équation 5.10 ;
        extraire le meilleur contour  $c'$  à partir de  $\overline{M}$  ;
        pour chaque  $\theta \in [[0; \theta_{\max}]]$  faire
             $M(\theta, c; (\theta)) = \infty$  ;
        fin
        jusqu'à  $d(c, c') \leq d_{\min}$  ;
         $c = c'$  ;
         $E = E \cup \{c\}$  ;
    fin
retourner  $E$ 

```

Cette expression permet de juger à quel point le contour considéré passe dans des zones denses de la carte de densité. Ainsi lorsque le contour possède la plus forte densité cumulée parmi tous les contours considérés, son degré d'appartenance va se rapprocher de 1 et inversement, lorsqu'il passe par des zones de faible densité, son degré va tendre vers 0.

Remarquons que le paramètre K doit être assez important pour obtenir une carte de densité D représentative. Un exemple de calcul d'une telle carte est proposé à la figure 5.10.

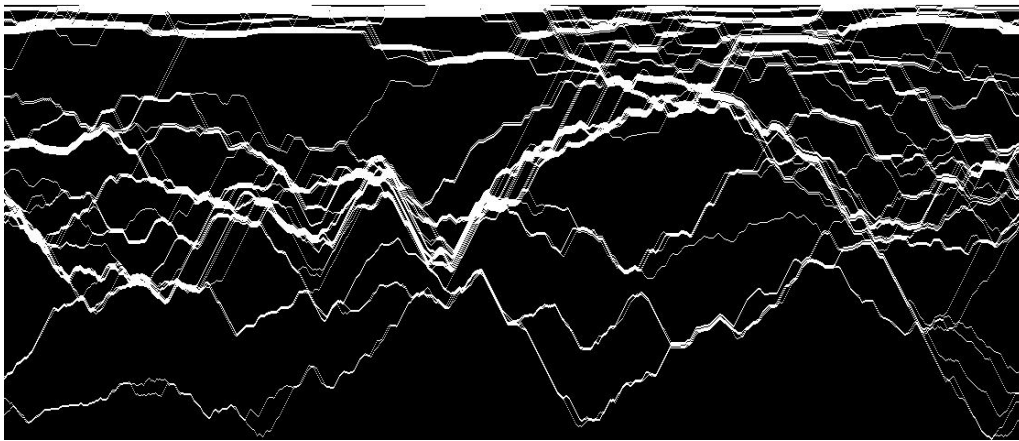
Extraction de l'incertitude et de l'imprécision

Il faut faire attention à ne pas comparer directement les méthodes de pénalisation par bande et de distance. En effet ces deux approches sont complémentaires et nullement en compétition. La première méthode permet d'extraire des contours emboîtés traduisant l'imprécision possible d'une segmentation. La deuxième méthode permet quant à elle d'extraire différentes segmentations possibles traduisant l'incertitude quant à la localisation de la lésion. Pour tirer parti des deux approches, il faut un moyen d'extraire un contour flou répondant aux critères de la première méthode pour chacun des contours extraits par la seconde méthode. Ainsi on se ramène au principe exposé à la section 5.1 qui évoque l'extraction d'un contour flou pour chaque lésion possible. Pour se faire on peut combiner de manière directe les deux approches, c'est-à-dire en utilisant l'algorithme 5.1 dans l'algorithme 5.2 : pour chaque contour possible issu de l'algorithme 5.2, on extrait un contour flou en utilisant l'algorithme 5.1. Cette approche est formalisée à l'algorithme 5.3. Le calcul des degrés d'appartenance peut se faire de manière indépendante en utilisant l'une des deux méthodes proposées précédemment. Cependant, comme on le verra dans la présentation des résultats, la méthode par carte de densité donne de meilleurs résultats.

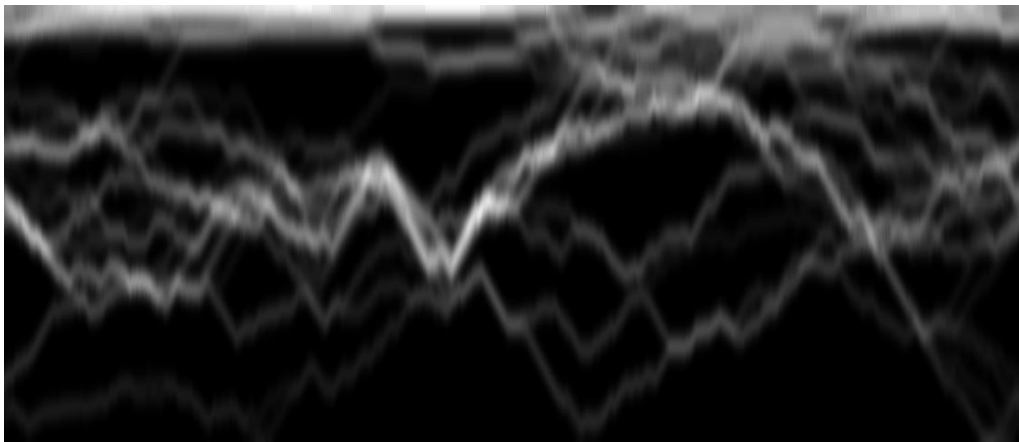
5.4.3 Résultats**Données**

Les méthodes de segmentation nette et floue reposant sur le formalisme de chemins minimaux ont été évaluées sur une base de données contenant des régions d'intérêt centrées sur des lésions circonscrites et spiculées issues de coupes reconstruites de volumes de tomosynthèse du sein. Les lésions spiculées étaient au nombre de 32, et les lésions circonscrites au nombre de 16.

Pour chaque image, un contour de référence a été tracé pour faciliter une comparaison objective des performances de la segmentation nette. La constitution d'une base contenant plusieurs contours est une



(a)



(b)

FIG. 5.10 – Calcul d’une carte de densité. (a) Accumulation des K meilleurs contours. (b) Carte de densité après filtrage.

Algorithme 5.3 : Extraction de contours flous à partir d’une image.

```

soit cf un ensemble de contours flous;
calculer la matrice  $M$  ;
obtenir l’ensemble de contour  $E_1$  en utilisant l’algorithme 5.2 ;
pour chaque  $c \in E_1$  faire
    soit  $M'$  la matrice  $M$  où  $c$  est pénalisé sur une bande de pixels ;
    soit  $E_2$  l’ensemble des contours obtenus en utilisant l’algorithme 5.1 sur  $M'$  ;
     $E_2 = E_2 \cup \{c\}$ 
    pour chaque  $c' \in E_2$  faire
        calculer le degré d’appartenance  $\mu$  de  $c'$  ;
        insérer  $(c', \mu)$  dans cf en respectant l’ordre d’inclusion des contours ;
    fin
fin
retourner cf

```

tâche plus délicate : lorsque l’on trace un premier contour, ce dernier nous influence pour trouver les suivants à moins de laisser s’écouler assez de temps pour l’oublier. Pour cette raison l’évaluation des contours flous est beaucoup plus qualitative.

Segmentation nette

La méthode de Timp et Karssemeijer (2004) étant originellement proposée pour la segmentation de lésions en mammographie, nous avons voulu l'évaluer sur des images issues de coupes de volumes de tomosynthèse de seins. Pour comparer les résultats de segmentation aux contours vérité à disposition, nous avons utilisé le même critère que celui proposé par Domínguez et Nandi (2007). Ce critère est en fait une composition de trois critères :

$$\begin{aligned} P_1 &= \frac{|A \cap R|}{|A \cup R|} && \text{mesure de similarité} \\ P_2 &= \frac{|R \setminus (A \cap R)|}{|R|} && \text{mesure de sous-segmentation} \\ P_3 &= \frac{|A \setminus (A \cap R)|}{|A|} && \text{mesure de sur-segmentation} \end{aligned}$$

où $|\cdot|$ est le cardinal d'un ensemble, et A et R correspondent à l'ensemble des pixels contenus respectivement dans le contour à évaluer et dans le contour de référence.

Ces trois mesures sont ensuite moyennées pour obtenir un critère global de performance de segmentation :

$$P = \frac{P_a + 1 - P_2 + 1 - P_3}{3}$$

Rappelons que l'approche dépend de plusieurs paramètres. Nous avons considéré les trois paramètres définissant l'a priori sur les structures recherchées, c'est-à-dire α , β et f . Les paramètres d'échantillonnages étant liés au paramètre f , ces derniers sont restés fixes. Les performances ont été évaluées en utilisant un procédé de validation croisée : pour chaque élément de la base de données, les éléments restants ont été utilisés pour apprendre les paramètres α , β et f . Les deux premiers sont obtenus par recherche exhaustive sur une grille de 110 couples $((\alpha, \beta))$ de celui qui optimise le critère P moyen sur la base d'apprentissage et le paramètre f a été obtenu à l'aide des contours de vérité.

Le tableau 5.1 résume les performances de segmentation pour différentes configurations d'apprentissage. Ainsi on remarquera que lorsque l'on utilise que le sous-ensemble de la base de données ne contenant que les lésions circonscrites, on obtient un très bon critère de performance moyen (0,87). Cela peut s'expliquer par le fait que la méthode de segmentation utilisée fait l'hypothèse d'une lésion compacte. Dans le cas où l'on ne considère que les lésions spiculées, l'indice de performance est moins bon (0,71). Néanmoins, il faut remarquer que c'est ce genre de lésions qui présente généralement le plus d'incertitude. Ainsi la segmentation de référence est beaucoup plus sujette à controverse, et c'est sur ce type de lésions que l'approche par contour flou est la plus intéressante. Remarquons enfin que lorsque on utilise les paramètres optimaux appris sur les lésions spiculées pour segmenter des lésions circonscrites, les performances ne sont que légèrement dégradées. Ainsi il semble raisonnable d'utiliser ce jeu de paramètres dans un système de détection n'utilisant pas les contours flous.

Base d'apprentissage/de test	Similarité	Sous-segmentation	Sur-segmentation	Critère de performance
Lésions cir./cir.	0,80	0,10	0,11	0,87
Lésions spi./spi.	0,59	0,32	0,10	0,71
Lésions spi./cir.	0,75	0,10	0,17	0,83

TAB. 5.1 – Performances de la segmentation nette par extraction de chemins minimaux.

La figure 5.11 illustre les résultats de segmentation pour une lésion circonscrite. On remarquera que visuellement la qualité du contour obtenu avec des paramètres optimaux pour les lésions spiculées est un peu moins bonne que lorsque l'on apprend les paramètres sur des lésions circonscrites. Néanmoins la qualité du contour obtenu est relativement correcte.

La figure 5.12 illustre aussi la comparaison entre les contours obtenus en apprenant les paramètres de la méthode de segmentation sur différents types de lésions, mais cette fois ci sur une lésion spiculée. On remarque que lorsque l'on se sert des lésions circonscrites pour l'apprentissage, la qualité du contour est grandement dégradée.

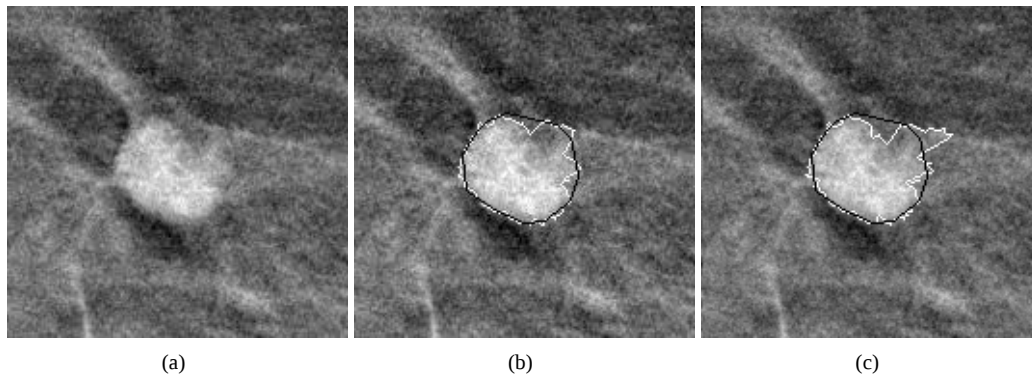


FIG. 5.11 – Exemple de segmentation d’une lésion circonscrite (a) avec des paramètres optimaux pour des lésions circonscrites (b) et spiculées (c). Le contour en noir est celui de référence.

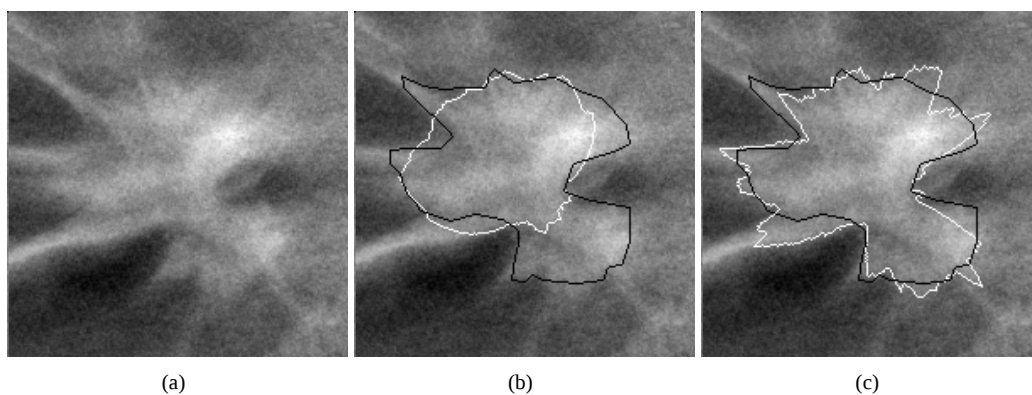


FIG. 5.12 – Exemple de segmentation d’une lésion spiculée (a) avec des paramètres optimaux pour des lésions circonscrites (b) et spiculées (c). Le contour en noir est celui de référence.

Les remarques d’ordre visuel faites sur les figures 5.11 et 5.12 vont dans le même sens que les mesures plus objectives présentées au tableau 5.1.

Segmentation floue

La figure 5.13 présente un exemple de segmentation d’une lésion circonscrite par la méthode d’extraction de contours flous par pénalisation par bande. Les degrés d’appartenance sont calculés à l’aide de l’équation 5.11. On remarquera qu’un seul contour est vraiment considéré comme représentatif, ce qui est cohérent avec l’image segmentée : la lésion considérée est assez bien définie.

Néanmoins, lorsque la lésion présente une plus grande incertitude, cette méthode de pénalisation par bande peut ne pas être adaptée. En effet si on considère l’image synthétique présenté à la figure 5.17(h), on ne peut espérer pouvoir segmenter les deux disques simultanément. Cette limitation vient du fait que la pénalisation par bande implique l’inclusion des contours constituant le contour flou. De manière plus haut niveau, la méthode de pénalisation par bande permet d’extraire en ensemble de contours traduisant essentiellement de l’imprécision, alors que dans l’exemple qui pose problème, on est devant de l’incertitude.

En regardant comment se comporte la méthode des distances sur ce même exemple de la figure 5.14, on s’aperçoit que cette méthode manipule en fait l’incertitude contenue dans cette image. Ainsi on voit que les deux méthodes ne sont pas en compétition mais plutôt complémentaires. La figure 5.15 illustre cette même représentation de l’incertitude sur une lésion spiculée.

Les figures 5.16 et 5.17 illustrent enfin l’apport du calcul de degré d’appartenance grâce à une cadre de densité comme proposé à l’équation 5.12. En effet comme on peut le voir sur cet exemple, l’ordre dans

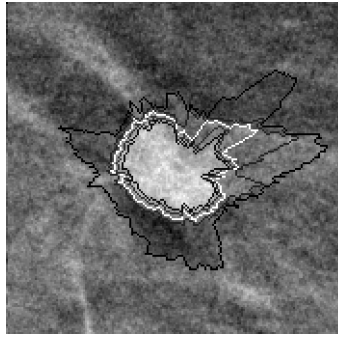


FIG. 5.13 – Segmentation floue par la méthode de pénalisation par bande. L'intensité des contours est proportionnelle au degré d'appartenance.

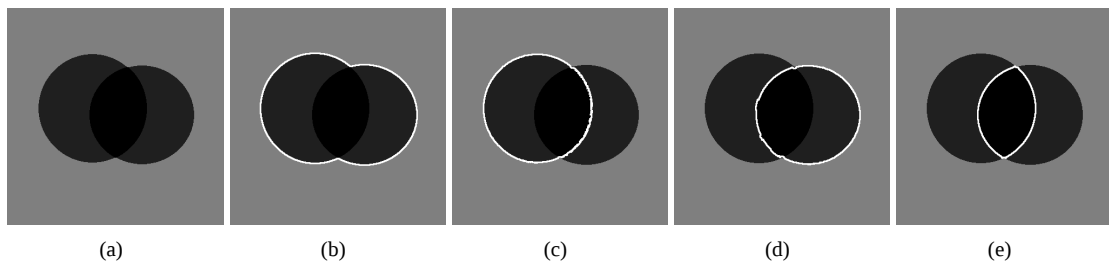


FIG. 5.14 – Image synthétique illustrant la manipulation de l'incertitude par l'algorithme 5.2. (a) Image originale. (b) Premier contour extrait. (c) Deuxième contour extrait. (d) Troisième contour extrait. (e) Quatrième contour extrait.

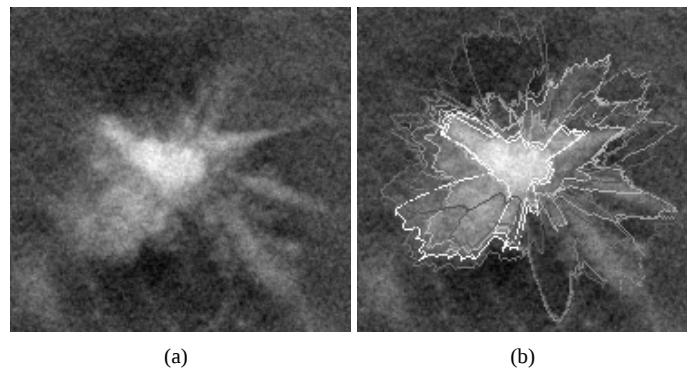


FIG. 5.15 – Exemple de segmentation floue d'une lésion spiculée comportant de l'incertitude.

lequel on trouve les contours n'est pas forcément est bon indicateur de la qualité de ces derniers. Ainsi le contour le plus mauvais selon l'équation 5.11 se retrouve le meilleur si l'on utilise l'équation 5.12. De manière visuelle, ce contour correspond au noyau dense de la lésion, et à ce titre peut avoir du sens. Cependant on remarquera que le fait de le considérer comme le meilleur contour est à relativiser : intuitivement, on pourrait avoir tendance à ajouter les spicules principales. Cela constitue une piste à explorer dans le but d'améliorer ce calcul de degré d'appartenance.

La figure 5.18 illustre enfin l'algorithme 5.3. Cet exemple reprend l'image synthétique de la figure 5.14(a) pour montrer l'extraction de plusieurs contours flous pour une même image. En pratique on extrait quatre contours flous correspondant aux quatre contours extraits dans la figure 5.14. Les contours des figures 5.18(b) et 5.18(c) contiennent bien respectivement les deux cercles constituant l'image synthétique.

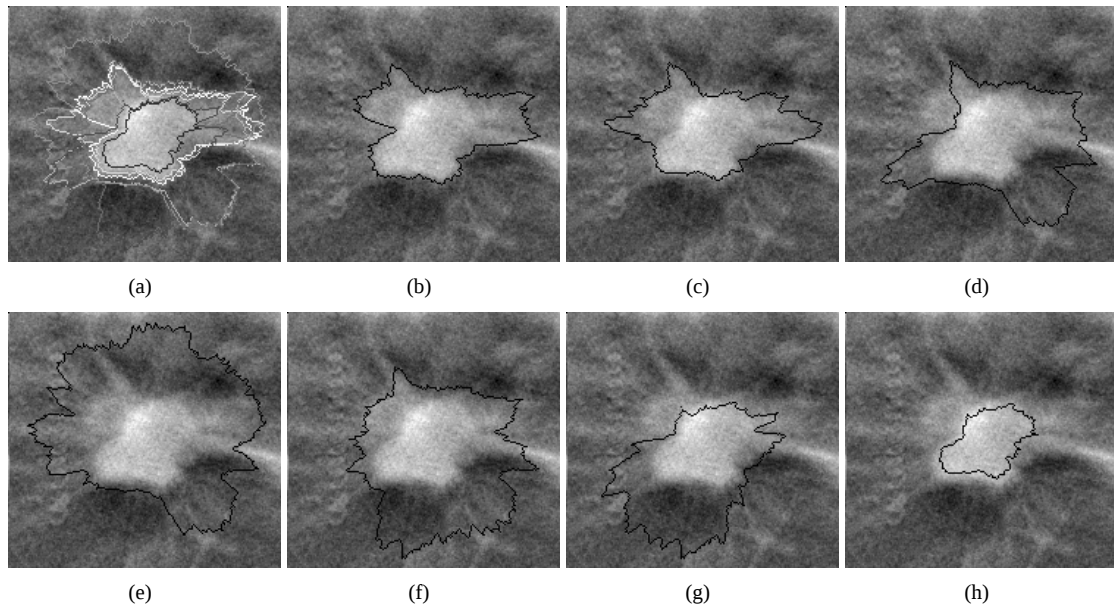


FIG. 5.16 – Illustration du calcul des degrés d'appartenance à chacun des contours extraits par la méthode des distances en utilisant l'équation 5.11. La figure 5.16(a) représente tous les contours avec un niveau de gris représentatif de leur degré d'appartenance. Les autres figures illustrent de manière isolée les différents contours par ordre de degré d'appartenance décroissant.

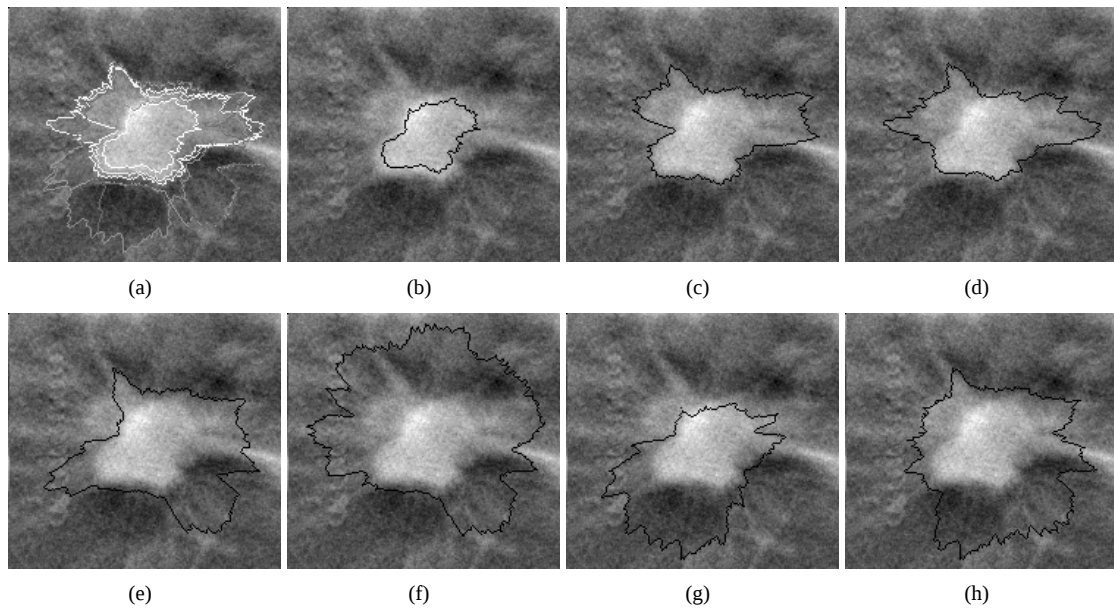


FIG. 5.17 – Illustration du calcul des degrés d'appartenance à chacun des contours extraits par la méthode des distances en utilisant l'équation 5.12. La figure 5.17(a) représente tous contours avec un niveau de gris représentatif de leur degré d'appartenance. Les autres figures illustrent de manière isolée les différents contours par ordre de degré d'appartenance décroissant.

L'imprécision se traduit par l'extraction dans ces deux cas de deux autres contours : l'un englobant les deux cercles et le second délimitant l'intersection des deux cercles. Les contours flous des figures 5.18(a)

et 5.18(d) sont assez similaires. En fait, l'algorithme 5.3 peut induire une certaine redondance dans les contours flous extraits.

Remarquons enfin que l'on pourrait être tenté de rechercher deux contours de plus qui correspondraient aux deux croissants de lune formés par les parties non superposées des disques. De tels contours ne peuvent pas être obtenus grâce à cette méthode de segmentation car le centre de la lésion a priori doit obligatoirement être contenu dans chaque contour potentiel.

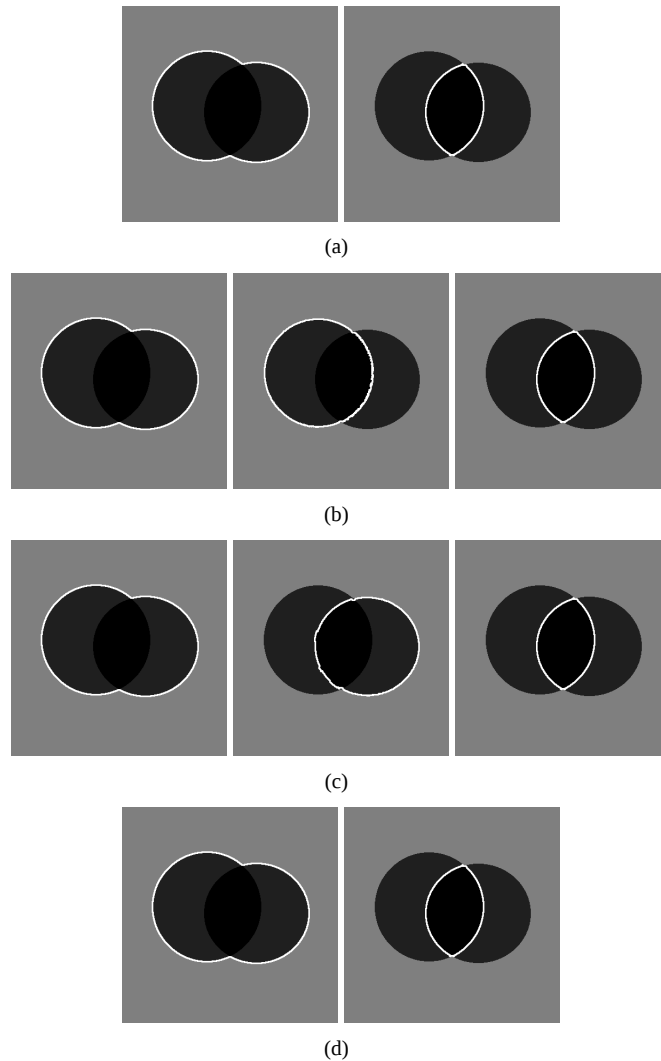


FIG. 5.18 – Exemple d'extraction de plusieurs contours flous en utilisant l'algorithme 5.3 sur l'image de la figure 5.14(a). Seuls les contours représentatifs des contours flous extraits sont montrés. (a) Contour flou correspondant au contour de la figure 5.14(b). (b) Contour flou correspondant au contour de la figure 5.14(c). (c) Contour flou correspondant au contour de la figure 5.14(d). (d) Contour flou correspondant au contour de la figure 5.14(e).

5.5 Conclusions et perspectives

Dans ce chapitre, nous avons abordé le problème de la segmentation de lésion en tomosynthèse du sein. En plus des méthodes classiques de segmentation, nous avons abordé le problème de la manipulation de l'imprécision et de l'incertitude qui peuvent être présente dans les images à segmenter. Pour cela nous nous

sommes appuyé sur le formalisme des contours flous (Peters, 2007).

Nous avons étudié une méthode d'extraction de contours flous reposant sur différents seuillages de l'image en appuyant sur le lien avec le formalisme des filtres connexes flous proposés au chapitre 2. Nous avons aussi étudié l'extension d'une méthode de segmentation reposant sur les ensembles de niveaux. Enfin, en dernière contribution sur les contours flous nous avons proposé l'extension d'une méthode de segmentation par programmation dynamique pour obtenir une segmentation floue.

La version originale (non floue) de cette dernière méthode de segmentation étant dédiée à la segmentation de lésions en mammographie standard, nous l'avons évaluée sur des lésions issues de coupes reconstruites de tomosynthèse. Les résultats obtenus laissent penser que cette méthode de segmentation nette est viable pour l'utilisation dans une chaîne de détection automatique d'opacités, cela sera illustré au chapitre 8.

Les résultats sur les contours flous obtenus par l'extension proposée de cette approche semblent prometteurs. Néanmoins, on pourrait imaginer apporter quelques améliorations quant à l'extraction de contours flous composés de contours qui partagent une même portion de contour. En effet notre approche de pénalisation par bande empêche un tel comportement. Une autre piste d'investigation à suivre porterait sur l'élaboration d'une méthode formelle de validation des résultats de segmentations floues. En effet ce type de validation est très difficile de par la difficulté à prendre en compte et à expliciter l'incertitude et l'imprécision présentes dans les images pour l'obtention d'une vérité. Une idée pour obtenir plusieurs contours de référence serait de combiner la segmentation de plusieurs experts travaillant indépendamment.

Chapitre 6

Utilisation des contours flous pour la caractérisation de lésions

En mammographie, certains types de lésion sont parfois difficiles à délimiter précisément et de manière unique. Dans le chapitre 5, nous avons introduit des méthodes permettant d'extraire plusieurs contours pour une même structure. Ces approches de segmentation, formalisées grâce à la théorie des ensembles flous, peuvent potentiellement permettre de représenter l'imprécision sur la délimitation d'une lésion pour la propager dans la chaîne de traitement. Il devient alors possible d'attendre la disponibilité d'information pertinente non disponible lors de l'étape d'extraction de contours et ainsi obtenir une meilleure décision. Ce chapitre a pour but de détailler un exemple concret sur l'utilisation et la manipulation de telles méthodes de segmentation. Nous prendrons l'exemple d'un système de caractérisation du type d'une lésion à partir des projections de tomosynthèse servant à la reconstruction du volume.

Dans la méthode que nous proposons, nous différencions les lésions spiculées des lésions circonscrites. Un tel outil est intéressant en pratique dans la mesure où ces deux classes de contours sont généralement fortement corrélées à la malignité et respectivement à la bénignité de la lésion. Remarquons enfin que ce qui est présenté ici n'utilise pas tous les outils théoriques proposés dans les premiers chapitres. Ici, l'idée est principalement d'illustrer, par une application concrète, l'intérêt des méthodes de segmentation floue introduites au chapitre 5.

Dans un premier temps, nous présenterons la chaîne de traitement de manière générale. Cette dernière repose sur l'établissement de deux hypothèses (i.e. la lésion est spiculée ou circonscrite) qui doivent être validées ou invalidées. Nous détaillerons dans les deux sections suivantes les différentes étapes de cette chaîne de traitement, notamment le conditionnement de l'information résultante de l'étape de segmentation, puis l'utilisation d'arbres de décision flous à entrée floue. Enfin nous discuterons des résultats obtenus avec l'approche proposée.

6.1 Schéma global

Notre approche est fondée sur la détection et la segmentation de structures directement à partir des projections. Ce type d'approche a l'avantage de nécessiter un temps de traitement réduit en comparaison avec un traitement direct du volume reconstruit. Néanmoins, le traitement de manière indépendante des projections peut facilement aboutir à une prise de décision prématurée, comme une mauvaise segmentation, mettant en péril la validité de toute la chaîne de traitement. L'utilisation d'une segmentation floue permet de garder l'incertitude émanant de la segmentation dans chacune des projections et le cadre de la logique floue nous fournit les outils pour agréger l'information provenant de ces segmentations permettant ainsi de prendre une décision.

6.1.1 Marqueurs

Tout d'abord, des marqueurs doivent être positionnés sur les localisations des opacités dans les projections. Cette étape peut être effectuée de manière manuelle, semi-manuelle ou complètement automatique. Dans le cadre des résultats présentés dans ce chapitre, les localisations ont été fournies par un expert. En effet, le but de ce chapitre est principalement de différencier les masses spiculées des masses circonscrites, et d'illustrer l'apport du flou pour accomplir cette tâche.

6.1.2 Hypothèses multiples

Dans le but de décrire le type d'une lésion donnée, nous considérons deux hypothèses (c.f. la figure 6.1) : il s'agit soit d'une lésion spiculée, soit d'une lésion circonscrite. Cela résulte en deux types d'informations a priori sur la forme du contour de la lésion qui seront utilisés pour l'étape de segmentation. Le point clé abordé dans ce chapitre sera de valider ou invalider ces hypothèses dans l'étape finale de décision.

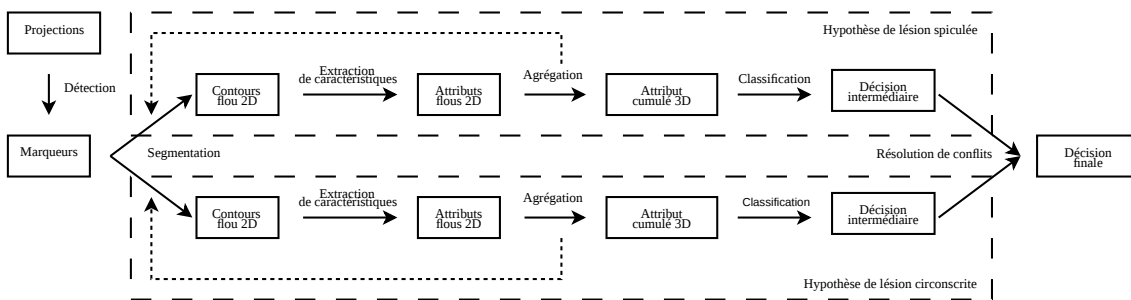


FIG. 6.1 – Schéma global de l'algorithme.

Pour atteindre ce but, chaque structure d'intérêt est segmentée en deux ensembles de contours en utilisant des jeux de paramètres correspondant à chacune des hypothèses. Pour cette étape, nous utilisons la méthode de contour actifs flous proposée par Peters *et al.* (2007) et détaillée au chapitre 5. Ces segmentations sont effectuées pour chacune des projections, ainsi pour chaque hypothèse, une étape d'agrégation est mise en œuvre. Les données résultant de cette dernière étape sont alors utilisées en entrée d'arbres de décision flous comme il sera décrit à la section 6.3.

6.2 Segmentation et conditionnement de l'information

Une partie importante de la chaîne de traitement présentée à la figure 6.1 consiste à segmenter et à extraire l'information des lésions à partir de ces segmentations. Nous allons décrire ici successivement les deux processus.

6.2.1 Extraction de contours flous

L'extraction de contours flous peut se faire de différentes manières comme on a pu le voir au chapitre 5. Dans les expérimentations présentées dans ce chapitre, la méthode des contours actifs flous a été utilisée. Les différentes constantes utilisées dans l'énergie posée sur la fonction de distance signée ont été réglées de manière à traduire les informations a priori sur les lésions circonscrites et les lésions spiculées. Plus concrètement, la contrainte sur la régularité a été relâchée dans le cas d'une hypothèse de lésions spiculées.

6.2.2 Conditionnement de l'information

Dans le but de classifier une lésion, des attributs doivent être extraits à partir des contours. Dans la mesure où le résultat de la segmentation n'est pas net, et qu'une lésion potentielle est à l'origine de plusieurs segmentations (une par projection), il est assez difficile d'utiliser un classifieur classique. Dans ce contexte,

la théorie des ensembles flous peut nous permettre de nous abstraire de ces contraintes : on peut exprimer des valeurs d'attribut sous forme de quantités floues grâce au principe d'extension proposé par Zadeh (1975), puis agréger chacun des attributs flous obtenus pour chaque projection. Il ne reste plus qu'à utiliser un classifieur capable de manipuler des entrées floues comme un arbre de décision flou.

Extraction d'attributs flous

Pour chaque contour net d'un contour flou, des attributs classiques (compacité, homogénéité, orientation des gradients, etc.) sont calculés (Bothorel, 1996). En utilisant le principe d'extension (Zadeh, 1975), des attributs flous peuvent être calculés pour le contour flou. La mise en œuvre du principe d'extension pour cet exemple est illustrée à la figure 6.2.

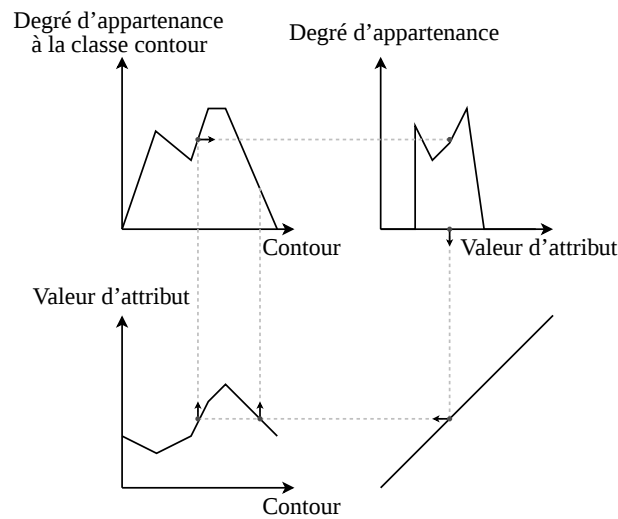


FIG. 6.2 – Principe d'extension : pour chaque valeur d'attribut possible, les contours qui ont cette valeur d'attribut sont considérés et le degré d'appartenance maximal de ces contours est associé à cette valeur d'attribut.

Agrégation des attributs extraits des différentes projections

Puisqu'un marqueur dans l'espace est associé à différents contours flous dans les différentes projections, les attributs flous de ces contours doivent être agrégés. Le résultat de cette agrégation se nomme attribut cumulé flou (Peters, 2007). Cette valeur est représentative de la lésion potentielle vue dans les différentes projections. C'est à partir de ces attributs flous cumulés que le classifieur peut prendre une décision sur la validité de l'une ou l'autre des hypothèses. L'agrégation est faite de manière disjonctive en utilisant une t-conorme (union floue : e.g. $\max$) comme illustré à la figure 6.3.

6.3 Classification à l'aide d'arbres de décision flous

Après l'extraction d'information, des arbres de décision flous sont utilisés pour vérifier si les hypothèses précédentes tiennent. Contrairement aux arbres flous standard, ces arbres ne propagent pas seulement leur entrée dans toutes les branches avec un degré d'appartenance, ils sont aussi capables de manipuler des données d'entrée qui sont floues (Bothorel, 1996). Ces entrées sont composées d'un ensemble de quantités floues (les attributs flous cumulés précédemment calculés) associées aux différents attributs.

Algorithme 6.1 : Propagation d'une particule dans un arbre flou.

faire entrer la particule p dans le nœud racine (associé à un attribut a) avec un degré d'appartenance μ ;

pour chaque fonction de densité f_i^a de ce nœud **faire**

calculer la similitude $S(f_i^a, p_a)$, avec p_a la valeur d'attribut de p ;

propager la particule dans le sous-arbre correspondant à f_i^a avec le degré d'appartenance :

$$\mu \top S(f_i^a, p_a) \quad (6.1)$$

où $\top$ est une t-norme (ici, le min a été choisi) ;

itérer l'algorithme ;

fin

Les résultats présentés dans ce chapitre ont été obtenus avec la mesure de similitude entre deux ensembles flous f et h proposée par Dubois et Prade (1980) et reprise par Bothorel (1996) :

$$S(f, h) = \frac{\int_x \min(f(x), h(x)) dx}{\int_x h(x) dx}$$

Ainsi quand une particule passe à travers l'arbre en entier, on obtient pour chaque feuille le degré avec lequel la particule arrive dans celle-ci (μ_l). Pour prendre une décision finale, les résultats de chaque feuille doivent être agrégés en utilisant une t-conorme $\perp$. Ainsi, on obtient deux degrés μ_A et $\mu_{\bar{A}}$:

$$\begin{aligned} \mu_A &= \perp_{l \in \mathcal{L}} \mu_l \\ \mu_{\bar{A}} &= \perp_{l \in \bar{\mathcal{L}}} \mu_l \end{aligned} \quad (6.2)$$

avec $\mathcal{L}$, l'ensemble des feuilles étiquetées comme A , et $\bar{\mathcal{L}}$ l'ensemble des feuilles étiquetées comme $\bar{A}$.

Dans notre cas, deux arbres utilisant directement les attributs flous précédemment calculés sont construits (un par hypothèse). Pour chaque lésion potentielle, les agrégations des différentes segmentations associées aux différents a priori sont traitées par les arbres correspondant. Ainsi, chaque lésion potentielle est traitée en parallèle par les deux arbres, donnant ainsi accès à quatre degrés de satisfaction pour les propriétés suivantes : la lésions est/n'est pas circonscrite, et la lésion est/n'est pas spiculée.

6.3.3 Construction d'un arbre

Pour construire les arbres évoqués précédemment, on utilise un algorithme récursif sur chaque nœud. Ce dernier repose sur une base d'apprentissage composée de particules étiquetées comme appartenant à la classe A ou à la classe $\bar{A}$. Cet algorithme est initialisé à partir de la racine est fonctionne de la manière décrite à l'algorithme 6.2.

Algorithme 6.2 : Algorithme récursif de construction d'un arbre flou.

tant que le critère de pureté n'est pas vérifié et qu'il reste des attributs non utilisés **faire**

choisir l'attribut le plus discriminant ;

pour chaque future sous-branche **faire**

calculer la fonctions de densité pour le test au nœud en utilisant la base de données ;

faire passer chaque particule de la base de donnée dans le nouveau sous-arbre (en utilisant l'équation 6.1) ;

itérer l'algorithme sur ce dernier ;

fin

fin

La pureté d'une feuille correspond à la proportion d'éléments d'une classe par rapport au nombre total d'éléments qu'elle contient. Ainsi une feuille est pure lorsqu'elle ne contient que des éléments d'une même classe. Cette notion de pureté peut s'exprimer formellement de la manière suivante :

$$\frac{\sum_{p/ \text{classe}(p)=A} \mu(p)}{\sum_p \mu(p)} \geq t \quad \text{ou} \quad \frac{\sum_{p/ \text{classe}(p)=\bar{A}} \mu(p)}{\sum_p \mu(p)} \geq t \quad (6.3)$$

avec t le seuil de pureté, p les éléments de la base d'apprentissage et $\mu(p)$ le degré avec lequel ils arrivent dans la feuille.

Ainsi, on a principalement besoin d'un moyen de sélectionner le meilleur attribut, ainsi qu'une procédure pour calculer les fonctions de densité pour un certain nœud. Pour accomplir la dernière tâche, pour chaque classe, un histogramme flou est calculé : chaque valeur d'attribut est normalisée par l'aire sous sa courbe, et pondérée par le degré d'appartenance avec lequel la particule entre dans le nœud. L'aire de l'attribut flou est aussi normalisée. Cela est équivalent à dire que l'on impose que les histogrammes contiennent un nombre égal de valeurs d'attribut dans le but de pouvoir les comparer. Pour être plus robuste au bruit, et pour être capable de généraliser les données d'apprentissage, cet histogramme est filtré à l'aide d'un filtre médian et d'un filtre moyennneur. Il est ensuite reconstruit par rapport à son maximum et ce dernier est étendu vers la limite du domaine (c.f. figure 6.5). Il en résulte une fonction de densité adéquate pour être utilisée dans un arbre flou. La dernière étape d'extension est importante car elle permet de prendre une décision pour les particules qui ont des valeurs hors de l'intervalle utilisé pour l'apprentissage comme illustré à la figure 6.6.

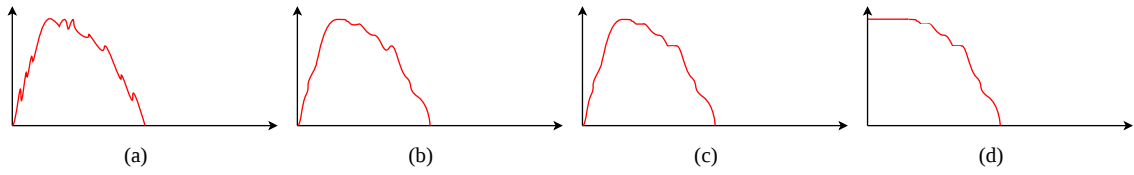


FIG. 6.5 – Construction d'un histogramme flou. Histogramme original (a) filtré par un filtre median et un filtre moyennneur (b), puis reconstruit (c) et finalement étendu (d).

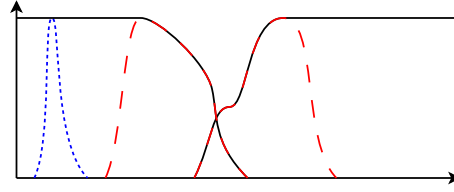


FIG. 6.6 – Exemple de fonctions de test au nœud. Les deux histogrammes non étendus (tirets) ne sont pas adaptés pour classer une particule (pointillés) qui a une valeur non représentée par la population utilisée lors de leur construction. La version étendue de ces histogrammes (trait plein) ne souffre pas de cette limitation.

En utilisant ces fonctions de densité, il est possible d'évaluer à quel point un attribut a est discriminant en fonction des éléments arrivant dans un nœud donné S_k avec un degré d'appartenance μ_{S_k} en utilisant une mesure de gain d'entropie (Lee *et al.*, 1999) définie comme :

$$Gain(S_k, a) = Entropie(S_k) - \sum_{b \in \text{sous-branches}} \frac{|C_{S_k}^{b|a}|}{|C_{S_k}^a|} Entropie(S_k^{b|a}) \quad (6.4)$$

avec :

$$\begin{aligned} C_{S_k}^i &= \sum_{\text{classe}(p)=i \wedge p \in \text{supp}(S_k)} \mu_{S_k}(x) & C_{S_k} &= \sum_i C_{S_k}^i \\ P_i^{S_k} &= \frac{C_{S_k}^i}{C_{S_k}} & Entropie(S_k) &= - \sum_i P_i^{S_k} \log_2 P_i^{S_k} \end{aligned}$$

et pour une branche b , $S_k^{b|a}$ les éléments de la base d'apprentissage évalués par leur degré d'appartenance quand ils arrivent dans la sous-branche b ($\mu_{S_k}(p) \top S(f_b^a, p_a)$).

D'autres approches existent comme des mesures de contraste (Bothorel, 1996; Peters, 2007), mais elles semblent moins adaptées dans le cas d'attributs bruités ou non discriminants. Les étiquettes des feuilles

sont simplement assignées en fonction de la classe la plus représentative (somme pondérée par les degrés d'appartenance) des éléments arrivant dans ces dernières.

6.3.4 Prise de décision

Pour chaque lésion potentielle, nous avons deux particules (une par hypothèse) traitées par deux arbres de décision distincts. Il est donc possible que ces deux arbres ne prennent pas des décisions cohérentes : par exemple, une lésion peut à la fois être classée comme spiculée et circonscrite.

En cas de conflit entre les deux modèles, un degré de confiance est calculé sur les sorties (A et $\bar{A}$) de chaque arbre (Bothorel, 1996) :

$$D_A = \frac{|\mu_A - \mu_{\bar{A}}|}{\mu_A + \mu_{\bar{A}}}$$

Un tel degré représente à quel point un arbre est confiant dans sa décision : si μ_A et $\mu_{\bar{A}}$ sont très proches l'un de l'autre, il est très probable que l'arbre soit incapable de prendre une décision pour la particule considérée. En utilisant ces degrés de confiance, (D_{cir} et D_{sp}), les sorties (degrés d'appartenance μ_{sp} , $\mu_{\bar{sp}}$, μ_{cir} et $\mu_{\bar{cir}}$ aux classes *spiculé*, *non spiculé*, *circonscrit*, *non circonscrit*) des arbres peuvent être pondérés dans l'optique d'obtenir la décision la plus adéquate :

$$\begin{cases} \text{la lésion est circonscrite si } \max(D_{cir}\mu_{cir}, D_{sp}\mu_{\bar{sp}}) > \max(D_{cir}\mu_{\bar{cir}}, D_{sp}\mu_{sp}) \\ \text{sinon la lésion est spiculée} \end{cases}$$

6.4 Résultats

Comme dit précédemment, seules les étapes de segmentation et classification ont été évaluées. Après avoir décrit la base de données utilisée, nous discuterons les résultats obtenus pour chaque partie de la méthode.

6.4.1 Base de données

La base de données utilisée dans cette étude a été acquise avec un appareil expérimental au Massachusetts General Hospital (MGH), Boston, MA, USA. Pour chaque sein, 15 images de projection ont été acquises avec une ouverture angulaire de 30 degrés. Pour garantir une dose totale au patient non supérieure à la dose délivrée lors d'un examen standard de mammographie, la dose par projection est approximativement égale à 10% de la dose délivrée au patient lors d'un examen de mammographie. Toutes les acquisitions ont été prises dans une position médio-latérale oblique (MLO). Aucune vérité terrain n'était disponible pour ces données (pas de rapport de radiologue ou d'histologie). Nous avons donc effectué une revue des cas avec plusieurs experts en imagerie médicale. Cette revue a tout d'abord été effectuée sur les coupes des volumes reconstruits avec la méthode de reconstruction itérative SART (Andersen et Kak, 1984). Nous avons ensuite identifié les régions d'intérêt correspondantes dans les images de projection. Un ensemble de 23 opacités a été identifié à partir de 9 jeux de données de tomosynthèse du sein. La base de données est constituée de 16 lésions circonscrites et de 7 lésions spiculées. A cause de la dose réduite dans les images de projection, certaines des lésions sont difficilement discernables dans ces images. Pour cette même raison, la base de données présente un challenge considérable pour tester un tel système de CAD.

De par la taille de la base de données, une approche de validation croisée (leave-one-out) a été retenue pour évaluer la méthode : chaque élément de la base de données est classifié avec des arbres construits à partir des éléments restants.

6.4.2 Extraction de contours

Des résultats de segmentation pour les deux modèles de contours actifs sont présentés à la figure 6.7 sur une lésion circonscrite et une lésion spiculée. Dans le cas de la lésion circonscrite, même si les contours obtenus en utilisant l'hypothèse spiculée sont moins réguliers, un grand nombre de contours se rapprochent bien du contour réel. Dans le cas d'une lésion spiculée, le contour donné par un expert est plus subjectif.

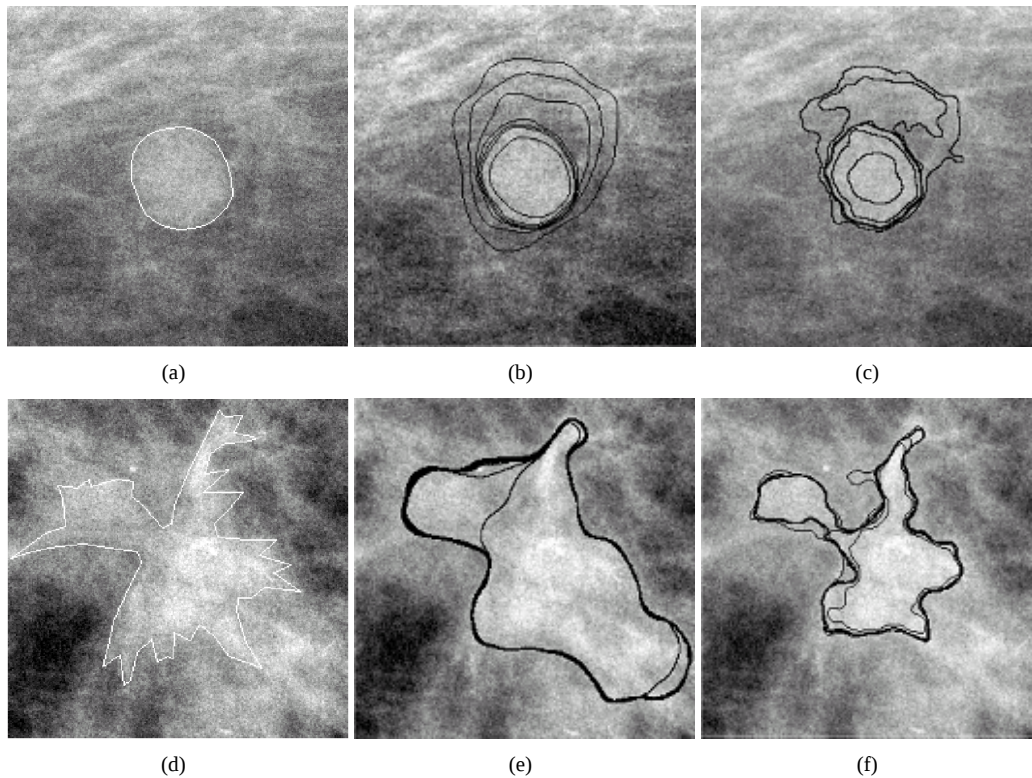


FIG. 6.7 – Résultat de segmentation : une lésion circonscrite segmentée par un expert (a), en utilisant un a priori circonscrit (b) et spiculé (c), et une lésion spiculée segmentée par un expert (d), en utilisant un a priori circonscrit (e) et spiculé (f).

Cela justifie l'extraction de plusieurs contours, et ainsi l'utilisation du flou. Les contours obtenus pour les deux hypothèses sont assez différents : le modèle circonscrit n'est pas capable d'extraire les spicules alors que l'autre modèle y arrive un peu mieux.

6.4.3 Extraction de caractéristiques

Des exemples de calculs d'attributs cumulés sur des régions d'intérêt contenant une lésion spiculée ou une lésion circonscrite sont illustrés à la figure 6.8. Pour la lésion circonscrite, une forte compacité (autour de 0,8) a été mesurée pour la plupart des contours candidats dans les images de projection. Cette information est transférée dans l'attribut flou cumulé où des valeurs proches de 0,8 pour la compacité de la particule sont marquées comme les plus représentatives. Pour la lésion de la figure 6.7(d), des valeurs de compacité plus faibles ont été obtenues. De plus la variance des mesures à travers les différentes images de projections est plus grande que celle obtenue pour la lésion circonscrite. Cela conduit à une plus grande ambiguïté dans la fonction d'attribut cumulé. Des valeurs comprises entre 0,2 et 0,45 semblent être les plus représentatives pour la particule. Les attributs flous cumulés obtenus expriment qu'une forte compacité est caractéristique pour la lésion circonscrite de la figure 6.7(a) alors qu'une faible compacité est plus caractéristique de la lésion spiculée de la figure 6.7(d). Cela correspond bien à l'interprétation intuitive que l'on peut faire de ces deux exemples.

6.4.4 Classification

La figure 6.9(a) présente les taux d'erreur pour l'étape de classification : dans 57% des cas, les deux arbres donnent la bonne décision, dans 39% des cas ils sont en désaccord et dans 4% ils ont tous les deux tort. En utilisant les degrés de confiance, certains des cas portant à conflit peuvent être résolus, donnant

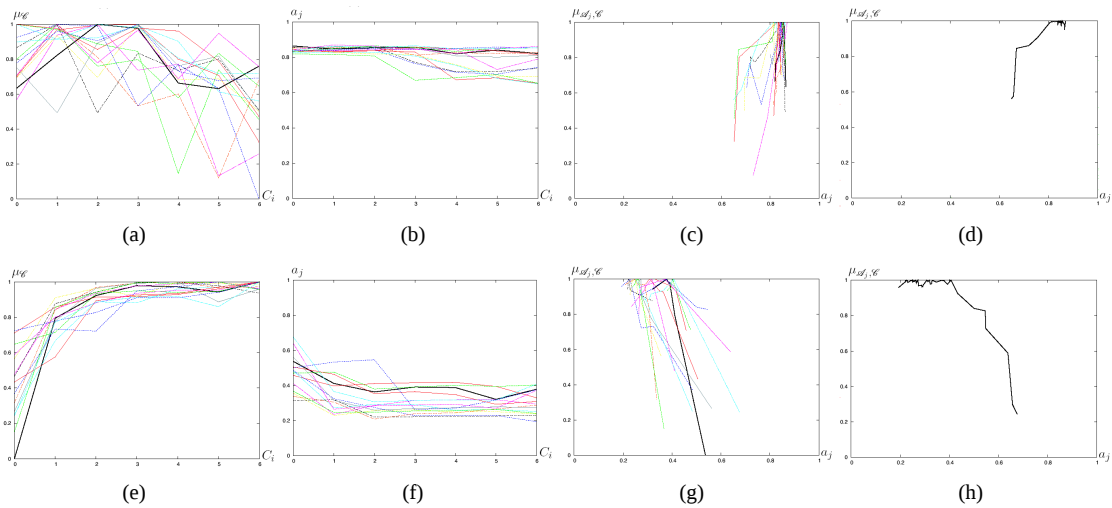


FIG. 6.8 – Agrégation d'attributs flous extraits à partir de régions d'intérêt contenant la lésion circonscrite (resp. spiculée) présentée à la figure 6.7(a) (resp. figure 6.7(d)) : (a) (resp. (e)) correspondant aux fonctions d'appartenance à la classe contour, (b) (resp. (f)) valeurs d'attributs flous pour l'ensemble des contours candidats pour l'attribut compacité, (c) (resp. (g)) la fonction d'appartenance résultante pour l'attribut flou cumulé. Dans (a), (b) et (c) (resp. (e), (f), et (g)), les différentes courbes représentent les valeurs pour différentes images de projection. Le trait plein noir correspond à la projection à 0 degré représentée à la figure 6.7(b) (resp. figure 6.7(e)) et à la figure 6.7(c) (resp. figure 6.7(f)).

ainsi un taux d'erreur de classification total de 9% (c.f. la figure 6.9(b)). Cela montre clairement le gain apporté par l'utilisation de deux hypothèses différentes pour caractériser les lésions : seulement un modèle de segmentation n'aurait pas été suffisant (les modèles circonscrits et spiculés ont tort dans respectivement 13% et 30% des cas).

Tout au long des itérations de la procédure de validation croisée, l'arbre correspondant à l'hypothèse de lésion circonscrite était principalement constitué de trois attributs : la compacité, la moyenne de l'amplitude du gradient sur le contour et l'orientation moyenne sur le contour du gradient relativement au centre de la lésion. En ce qui concerne le second arbre correspondant à l'hypothèse de lésion spiculée, il utilisait principalement la moyenne du gradient dans la région d'intérêt.

La figure 6.10 présente un exemple des différentes décisions possibles. Dans le cas (a), les deux modèles ont donné une mauvaise réponse. Cela est principalement dû à une mauvaise segmentation pour les deux modèles provoquée par un fond non uniforme (lésion en bordure de sein). Dans le cas de la figure 6.10(b), les deux modèles ont donné une réponse contradictoire. L'étape de résolution de conflit a permis de bien classer cette lésion. Le cas de la figure 6.10(c) a aussi donné lieu à un conflit, néanmoins, l'étape de résolution a échoué à donner la bonne classe. Enfin dans le cas de la figure 6.10(d), les deux modèles ont proposé la bonne classe. De manière générale, dans la plupart des cas problématiques impliquant une lésion au bord du sein, certains des contours flous obtenus étaient de mauvaise qualité dans au moins un des deux modèles (généralement celui spiculé car il offre une plus grande liberté de contour).

6.5 Conclusion

Nous avons proposé une chaîne de traitement pour la caractérisation de lésions en tomosynthèse du sein en utilisant la théorie des ensembles flous. Cette approche repose sur deux idées. Tout d'abord, les ensembles flous sont utilisés pour manipuler l'imprécision jusqu'à ce que l'information provenant des différentes projections soit disponible (agrégation des attributs flous). Deuxièmement, deux hypothèses sont faites dans le but d'être validées ou invalidées dans des processus de traitement assez similaires (seul l'information a priori introduit dans l'étape de segmentation change). Cette dernière étape peut induire des conflits

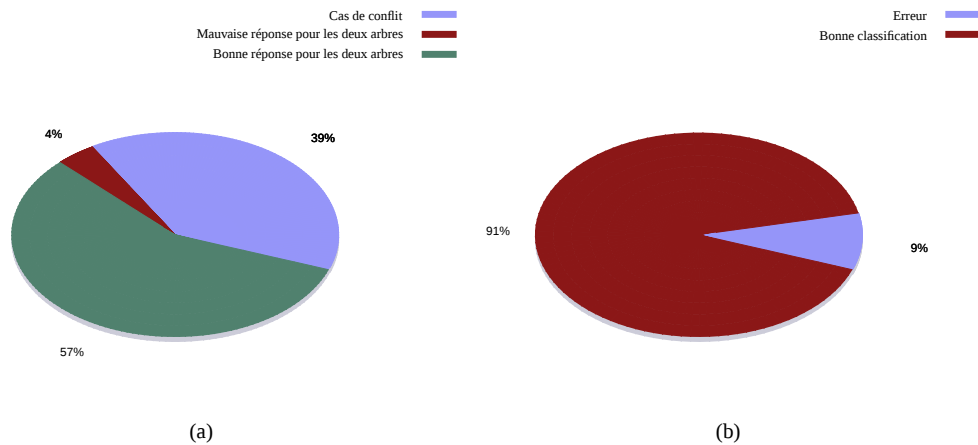


FIG. 6.9 – Taux d'erreurs de classification. (a) Résultats non fusionnés pour les deux classifieurs. (b) Taux d'erreurs final après résolution de conflits.

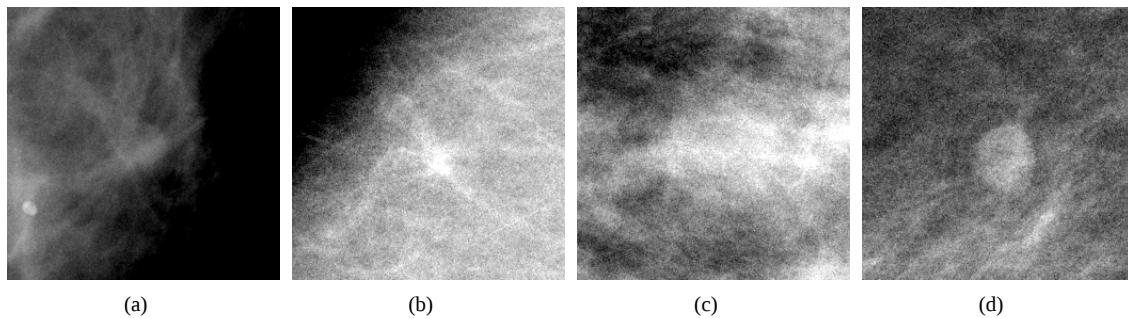


FIG. 6.10 – Exemple de résultats pour le modèle à deux hypothèses. Les images présentées sont extraites des images de projections. (a) Les deux modèles ont donné la mauvaise classe. (b) Les deux modèles étaient en désaccord et l'étape de résolution de conflits a permis de prendre la bonne décision. (c) Même configuration mais avec un échec de la procédure de résolution de conflits. (d) Cas où les deux modèles ont donné la bonne réponse.

dans la décision à prendre dans la mesure où les deux chaînes de traitement associées aux précédentes hypothèses sont indépendantes. Pour cette raison, nous avons introduit une méthode pour résoudre ce genre de problème.

Cette étape de segmentation/classification a été évaluée sur une base de données composée de 23 lésions réelles. Bien que les résultats statistiques obtenus avec une base de données de taille aussi faible ne sont pas complètement fiables, ces résultats tendent à valider l'idée de faire deux hypothèses différentes et d'essayer de les valider/invalides.

Le but de cette application est essentiellement de montrer l'utilité des méthodes de segmentation floues introduites au chapitre 5 dans un système de caractérisation de lésion. Les résultats présentés ici, ainsi que les approches mises en œuvre sont assez différentes de ce qui est proposé dans la chaîne complète de détection automatique qui sera détaillée au chapitre 8. Néanmoins, dans la mesure où les objectifs des deux approches ne sont pas les mêmes, elles ne sont pas comparables mais plutôt complémentaires.

Chapitre 7

Détection de motifs de convergence

Dans les volumes de tomosynthèse numérique du sein, la convergence de structures est souvent une preuve de la présence d'une lésion maligne (lésion spiculée ou distorsion architecturale). Nous avons vu dans les chapitres précédents des outils permettant de détecter les sur-densités qui sont un autre signe de suspicion. Néanmoins ces deux signes radiologique ne sont pas toujours concomitants, et certaines lésions malignes peuvent ne se traduire que par un motif de convergence. Des méthodes efficaces existent dans la littérature pour résoudre ce problème en mammographie conventionnelle (Karssemeijer et te Brake, 1996; te Brake et Karssemeijer, 1999).

Le formalisme de modélisation a contrario (Desolneux *et al.*, 2000) a depuis un certain temps été introduit et a prouvé son efficacité dans des applications de détection en traitement d'images. Son utilisation a déjà été proposée en mammographie, notamment pour la détection de sur-densités en présence de fond texturé (Grosjean, 2007).

Dans ce chapitre nous proposons une méthode inspirée d'une approche originellement proposée par Karssemeijer et te Brake (1996) exprimée à l'aide du formalisme a contrario précédemment évoqué. Cette nouvelle méthode est adaptée à la détection de motifs de convergence dans des données de tomosynthèse du sein. Dans la mesure où les objets contenus dans ces images souffrent d'une distorsion dans l'axe de la profondeur, l'approche proposée se focalisera sur le traitement des coupes du volume. L'information tridimensionnelle sera néanmoins prise en compte dans un second temps comme on le verra au chapitre 8.

Dans un premier temps nous rappellerons quelques généralités sur les lésions qui se traduisent par un motif de convergence dans les images, puis nous verrons les approches existantes pour ce type de problème. Nous discuterons ensuite des différentes définitions mathématiques possibles pour la notion de convergence d'un point vers un autre. Suite à cela, nous proposerons une modélisation a contrario du problème, ainsi qu'un algorithme permettant une mise en œuvre rapide. Nous finirons par une illustration de la méthode sur des images réelles ainsi qu'une discussion sur les moyens de réduire les faux positifs produits par cette approche.

7.1 Motifs de convergence et cancers

Comme on a pu le voir dans le chapitre 1, il existe deux types de lésions qui se traduisent par un motif stellaire. Le premier correspond aux lésions spiculées. Ces dernières se composent généralement d'un noyau dense associé à des spicules qui lui donne une forme en étoile. Le deuxième type de lésions correspond aux distorsions architecturales. Dans ce cas, il n'y a plus de densités au niveau du centre de la lésion. Le signe radiologique permettant de reconnaître ce type de lésion se traduit par une modification locale de la structure du sein laissant paraître un motif en étoile. La figure 7.1 présente un exemple de ces deux types de lésions.

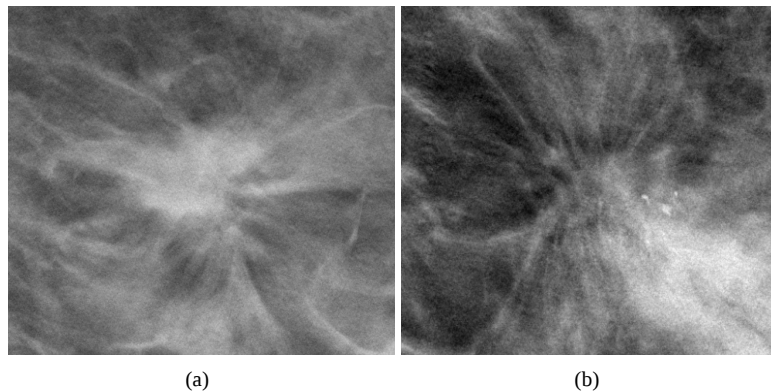


FIG. 7.1 – Exemple d’une lésion spiculée (a) et d’une distorsion architecturale (b) en tomosynthèse du sein.

7.2 Approches courantes

La détection des lésions ayant une structure stellaire peut se faire de deux manières. La première, consiste à détecter les zones de sur-densité, puis à examiner leur périphérie pour vérifier si on est en présence d’une lésion spiculée. La deuxième méthode consiste à détecter quelles sont les zones dans l’image où un motif de convergence apparaît.

7.2.1 Détection de densité et analyse de la périphérie

Ce type d’approche repose sur une première étape de détection des zones suspectes par l’intermédiaire de mesures modélisant la notion de sur-densité. Comme on l’a vu dans le chapitre 1, plusieurs méthodes peuvent être utilisées. Dans les chapitres 2, 3 et 4, nous avons introduit un formalisme permettant la détection de ces zones.

La deuxième étape de ce type d’approche repose sur l’analyse du contour de la lésion. Ainsi, il est nécessaire de la segmenter pour l’analyser. Certaines méthodes de segmentation ont été discutées dans le chapitre 5 et la caractérisation de ce que les contours représentent a été abordée au chapitre 6.

Le problème majeur de ce type d’approche est l’hypothèse de noyau dense vers lequel des spicules convergent. En effet, on ne peut espérer détecter que les lésions spiculées, laissant ainsi de côté les distorsions architecturales.

7.2.2 Détection de convergences

Une deuxième classe d’approche repose sur la détection de motifs stellaires, relaxant ainsi l’hypothèse de présence d’une densité. L’une des méthodes les plus connues a été proposée pour le traitement de mammographies 2D par Karssemeijer et te Brake (1996). Elle repose sur un critère de mesure de convergence défini à partir d’une carte d’orientations des structures présentes dans l’image.

Les orientations sont obtenues par filtrages avec des dérivées secondes de gaussiennes selon trois directions comme illustré à la figure 7.2. En effet en utilisant les travaux de Koenderink et van Doorn (1992), on peut recombinaison les résultats du filtrage d’une mammographie par ces trois noyaux pour obtenir l’orientation des structures linéaires la composant.

En utilisant une telle carte d’orientation, Karssemeijer et te Brake (1996) proposent des mesures de convergence vers un point donné reposant sur l’analyse des orientations dans un anneau centré en ce point comme illustré à la figure 7.3(a). L’idée est de positionner trois cercles de rayons R_1 , R_2 , R_3 différents. Le premier va modéliser la zone focale de convergence autour du centre commun c , et la zone entre les deux derniers va délimiter les pixels pour lesquels on va vérifier s’ils convergent vers la zone focale ou non. Plus concrètement, on va compter le nombre de points q dans cette zone, tels que leur orientation pointe vers le cercle de plus petit rayon. Plus mathématiquement, la définition de convergence peut s’exprimer de la

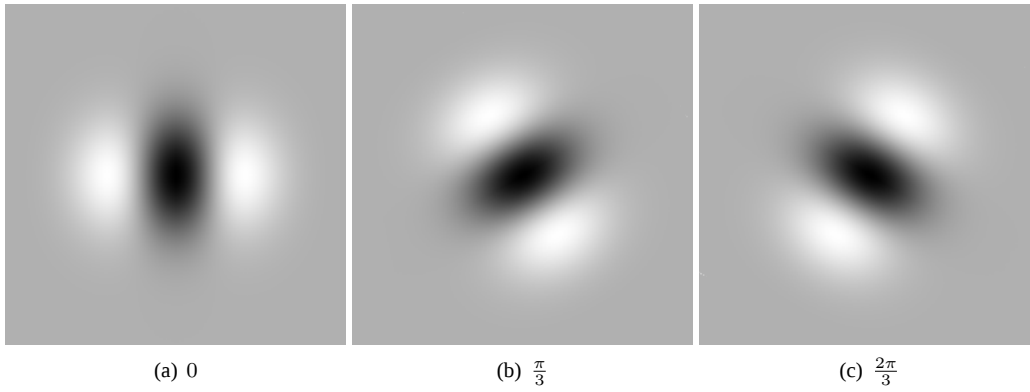


FIG. 7.2 – Dérivée seconde de gaussienne selon trois directions.

manière suivante :

$$\tilde{K}_{c,q} = \begin{cases} 1 & \text{si } \theta ||\vec{c}\vec{q}|| < R_1 \\ 0 & \text{sinon.} \end{cases} \quad (7.1)$$

où θ est l'angle entre le vecteur $\vec{c}\vec{q}$ et la droite portée par q orientée selon la direction encodée en ce point.

Dans la section 7.3, nous détaillerons cette notion de convergence d'un point vers un autre, et nous étudierons les différentes variantes.

Plus concrètement, deux mesures utilisant l'équation 7.1 sont proposées. Ces mesures reposent sur la quantité suivante :

$$n_{c,j} = \sum_{q \in N_{c,j} \cap S} \tilde{K}_{c,q} \quad (7.2)$$

où S représente les pixels ayant une orientation significative, c'est-à-dire dont l'amplitude du filtrage est assez grande, et $N_{c,j}$ représente un sous-domaine de l'anneau défini par R_2 et R_3 comme illustré à la figure 7.3(b).

Cette mesure correspond au nombre d'individus qui convergent vers c et qui proviennent d'une localisation particulière identifiée par l'indice j . Cette subdivision de l'espace permet de différencier les cas où la convergence est plutôt uniformément distribuée des cas où elle ne provient que de quelques directions particulières. L'intérêt d'une telle discrimination est qu'une lésion se trouve souvent dans le premier cas et très rarement dans le second.

La première mesure correspond au comptage de tous les points qui convergent vers le centre considéré indépendamment de leur appartenance aux différents $N_{c,j}$:

$$n_c = \sum_j n_{c,j} \quad (7.3)$$

Puisque le nombre d'éléments dans chaque $N_{c,j} \cap S$ est variable d'un point c à un autre, il est difficile d'interpréter le résultat de cette mesure. Pour faciliter ce point, les auteurs proposent de la normaliser et de la centrer en faisant l'hypothèse que les orientations sont uniformément distribuées :

$$f_{1,c} = \frac{n_c - p|N_i \cap S|}{\sqrt{|N_i \cap S|p(1-p)}} \quad (7.4)$$

où p est la probabilité moyenne qu'un pixel pointe vers le centre c .

La deuxième mesure correspond à l'analyse des provenances des orientations. Sous les mêmes hypothèses d'orientation uniforme, on compte le nombre de fois (n_+) où la valeur $n_{c,j}$ normalisée est supérieure à la valeur médiane pour les différentes valeurs de j , et on définit la mesure suivante :

$$f_{2,c} = \frac{n_+ - J/2}{\sqrt{J/4}} \quad (7.5)$$

avec J le nombre de subdivisions du voisinage autour de c .

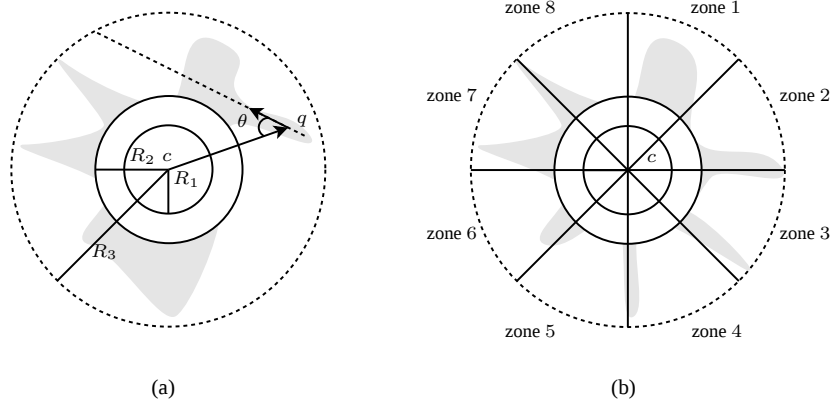


FIG. 7.3 – (a) Mesure de convergence d'un point q vers un point c utilisée par Karssemeijer et te Brake (1996). (b) Découpage des zones autour du point c en fonction de leur orientation par rapport à ce point. Les zones en gris dans les deux figures correspondent aux points q situés dans l'anneau défini par R_2 et R_3 pour lesquels l'amplitude de l'information d'orientation est assez importante pour qu'ils soient considérés.

7.3 Discussion sur la définition de convergence d'un point vers un autre

La méthode proposée dans ce chapitre ainsi que celle avancée par Karssemeijer et te Brake (1996) repose sur différentes manières mathématiques de définir la convergence d'un point vers un autre. Nous allons discuter ici des différentes variantes possibles ainsi que des conditions dans lesquelles elles sont valides.

7.3.1 Définitions possibles

Pour un centre c et un point q donnés, la définition directe que l'on peut donner pour la convergence de q vers c est :

$$K_{c,q}^1 = \begin{cases} 1 & \text{si } (\|\vec{cq}\| > R_1) \wedge \sin(\theta)\|\vec{cq}\| \leq R_1 \\ 0 & \text{sinon.} \end{cases} \quad (7.6)$$

Cette équation revient à comparer la distance minimale entre la droite passant par q , orientée selon la direction associée à ce point, au rayon R_1 du disque représentant la zone de convergence. Ce calcul est illustré à la figure 7.4(a)

Cette formulation, bien qu'intuitive, pose un problème lorsque le point q considéré est très proche du disque de rayon R_1 . En effet, dans le cas extrême où q est sur le cercle, même avec un angle de θ égal à $\frac{\pi}{2}$, le critère de convergence est vérifié comme illustré à la figure 7.4(b). Une solution pour éviter ce problème est de s'interdire de considérer les points qui sont trop proches de ce cercle. Cela peut être fait en considérant un second rayon R_2 plus grand que R_1 et en considérant que les points se situant dans le nouveau disque ne convergent pas :

$$K_{c,q}^2 = \begin{cases} 1 & \text{si } (\|\vec{cq}\| > R_2) \wedge (\sin(\theta)\|\vec{cq}\| \leq R_1) \\ 0 & \text{sinon.} \end{cases} \quad (7.7)$$

Ce cas se rapproche de ce que propose Karssemeijer et te Brake (1996) à l'équation 7.1. La seule différence se situe au niveau du sinus qui est substitué par la fonction identité :

$$K_{c,q}^3 = \begin{cases} 1 & \text{si } (\|\vec{cq}\| > R_2) \wedge (\theta\|\vec{cq}\| \leq R_1) \\ 0 & \text{sinon.} \end{cases} \quad (7.8)$$

En pratique on remarquera que le comportement de ces deux équations est assez similaire pour des angles assez petits. En effet il est tout à fait envisageable d'approcher un sinus de cette manière comme on peut le voir à la figure 7.5.

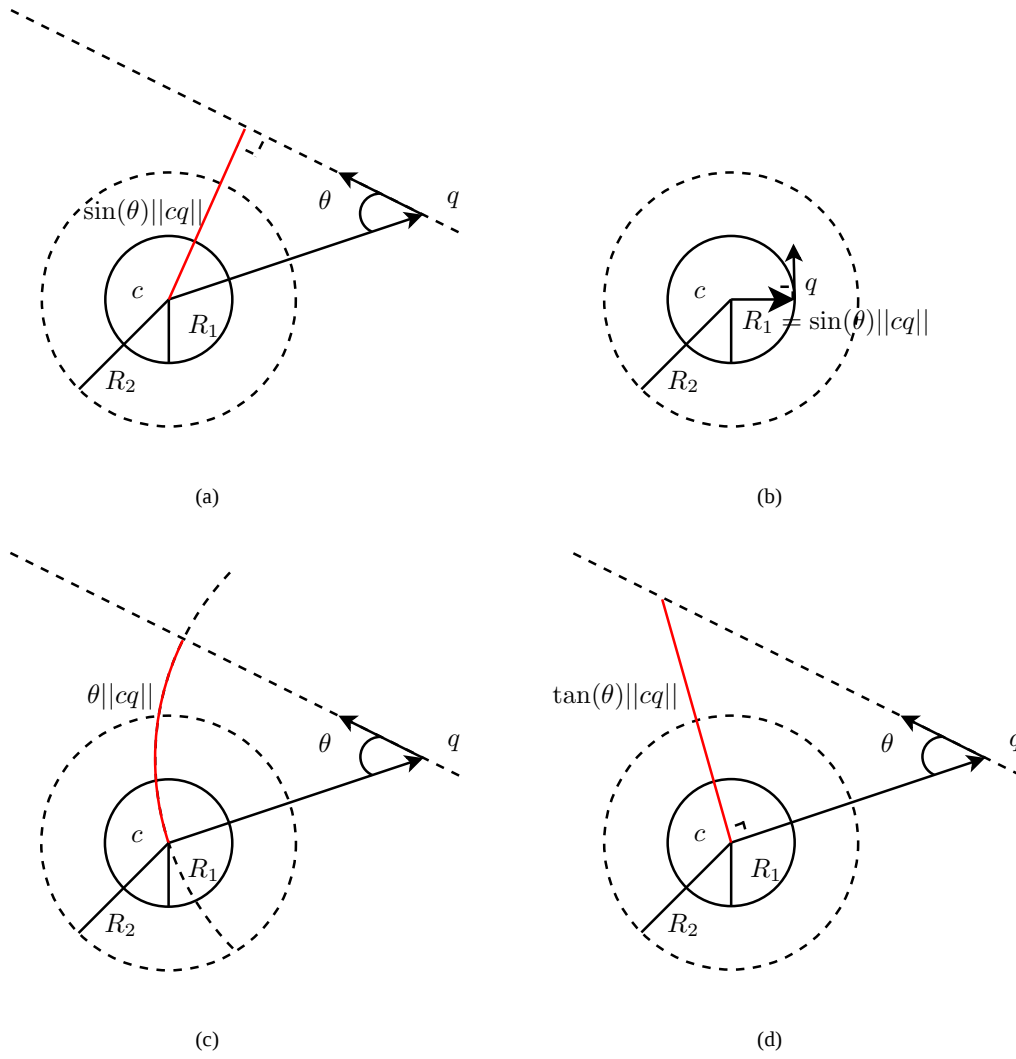


FIG. 7.4 – Comparaison des différentes définitions de convergence d'un point c vers un point q . (a) Illustration de la mesure K^2 . (b) Cas dégénéré possible avec la mesure K^1 . (c) Illustration de la mesure K^3 . (d) Illustration de la mesure K^4 .

Originellement, Karssemeijer et te Brake (1996) n'expliquent pas cette approximation. Dans le brevet traitant de cette même approche (Karssemeijer et Te Brake, 1997), il est expliqué que la fonction identité est une approximation pour l'opérateur tangente. Ainsi, l'équation précédente est équivalente à :

$$K_{c,q}^4 = \begin{cases} 1 & \text{si } (||\vec{cq}|| > R_2) \wedge (\tan(\theta)||\vec{cq}|| \leq R_1) \\ 0 & \text{sinon.} \end{cases} \quad (7.9)$$

Néanmoins, pour un rapport $\frac{R_2}{R_1} = 2$ comme proposé dans Karssemeijer et Te Brake (1997) et pour les angles considérés, les trois opérateurs sont équivalents (c.f. figure 7.5).

Remarque 7.3.1. On peut remarquer que si, à partir de cette équation, on relâche la contrainte sur R_2 , on obtient un système qui ne souffre plus du problème de convergence non intuitive lorsqu'un point q est trop proche du disque de rayon R_1 .

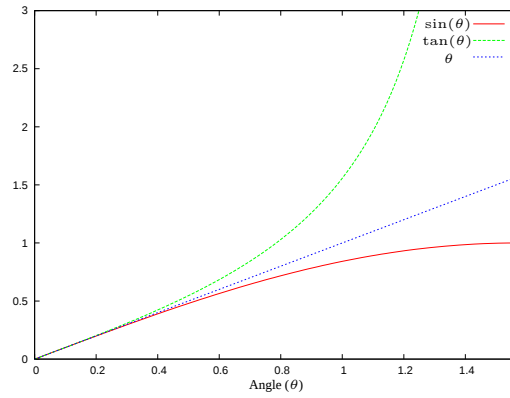


FIG. 7.5 – Approximation d'un sinus par la fonction identité et la fonction tangente pour de petits angles.

7.3.2 Comparaison des définitions

Pour comprendre ces approximations, on peut considérer le cas critique où le point q se situe à la limite définie par R_2 (ou R_1 le cas échéant) et se reporter à la figure 7.6. Pour s'abstraire de l'échelle du rayon de convergence utilisé en pratique, on peut, sans perte de généralité, introduire un coefficient β tel que :

$$\beta = \frac{R_2}{R_1} \quad (7.10)$$

Dans la figure 7.6, on trace la valeur de β lorsque q est à la distance critique R_2 de c pour les trois modèles exposés jusqu'à présent (sinus, identité et tangente) en fonction de l'angle maximal telle que la propriété de convergence correspondante soit vérifiée. Deux points de fonctionnement sont mis en évidence pour comparer les trois approches : un rapport $\beta = 2$ et un rapport $\beta = 1$ qui correspondent respectivement à la configuration proposée par Karssemeijer et Te Brake (1997) et à la non utilisation d'une zone de restriction pour la localisation des points q (lorsque $\beta = 1$, $R_1 = R_2$). La première chose à remarquer est que dans le cas où $\beta = 2$, les angles θ maximaux vérifiant les différents critères de convergence sont assez proches : ils varient de 0,46 à $\frac{\pi}{6}$. Cela laisse penser que les trois modèles sont dans ce cas assez similaires. Néanmoins lorsque l'on considère des valeurs de β plus petites, les différents modèles ne sont plus équivalents, ainsi lorsque l'on considère un $\beta = 1$, l'intervalle d'angle maximal pour les différents modèles va de $\frac{\pi}{4}$ à $\frac{\pi}{2}$. Dans tous les cas, c'est-à-dire pour toutes les valeurs de β , on peut remarquer que la mesure reposant sur l'opérateur tangente est plus stricte que les autres. Cela justifie la remarque 7.3.1.

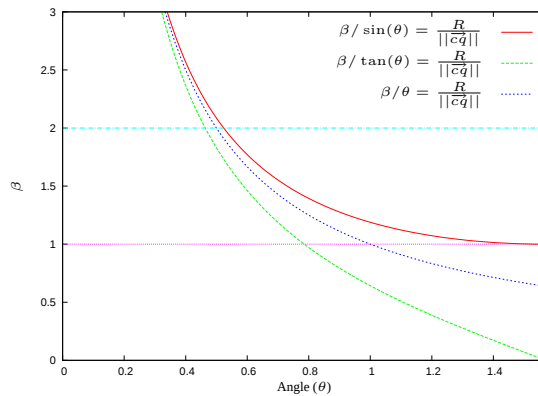


FIG. 7.6 – Rapport (β) entre $\|\vec{c}q\|$ et R_1 tel que $\tan(\theta) = \frac{R}{\|\vec{c}q\|}$, $\theta = \frac{R}{\|\vec{c}q\|}$ et $\sin(\theta) = \frac{R}{\|\vec{c}q\|}$.

Graphiquement, les équations précédentes ont un sens assez différent (c.f. figure 7.4). Dans le cas d'un

sinus on compare le rayon R_1 à la distance D_1 la plus courte entre le point c et la droite portée par l'orientation au point q . Dans le cas où la tangente est utilisée, on compare R_1 à la distance D_2 entre le point c et l'intersection entre la droite portée par l'orientation sous q et la droite perpendiculaire à cq passant par c . Dans le cas où on utilise la fonction identité, la distance D_3 mesurée n'est plus celle d'une droite, mais d'un arc de cercle. Ce dernier correspond à la portion du cercle centré en q de rayon $\|cq\|$ et allant de c à l'intersection avec la droite portée par l'orientation de q .

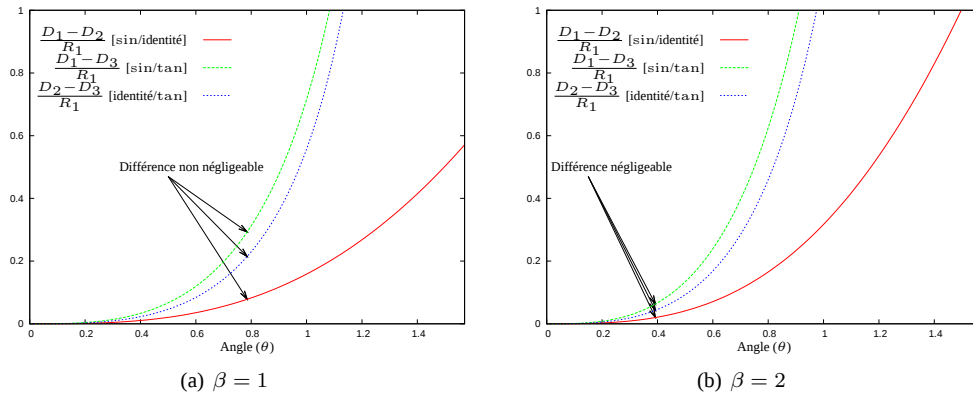


FIG. 7.7 – Différences entre les distances D_1 , D_2 et D_3 normalisées par R_1 en fonction de θ pour $\beta = 1$ et $\beta = 2$.

La comparaison entre ces intervalles d'angles pour les différentes valeurs de β évoquées lors de l'analyse de la figure 7.6 ne donne qu'une idée de l'interchangeabilité ou non des différents définitions de convergence dans la mesure où l'impact d'une variation d'un angle sur les mesures précédentes dépend de la distance $\|cq\|$ et donc de β . Pour être plus rigoureux, on peut comparer les différentes distances que représentent les valeurs $\tan(\theta)\|cq\|$, $\sin(\theta)\|cq\|$ et $\theta\|cq\|$. Encore une fois pour des raisons d'abstraction de l'échelle à laquelle on travaille en pratique, nous comparons ces mesures normalisées par R_1 (c.f. figure 7.7). Le choix de R_1 est judicieux dans la mesure où les différents critères de convergence comparent les précédentes valeurs à ce même rayon R_1 . La figure 7.7(a) présente les différences normalisées entre les trois mesures dans le cas où $\beta = 1$. La zone marquée par les flèches correspond à la comparaison des trois mesures pour l'angle maximal tel que l'on ait convergence pour le modèle le plus strict à base de tangente. En se plaçant sur cet angle, on remarque que les différences sont de l'ordre 8%, 29% et 21% respectivement pour les différences entre les modèles sinus/identité, sinus/tangente et identité/tangente. La figure 7.7(b) présente les différences normalisées entre les trois mesures dans le cas où $\beta = 2$. Comme précédemment, les flèches pointent différences pour l'angle maximal déduit de l'étude de la figure 7.6. Dans cette configuration, les différences entre les mesures sont de l'ordre de 3%, 10% et 7% respectivement pour les modèles sinus/identité, sinus/tangente et identité tangente, ce qui est plus raisonnable que dans le cas précédent.

On remarquera tout de même qu'ici, on compare les différentes méthodes pour l'angle maximal obtenu pour la méthode la plus stricte. Les différences évoquées pour les différentes méthodes sont donc valides si on considère des angles qui ont presque cette valeur. Dans le cas de $\beta = 2$, on a vu à la figure 7.6 que c'était le cas. Néanmoins, dans le cas où β est plus petit, cela devient de moins en moins vrai. Ainsi même si la variation entre la méthode utilisant le sinus et celle utilisant la fonction identité est d'environ 8% pour un tel β , le fait que la mesure de convergence utilisant le sinus soit beaucoup moins stricte implique la non équivalence des deux mesures pour de telles configurations.

On remarquera que globalement, la différence entre les mesures sinus/identité est toujours plus petite que celle entre les mesures tangente/identité. Cela revient à dire que la mesure proposée dans les travaux de Karssemeijer et te Brake (1996) est plus proche de celle à base de sinus que d'une mesure à base de tangente.

On se référera enfin à la figure 7.8 pour voir de manière plus visuelle les différences entre les trois méthodes pour les cas où $\beta = 1$ et $\beta = 2$. Dans le premier cas, le modèle en sinus est beaucoup plus permissif que les deux autres modèles alors que dans le deuxième cas, les trois modèles sont à peu près

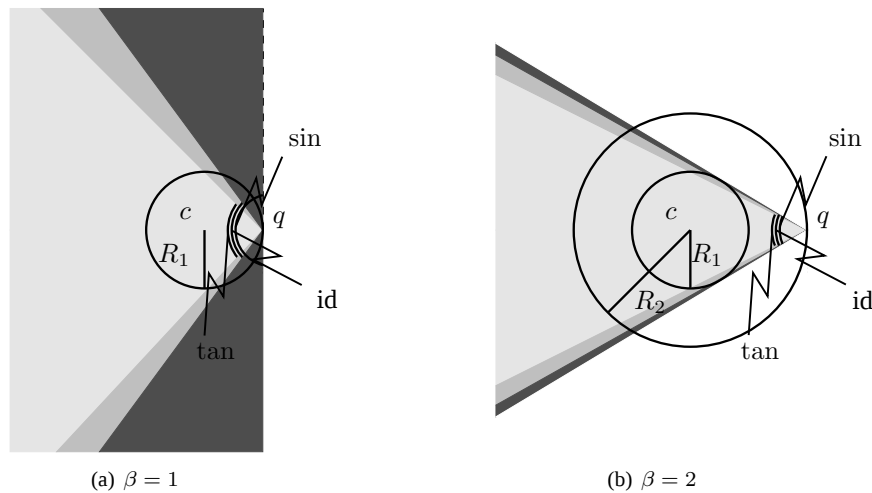


FIG. 7.8 – Comparaison visuelle des ouvertures angulaires maximales pour les trois modèles construits à partir des fonctions sinus, tangente et identité pour $\beta = 1$ et $\beta = 2$.

équivalents.

Dans la suite de ce chapitre on utilisera le modèle en tangente sans zone morte (anneau pour lequel on s'interdit de regarder la convergence des points qu'il contient). En effet, comme on vient de le voir, il se comporte de manière intuitive quelque soit la distance entre les deux points considérés.

7.4 Modélisation a contrario pour la détection de zones de convergence

Récemment, une nouvelle manière de poser le problème de détection d'une structure a été proposée par Desolneux *et al.* (2000). Originellement, la méthode a été proposée pour la détection d'alignements de points dans une image (modèle naïf). L'idée dans ce qui est proposé peut se résumer de la manière suivante : on définit ce qu'est le contenu normal d'une image, ainsi qu'une mesure locale correspondant au motif recherché, et on regarde les occurrences peu probables dans le modèle naïf. En pratique si la modélisation est correctement faite et que le critère est judicieusement choisi, cela permet de détecter les structures recherchées.

Plus concrètement, Desolneux *et al.* (2000) ont proposé des définitions formelles pour cette notion de réalisation peu probable d'un événement grâce à l'introduction du nombre de fausses alarmes dans une image. Grâce à cela, une définition d'événements ϵ -significatifs a été proposée.

Dans cette section nous utilisons le même formalisme en l'adaptant à la détection de zones de convergence. Dans un souci de clarté, et pour éviter toute redondance entre notre approche et les travaux originaux de Desolneux *et al.* (2000) nous réintroduisons le formalisme de raisonnement a contrario en l'illustrant sur notre problématique en mettant l'accent dès qu'une contribution est proposée.

Puisque les objets contenus dans les volumes de tomosynthèse du sein souffrent d'une distorsion verticale qui est due à l'angulation limitée de la géométrie du système d'acquisition, nous proposons de travailler de manière indépendante sur les tranches du volume reconstruit. L'information 3D donnée par la localisation spatiale des détections sera néanmoins utilisée pour réduire le nombre de faux positifs. Dans un premier temps, nous construisons un modèle statistique du motif de convergence dans les tranches reconstruites d'un sein ne contenant aucune lésion (modèle naïf). Ce modèle repose sur la définition de plusieurs variables aléatoires pour chaque pixel d'une coupe donnée dans le but de récupérer les zones de convergence qui sont peu probables. Ces variables correspondent à des orientations déterminées pour chaque pixel situé dans un voisinage en forme d'anneau centré sur le pixel courant. Dans la mesure où les motifs de convergence ne sont normalement pas contenus dans le modèle naïf, l'étape suivante consiste à trouver un seuil permettant de déterminer quand les valeurs de ces variables aléatoires deviennent très peu probables. Le cadre de la

modélisation a contrario permet de résoudre ce point en calculant des seuils à partir du nombre de fausses détections que l'on est prêt à accepter dans le modèle naïf.

7.4.1 Variables aléatoires et modèle naïf

Tout d'abord, un modèle naïf est proposé pour exprimer ce à quoi doit ressembler le contenu classique d'un volume de tomosynthèse du sein. Ainsi nous considérons chaque coupe comme étant constituée d'un champ d'angles uniformément et identiquement distribués. Ces angles correspondent à l'orientation de structures dans l'image (par exemple, les directions orthogonales au gradient). L'idée principale est d'identifier les zones de convergence qui ont une faible probabilité d'apparition dans le modèle naïf. En pratique ces zones devraient correspondre à des distorsions architecturales, à des lésions spiculées ou à des croisements de fibres.

Pour détecter ces zones, deux variables aléatoires sont construites. Elles reposent toutes deux sur la construction de deux cercles emboîtés (de rayon r et αr avec $\alpha \in]0, 1[$ une constante). Le cercle le plus petit définit le centre de convergence et l'anneau délimité par les deux cercles, la zone où l'on cherche les structures qui convergent vers le centre. Cela est illustré à la figure 7.9(a). La première variable aléatoire est notée $K_{c,q,r}$ et est définie comme suit :

$$K_{c,q,r} = \begin{cases} 1 & \text{si } (\alpha r < \|\vec{c\bar{q}}\| < r) \wedge (\tan(\theta)\|\vec{c\bar{q}}\| \leq \alpha r) \\ 0 & \text{sinon.} \end{cases} \quad (7.11)$$

où θ correspond à l'angle entre $\vec{c\bar{q}}$ et l'orientation au point q . Cette mesure est dérivée de la relaxation évoquée à la remarque 7.3.1 du modèle le plus strict K^4 présenté à l'équation 7.9.

Dans le cas d'un champ d'angles composés de variables uniformément distribuées, on a pour un point q situé entre les deux cercles :

$$P[K_{c,q,r} = 1] = \frac{2}{\pi} \tan^{-1} \left(\frac{\alpha r}{\|\vec{c\bar{q}}\|} \right)$$

On peut remarquer que cette probabilité dépend de la distance entre le point q et le centre c ainsi que de la taille et du rapport entre les deux cercles. En fait, quand les deux points se rapprochent l'un de l'autre, la probabilité de cette variable d'être égale à 1 croît.

On peut aussi remarquer que $K_{c,q,r}$ n'est pas toujours égal à un lorsque la ligne droite passant par le point q , dont l'orientation est donnée par l'angle en ce point, intersecte le cercle de rayon αr qui est centré en c comme on l'a vu à section 7.3. Cela aurait été le cas si la tangente avait été remplacée par un sinus dans l'équation 7.11. Néanmoins, un tel choix aurait été problématique dans la mesure où certains points trop près du cercle de rayon αr aurait pu être considérés comme pointant vers le centre de convergence alors qu'ils intersectent l'extrémité du disque comme illustré à la figure 7.9(c). Dans la mesure où on ne s'interdit pas de se rapprocher très près de la zone de convergence, cette formulation est mieux adaptée (c.f. section 7.3).

La seconde variable $Z_{c,r}$ est construite en utilisant la précédente variable pour tous les points q de Ω . Puisque $K_{c,q,r} = 0$ pour les q hors de la zone entre les deux cercles centrés sur c et de rayons respectifs αr et r , on ne compte que la valeur de $K_{c,q,r}$ pour les points q qui se situent dans ce sous-ensemble du domaine Ω comme illustré à la figure 7.9(b) :

$$Z_{c,r} = \sum_{q \in \Omega} K_{c,q,r} \quad (7.12)$$

La variable définie à l'équation 7.12 est un indicateur de convergence parce qu'elle représente le nombre de pixels q qui pointent vers un centre c donné.

Dans le but de calculer la probabilité pour cette variable d'être égale à un entier donné, nous pouvons utiliser la fonction génératrice de $K_{c,q,r}$. Plus généralement, pour une variable aléatoire discrète qui prend ses valeurs dans $\mathbb{Z}^+$, la fonction génératrice permet d'exprimer les probabilités d'être égale à n'importe quel entier en utilisant une représentation polynomiale :

$$G_X(x) = \sum_{n=0}^{\infty} P[X = n] x^n$$

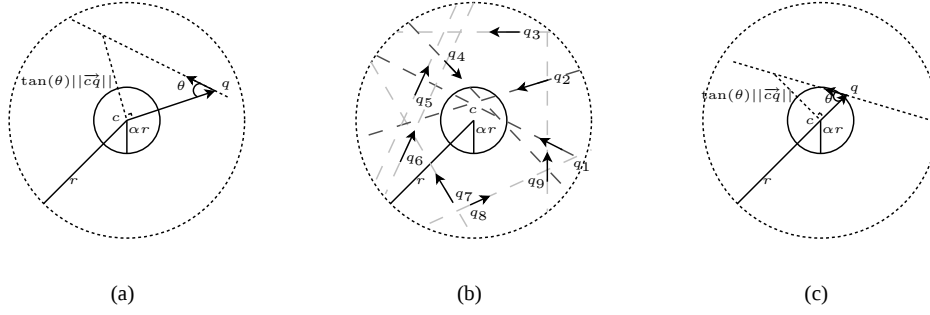


FIG. 7.9 – (a) Cercles emboîtés utilisés pour le calcul de $K_{c,q,r}$ qui est ici égal à 0. (b) Calcul de $Z_{c,r}$: les flèches sur les lignes en traits d'union (passant par q_1 , q_2 et q_4) convergent alors que les autres ne convergent pas. (c) Exemple de convergence non intuitive évitée.

Dans le cas de $K_{c,q,r}$, on a pour n'importe quel triplé (c, q, r) :

$$G_{K_{c,q,r}}(x) = P[K_{c,q,r} = 0] + P[K_{c,q,r} = 1]x$$

Maintenant, puisque nous avons fait l'hypothèse que les orientations à différents points q étaient indépendantes, il suit que les différentes variables $K_{c,q,r}$ sont elles aussi indépendantes. Cela nous permet de calculer la fonction génératrice de $Z_{c,r}$ en multipliant les $G_{K_{c,q,r}}$ associés aux différents points q qui se situent dans l'anneau centré en c défini par les cercles de rayons αr et r :

$$G_{Z_{c,r}}(x) = \prod_{q \in \Omega / \alpha r < ||\vec{c}q|| < r} G_{K_{c,q,r}}(x)$$

Ce qui peut être réécrit comme :

$$G_{Z_{c,r}}(x) = \prod_{q \in \Omega / \alpha r < ||\vec{c}q|| < r} (P[K_{c,q,r} = 0] + P[K_{c,q,r} = 1]x) = \sum_{k \in \mathbb{Z}} p_k x^k \quad (7.13)$$

avec $P[Z_{c,r} = k] = p_k$.

Le calcul de ces probabilités est requis pour permettre de définir à partir de quand une valeur de $Z_{c,r}$ devient improbable. L'étape suivante est la définition de cette notion d'*improbable*. En fait, le cadre général de la modélisation a contrario fournit une définition en utilisant le concept d'événement ϵ -significatifs.

7.4.2 Convergences ϵ -significatives

La modélisation a contrario repose essentiellement sur la détection d'événements qui sont peu probables dans le modèle naïf. Pour faire la distinction entre événement probable et improbable on peut utiliser la définition d'un événement ϵ -significatif proposée par Desolneux *et al.* (2000).

Définition 7.4.1. Un événement est ϵ -significatif si la moyenne de son nombre d'occurrences dans l'image est inférieur à ϵ .

Cette définition introduit un nouveau paramètre ϵ qui correspond à notre acceptation d'événements peu probables. Si on est capable de faire dépendre de ce paramètre la détection des zones de convergence, et si notre modèle naïf tient pour des données réelles, il est possible d'ajuster de manière intuitive les performances de notre détection. En effet, ϵ correspond au nombre de faux positifs que l'on est prêt à accepter dans le modèle naïf.

Revenons maintenant sur les variables aléatoires que nous avons introduites précédemment et plus spécifiquement sur $Z_{c,r}$, pour définir un événement qui correspond à la détection de zones de convergence. Pour se faire, on peut décider que c est le centre d'une telle zone définie par un rayon r si $Z_{c,r} \geq \lambda$ avec λ un seuil. Évidemment, la probabilité d'occurrence d'un tel événement pour un couple (c, r) dépendra de la valeur de ce seuil.

Théorème 7.4.1. L'événement $Z_{c,r} \geq \lambda_r$ est ϵ -significatif avec λ_r défini comme suit :

$$\lambda_r = \min \left\{ \lambda \in \mathbb{N} / P[Z_{c,r} \geq \lambda] \leq \frac{\epsilon}{M} \right\} \quad (7.14)$$

où M est le nombre total de couples (c, r) à considérer dans l'image (les valeurs de r sont comprises entre $R_{\min}$ et $R_{\max}$).

Démonstration. La moyenne du nombre de détections de zones de convergence est :

$$E_Z = \sum_{r \in [[R_{\min}, R_{\max}]]} \sum_{c \in \Omega} P[Z_{c,r} \geq \lambda_r]$$

Maintenant, en utilisant l'équation 7.14, on a :

$$\sum_{r \in [[R_{\min}, R_{\max}]]} \sum_{c \in \Omega} P[Z_{c,r} \geq \lambda_r] \leq \sum_{r \in [[R_{\min}, R_{\max}]]} \sum_{c \in \Omega} \frac{\epsilon}{M}$$

De plus, puisque M correspond au nombre de couples (c, r) considérés dans une image, on obtient :

$$\sum_{r \in [[R_{\min}, R_{\max}]]} \sum_{c \in \Omega} \frac{\epsilon}{M} = M \frac{\epsilon}{M} = \epsilon$$

Ainsi, on a $E_Z \leq \epsilon$. Pour cette raison, le théorème 7.4.1 est vérifié. $\square$

Remarquons finalement que le choix de λ_r dépend du rayon r . Cela s'explique par le fait que quand le rayon devient plus grand, le nombre de points q à considérer augmente, résultant potentiellement dans une plus grande valeur de Z . De plus, la position de c dans l'image n'impacte pas la valeur du seuil.

7.4.3 Détection de convergence ϵ -significatives

Dans le but de détecter les zones de convergence, les variables $Z_{c,r}$ sont calculées pour tous les $c \in \Omega$ et $r \in [[R_{\min}, R_{\max}]]$. La décision est alors faite en utilisant les seuils proposés au théorème 7.4.1. Dans le but de calculer ces seuils, pour chaque rayon r , les probabilités $P[K_{c,q,r} = 1]$ pour les points q situés dans l'anneau sont calculés permettant d'utiliser l'équation 7.13 qui est requise par le théorème 7.14. Puisque les probabilités associées aux $K_{c,q,r}$ ne dépendent pas de la localisation de c , mais plutôt de la distance $\|\vec{c}\vec{q}\|$, cela peut être fait une unique fois pour un centre arbitraire.

La carte d'angles à utiliser pour cette étape est déduite des directions orthogonales aux gradients de l'image d'entrée. Remarquons que ces gradients ne doivent pas être calculés à partir de filtrage linéaire (dérivées de Gaussienne, Sobel, etc.) parce que l'hypothèse d'indépendance faite pour le calcul des seuils serait invalidée. Cela peut paraître contre-intuitif, mais si on utilise de telles approches, les statistiques sur $Z_{c,r}$ seraient modifiées invalidant ainsi l'approche utilisée pour calculer les seuils à moins de considérer une carte d'orientations de variables aléatoires corrélées. Dans ce cas, la corrélation dépendrait de l'étape de convolution utilisé dans l'étape de calculs de la carte de gradients. Dans notre mise en œuvre, nous utilisons des blocs de quatre (2×2) pixels qui ne se recouvrent pas (c.f. les figures 7.10(a) et 7.10(b)). Pour chaque bloc, les gradients sont calculés par différences des moyennes des pixels le long des directions x et y comme illustré aux figures 7.10(c) et 7.10(d). Ainsi, on obtient un vecteur gradient par bloc de quatre pixels, ce qui se traduit par une carte de gradients sous-échantillonnée par rapport à l'image originale. Ce sous-échantillonnage permet d'avoir une estimation plus robuste de la carte de gradients. En utilisant des blocs plus grands, le sous-échantillonnage deviendrait plus important résultant en une carte de gradient plus régulière au détriment de la résolution. Bien que cela soit dramatique pour de petites lésions, le principe prend tout son sens pour des lésions plus grandes. Dans le but de combiner les avantages des différentes tailles de bloc, on peut utiliser une approche multi-résolution en doublant le côté des blocs à chaque étape et en réduisant la valeur de $R_{\max}$. Cette approche permet aussi d'améliorer les temps de calculs.

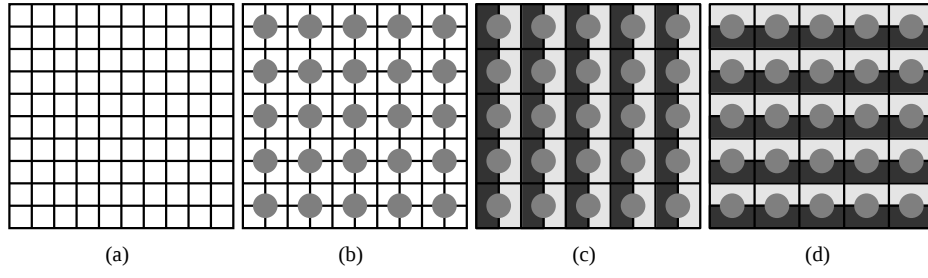


FIG. 7.10 – Extraction d'une carte d'orientations à partir d'une carte de gradients non corrélés. (a) Image originale où les pixels sont représentés par des carrés blancs. (b) Blocs de quatre pixels qui ne se recouvrent pas et qui sont utilisés lors du calcul de la carte de gradients. Les disques gris représentent les localisations des vecteurs de gradient. (c) Calcul du gradient dans la direction des x . (d) Calcul du gradient dans la direction des y .

7.5 Mise en œuvre rapide

Une mise en œuvre rapide de l'équation 7.12 pour un intervalle de rayons donnés peut être inefficace et en pratique inutilisable quand on veut détecter de grandes structures. Dans cette section, nous proposons une décomposition du problème permettant de réduire la complexité du calcul des $Z_{c,r}$. Dans un second temps nous étudions la complexité de la nouvelle approche.

7.5.1 Décomposition du problème

Dans l'optique d'accélérer les calculs des différents $Z_{c,r}$, on peut remarquer que pour un centre c donné, si on considère deux rayons successifs, $r - 1$ et r , certains points q dans Ω ont la même propriété de convergence ou non-convergence. Inversement, il existe un ensemble de points qui convergent pour r , alors qu'ils ne convergent pas pour $r - 1$. Considérons le cardinal, noté $|\cdot|$, de cet ensemble :

$$\delta_{c,r} = |\{q \in \Omega / (K_{c,q,r} = 1) \wedge (K_{c,q,r-1} = 0)\}| \quad (7.15)$$

De manière similaire, on peut définir un ensemble de points qui convergent pour $r - 1$ alors qu'ils ne convergent pas pour r . Son cardinal est :

$$\psi_{c,r} = |\{q \in \Omega / (K_{c,q,r} = 0) \wedge (K_{c,q,r-1} = 1)\}| \quad (7.16)$$

La figure 7.11 illustre les deux quantités $\delta_{c,r}$ et $\psi_{c,r}$ sur un exemple graphique.

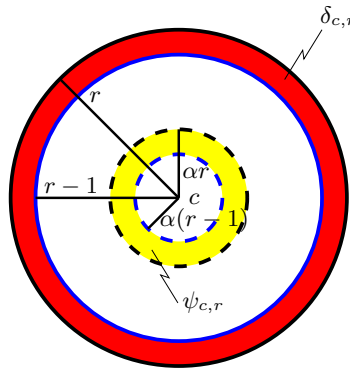


FIG. 7.11 – Illustration des quantités $\delta_{c,r}$ et $\psi_{c,r}$ pour une configuration r et $r - 1$ donnée.

Pour des raisons de clarté, introduisons la notation suivante : $\forall b \in \{0, 1\} \ C_r^b = \{q \in \Omega / K_{c,q,r} = b\}$, qui permet de réécrire les équations 7.15 et 7.16 comme suit :

$$\delta_{c,r} = |C_r^1 \cap C_{r-1}^0| \quad (7.17)$$

et

$$\psi_{c,r} = |C_r^0 \cap C_{r-1}^1| \quad (7.18)$$

Ces deux quantités permettent de décomposer $Z_{c,r}$ comme suit :

Théorème 7.5.1. $Z_{c,r} = \sum_{k=0}^r \delta_{c,k} - \sum_{k=0}^r \psi_{c,k}$

Démonstration. Prouvons le théorème par récurrence :

Initialisation : De manière évidente, on a $Z_{c,0} = 0 - 0 = \delta_{c,0} - \psi_{c,0}$ puisque pour tout $(c, q) \in \Omega^2$ et $r \leq 0$, on a $K_{c,q,r} = 0$.

récurrence : Supposons que le théorème 7.5.1 est vérifié pour un r donné. On peut écrire :

$$\begin{aligned} Z_{c,r+1} &= \sum_{q \in \Omega} K_{c,q,r+1} \\ &= |C_{r+1}^1| \\ &= |C_{r+1}^1 \cap C_r^1| + |C_{r+1}^1 \cap C_r^0| \\ &= |C_{r+1}^1 \cap C_r^1| + \delta_{c,r+1} \text{ (c.f. Equation 7.17)} \\ &= |C_r^1| - |C_{r+1}^0 \cap C_r^1| + \delta_{c,r+1} \\ &= Z_{c,r} - \psi_{c,r+1} + \delta_{c,r+1} \text{ (c.f. equation 7.18)} \end{aligned}$$

En utilisant l'hypothèse de récurrence, on a :

$$\begin{aligned} Z_{c,r+1} &= \left(\sum_{k=0}^r \delta_{c,k} + \delta_{c,r+1} \right) - \left(\sum_{k=0}^r \psi_{c,k} + \psi_{c,r+1} \right) \\ &= \sum_{k=0}^{r+1} \delta_{c,k} - \sum_{k=0}^{r+1} \psi_{c,k} \end{aligned}$$

Ainsi, la relation est aussi vérifiée pour $r + 1$.

Conclusion : Puisque le théorème 7.5.1 est vérifié pour un rayon nul, et puisqu'en supposant qu'il est vérifié pour un rayon r , on peut déduire qu'il l'est aussi pour $r + 1$, on obtient que le théorème 7.5.1 est vérifié pour tout rayon $r \in \mathbb{Z}^+$. $\square$

Ce théorème est très utile si on veut réduire le temps de calcul de la détection. En effet, on a juste besoin de calculer deux quantités $\delta_{c,r}$ et $\psi_{c,r}$ pour tous les points $c \in \Omega$ et les rayons $R_{\min} < r < R_{\max}$. En fait, ces quantités peuvent facilement être calculées en utilisant les orientations des différents pixels. Si on considère une orientation γ donnée, on peut construire les deux éléments structurants suivants :

$$B_\gamma = \{(d, r) \in \mathbb{Z}^2 \times \mathbb{Z}^+ / (K_{d,\tilde{0}\gamma,r} = 1) \wedge (K_{d,\tilde{0}\gamma,r-1} = 0)\}$$

et :

$$\overline{B}_\gamma = \{(d, r) \in \mathbb{Z}^2 \times \mathbb{Z}^+ / (K_{d,\tilde{0}\gamma,r} = 0) \wedge (K_{d,\tilde{0}\gamma,r-1} = 1)\}$$

avec $\tilde{0}\gamma$ le point de Ω égal à $(0, 0)$ dont l'orientation est γ .

Ces éléments structurants nous permettent pour différents points q d'orientation γ de lister les centres c et les rayons r qui doivent être considérés pendant les calculs respectifs de $\delta_{c,r}$ et $\psi_{c,r}$. Ainsi, si on pré-calcule ces éléments structurants pour différentes orientations quantifiées, on peut choisir pour chaque pixel q l'élément structurant dont l'orientation est la plus proche de celle de q , et ainsi pour chaque couple (d, r) , mettre à jour respectivement les valeurs de $\delta_{q+d,r}$ et $\psi_{q+d,r}$ en incrémentant leur valeur de 1. Faire cela est juste un moyen de propager l'impact de q aux centres qui considèrent q comme pointant vers eux. L'algorithme 7.1 résume la procédure de traitement d'un volume de tomosynthèse du sein.

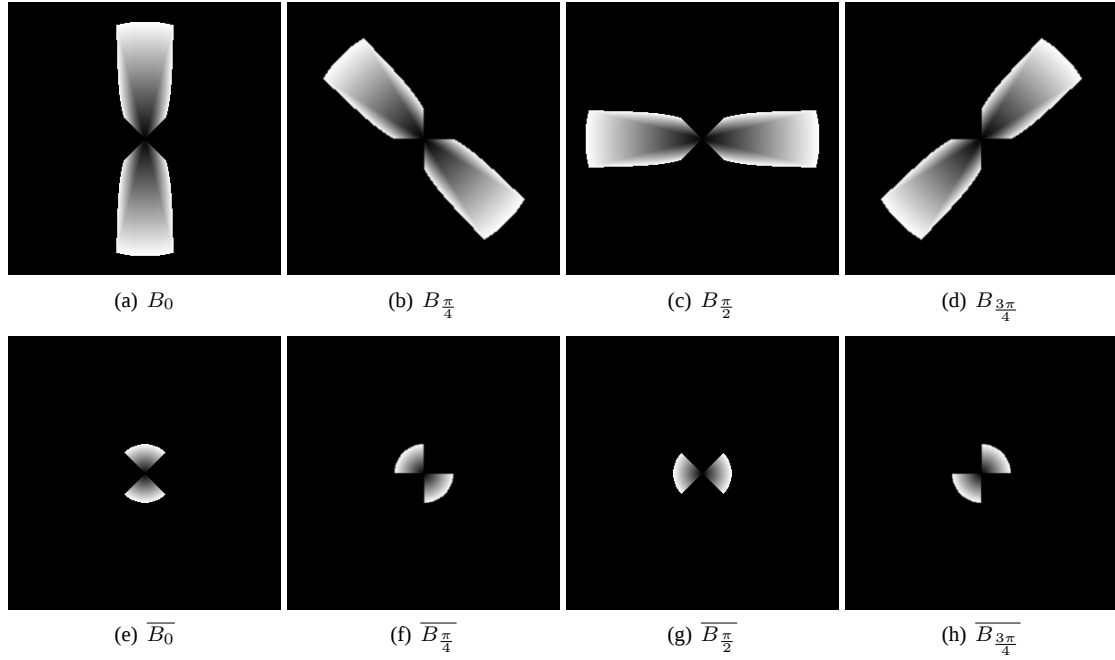


FIG. 7.12 – Exemple de B_γ et $\overline{B}_\gamma$ pour différentes orientations avec $R_{\max} = 130$ et $\alpha = 0.25$. Les niveaux de gris représentent le rayon associé à chaque différence de pixel. Les différences associées à un rayon nul n'appartiennent pas à l'élément structurant.

Théorème 7.5.2. $\forall d \in \mathbb{Z}^2, \forall \gamma \in [0, 2\pi[,$ il existe au plus un $r \in \mathbb{Z}^+$ tel que $(d, r) \in B_\gamma$ et au plus un $r' \in \mathbb{Z}^+$ tel que $(d, r') \in \overline{B}_\gamma$.

Démonstration. De part la définition de $K_{c,q,r}$, pour un couple (c, q) donné, soit $K_{c,q,r} = 0$ pour tout r , soit quand r croît, $K_{c,q,r}$ passe de 0 à 1, puis de 1 à 0 seulement une fois. $\square$

Les figures 7.12(d) et 7.12(h) illustrent respectivement B_γ et $\overline{B}_\gamma$, pour quatre orientations. Grâce au théorème 7.5.2, ils peuvent tous être deux représentés en utilisant une image à niveaux de gris où les niveaux de gris représentent les rayons associés aux décalages correspondants.

De plus, on peut rappeler que la détection de grandes lésions peut être effectuée en utilisant une approche multi-échelle, ce qui permet de réduire $R_{\max}$ et par conséquent la taille de ces éléments structurants.

7.5.2 Complexité

Intéressons-nous maintenant à la complexité de l'approche proposée dans le but de la comparer avec la mise en œuvre directe de l'équation 7.12.

Tout d'abord, pour chaque pixel q d'une coupe, nous avons besoin de traiter les points de B_γ et $\overline{B}_\gamma$ pour calculer de respectivement $\delta_{c,r}$ et $\psi_{c,r}$. En utilisant le théorème 7.5.2 et parce que B_γ est symétrique, comme montré à la figure 7.13(a), le nombre d'éléments de B_γ est donné par (c.f. annexe D) :

$$|B_\gamma| = 4 \left(\int_0^{\frac{\alpha R_{\max}}{\sqrt{\alpha^2 + 1}}} \sqrt{R_{\max}^2 - x^2} dx - \int_0^{\alpha R_{\max} \cos(\frac{\pi}{4})} x dx - \int_{\alpha R_{\max} \cos(\frac{\pi}{4})}^{\frac{\alpha R_{\max}}{\sqrt{\alpha^2 + 1}}} \sqrt{\frac{x^4}{(\alpha R_{\max})^2 - x^2}} dx \right) \quad (7.19)$$

$$\simeq O(R^2)$$

où le premier terme correspond à la sélection de points c qui peuvent atteindre le centre de l'élément structurant dans le pire cas, qui est la plus grande lésion possible (c.f. les zones non blanches dans la

Algorithme 7.1 : Mise en œuvre rapide de la détection de convergence dans un volume de tomosynthèse du sein.

```

pour chaque orientation quantifiée  $\gamma$  dans  $[0, 2\pi[$  faire
    calculer  $B_\gamma$  ;
    calculer  $\overline{B}_\gamma$  ;
fin
pour chaque coupe  $S$  du volume de tomosynthèse faire
    calculer les gradients non corrélés dans les directions  $x$  et  $y$  ;
    calculer la carte d'orientations à partir de vecteurs orthogonaux à ceux de la carte de gradients ;
    pour chaque  $p \in \Omega$  faire
        pour  $r = 0, r < R_{\max}, r++ = 1$  faire
             $\delta_{p,r} = 0$  ;
             $\psi_{p,r} = 0$  ;
        fin
    fin
    pour chaque  $q \in \Omega$  faire
        choisir  $B_\gamma$  et  $\overline{B}_\gamma$  en fonction de l'orientation sous  $p$  ;
        pour chaque  $(d, r) \in B_\gamma$  faire
             $\delta_{q+d,r}++ = 1$  ;
        fin
        pour chaque  $(d, r) \in \overline{B}_\gamma$  faire
             $\psi_{q+d,r}++ = 1$  ;
        fin
    fin
    pour chaque  $p \in \Omega$  faire
        pour  $r = 1, r < R_{\max}, r++ = 1$  faire
             $\delta_{p,r}++ = \psi_{p,r-1}$  ;
             $\psi_{p,r}++ = \psi_{p,r-1}$  ;
             $Z_{p,r} = \delta_{p,r} - \psi_{p,r}$  ;
            si  $Z_{p,r} > \lambda_r$  alors
                marquer  $(p, r)$  comme lésion potentielle ;
            fin
        fin
    fin
fin

```

figure 7.13(a)). Les deuxième et troisième termes permettent de supprimer les points c que le centre de l'élément structurant q ne peut atteindre soit parce que q serait dans le trou de l'anneau centré en c , soit parce qu'un trop grand rayon serait requis pour vérifier le critère de convergence. Ils sont tous deux représentés par les deux zones les plus sombres dans la figure 7.13(a).

Le nombre d'éléments dans $\overline{B}_\gamma$ est donné par (c.f. annexe D) :

$$\begin{aligned}
 |\overline{B}_\gamma| &= 4 \left(\int_0^{\frac{\alpha R_{\max}}{\sqrt{2}}} \sqrt{(\alpha R_{\max})^2 - x^2} dx - \int_0^{\frac{\alpha R_{\max}}{\sqrt{2}}} x dx \right) \\
 &= \frac{\pi \alpha^2 R_{\max}^2}{2} \\
 &\simeq O(R^2)
 \end{aligned} \tag{7.20}$$

Les deux termes sont illustrés dans la figure 7.13(b).

On obtient donc une complexité de $O(NR_{\max}^2)$, avec N le nombre de pixels dans l'image. Dans le cas d'une mise en œuvre directe de l'équation 7.12, pour chaque pixel, on a besoin de considérer $R_{\max}$ voisinages en forme d'anneau menant au nombre d'opérations suivant :

$$N \sum_{r=R_{\min}}^{R_{\max}} \pi r^2 - \pi(\alpha r)^2$$

soit une complexité de $O(NR_{\max}^3)$.

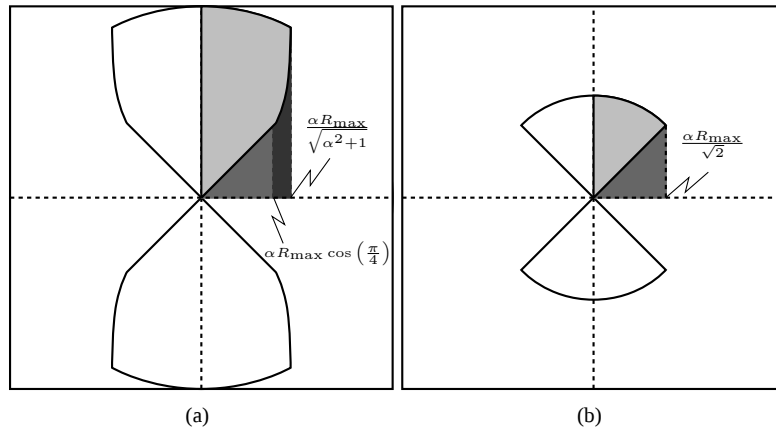


FIG. 7.13 – (a) Calcul de $|B_\gamma|$ pour une orientation verticale : la zone en gris clair correspond à $\frac{|B_\gamma|}{4}$ alors que les zones plus sombres correspondent aux deuxième et troisième termes de l'équation 7.19. (b) Même représentation pour $|\overline{B}_\gamma|$.

Finalement, on peut commenter l'impact de la constante α . Dans le cas d'une mise en œuvre directe, quand sa valeur est proche de 0, la surface de l'anneau augmente, alors que dans l'approche proposée, le comportement inverse est observé : $|B_\gamma|$ et $|\overline{B}_\gamma|$ décroissent. Cela impacte l'accroissement du temps de calcul pour les deux méthodes : quand α augmente le temps de calcul de la méthode naïve s'alourdit contrairement à celui de la méthode rapide.

7.6 Réduction de faux positifs

Les lésions et les distorsions architecturales ne sont pas les seules structures qui peuvent être décrites en utilisant ce motif de convergence dans les données de tomosynthèse du sein. Parfois, des fibres ou d'autres structures peuvent aussi être détectées par la précédente approche a contrario. Pour ne pas prendre en compte ces détections, les événements ϵ -significatifs sont groupés puis classifiés. Finalement, les zones suspectes restantes sont regroupées en 3D.

7.6.1 Agrégation des événements ϵ -significatifs

Parce qu'une zone de convergence n'est souvent pas détectée pour un unique centre c et un unique rayon r , certains événements détectés à la première étape de détection peuvent être associés. Cette étape vise à simplifier le processus de réduction de faux positifs. Cette étape d'agrégation est faite en calculant une carte A comme suit :

$$A = \bigcup_{(c,r)/Z_{c,r} \geq \lambda_r} D_{c,\alpha r} \quad (7.21)$$

avec $D_{c,r}$ le disque de rayon r centré au point c . Finalement, un étiquetage des composantes connexes pour chaque coupe permet d'extraire les zones suspectes pour vérifier si ce sont des vrais ou des faux positifs.

7.6.2 Différenciation lésion/croisement de fibres

Pour réduire le nombre de faux positifs, nous proposons l'analyse d'histogrammes construits à partir d'orientations de structures dans les régions agrégées dans l'étape précédente. Contrairement à l'étape de détection de convergences ϵ -significatives, les orientations sont obtenues cette fois à partir de dérivées secondes de gaussienne comme proposé par Karssemeijer et te Brake (1996), dans le but d'être plus robuste

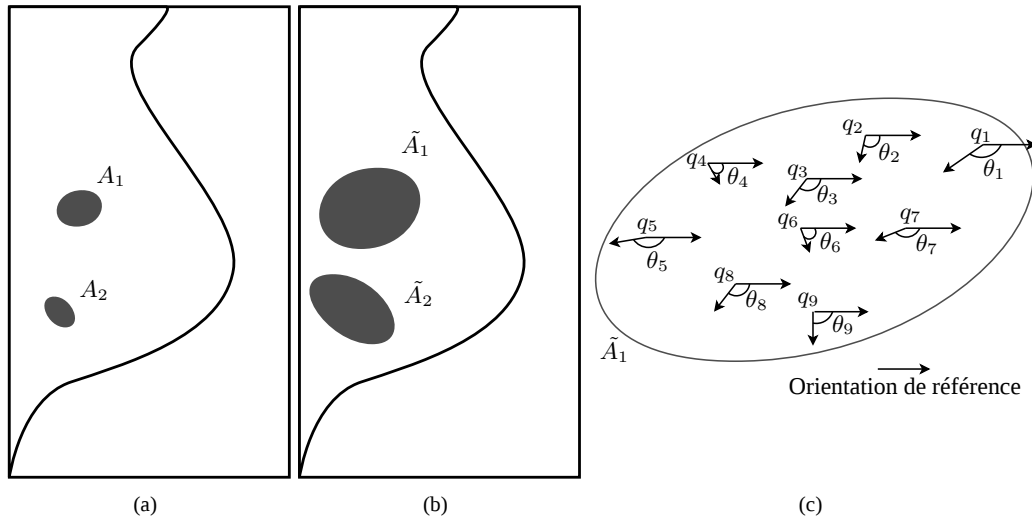


FIG. 7.14 – Analyse des orientations dans une zone de l’image. (a) Illustration de l’équation 7.21 : agrégation A des zones détectées par l’approche a contrario : deux composantes connexes A_1 et A_2 sont présentes dans la coupe de sein. (b) Illustration de l’équation 7.22 : zones $\tilde{A}_1$ et $\tilde{A}_2$ utilisées pour calculer les histogrammes d’orientations. (c) Illustration des orientations des structures représentées par les angles θ_k utilisés lors du calcul de l’histogramme pour un $\tilde{A}_i$ donné.

au bruit. La mesure que nous proposons est l’entropie des orientations des structures. Son calcul se fait dans la zone $\tilde{A}_i$ qui dépend de la composante considérée A_i de A :

$$\tilde{A}_i = \bigcup_{(c,r)/(c \in A_i) \wedge (Z_{c,r} \geq \lambda_r)} D_{c,r} \quad (7.22)$$

$\tilde{A}_i$ est ensuite utilisé pour construire un histogramme d’angles correspondant aux orientation pour tous les points $p \in \tilde{A}_i$. Les éléments de cet histogramme sont pondérés par l’amplitude de la carte d’orientations. Ce processus est illustré à la figure 7.14.

Cette analyse permet d’extraire les orientations principales des structures qui sont présentes dans $\tilde{A}_i$. Ainsi, on peut espérer qu’un croisement de fibres produira un histogramme contenant des pics alors qu’une lésion spiculée donnera une représentation plus homogène des orientations. Avec une telle interprétation, il semble naturel de choisir l’entropie pour distinguer les deux cas.

7.6.3 Illustration

La figure 7.15 illustre le résultat de la détection a contrario (c.f. figure 7.15(b)) suivi de l’utilisation de la précédente mesure pour réduire les faux positifs (c.f. figure 7.15(c)) sur une coupe de volume de tomosynthèse du sein (c.f. figure 7.15(a)). Les délimitations des zones détectées correspondent à l’agrégation des disques de rayon αr et r proposé par la méthode a contrario. Dans cet exemple la mesure d’entropie des orientations permet de supprimer la détection d’un croisement de fibres tout en conservant la lésion présente dans l’image.

Les histogrammes correspondant à la lésion et au faux positif sont respectivement présentés aux figures 7.16(b) et 7.16(a). La différence principale qui peut être remarquée entre ces deux graphiques est que celui du faux positif possède deux pics privilégiés dans son histogramme (un vers $-0,1$ et un second vers $\pm \frac{\pi}{2}$), alors que la lésion n’a pas de réelle orientation privilégiée. Dans cet exemple, l’histogramme vérifie bien ce à quoi on peut s’attendre en observant les données : le faux positif est dû à un croisement de deux fibres et les spicules de la lésion proviennent de directions uniformément distribuées.

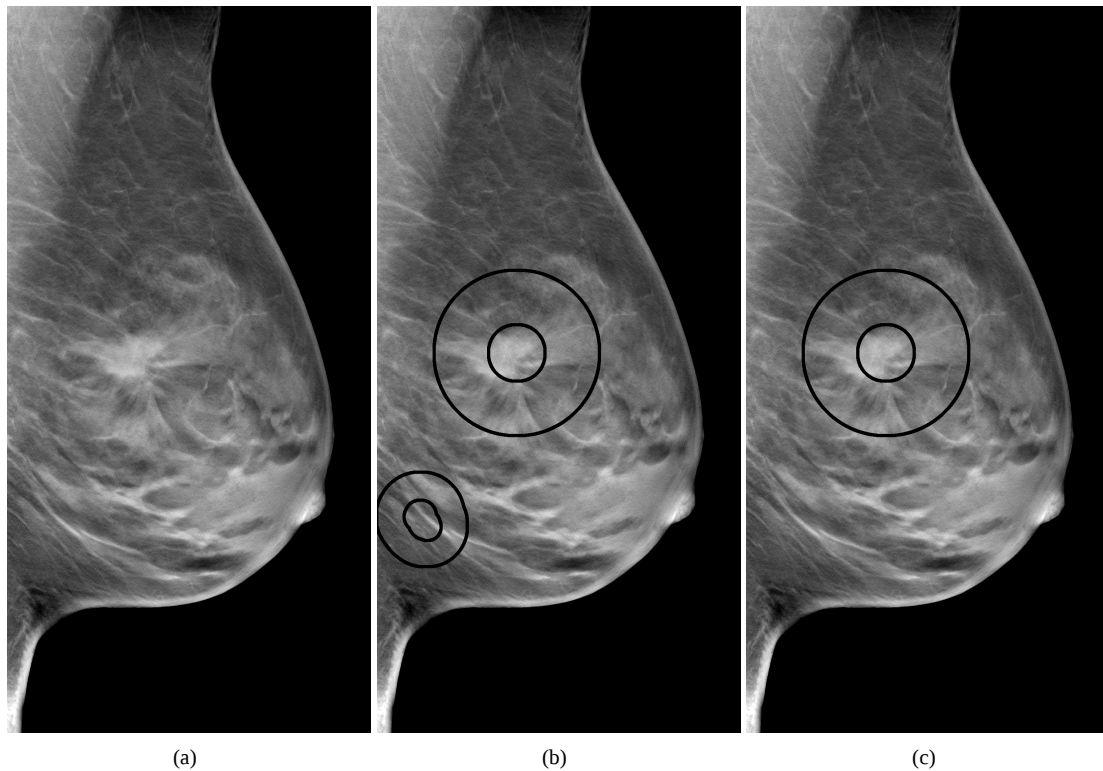


FIG. 7.15 – (a) Coupe d'un sein reconstruit contenant une lésion. (b) Résultat de la détection a contrario. (c) Résultat après réduction de faux positifs.

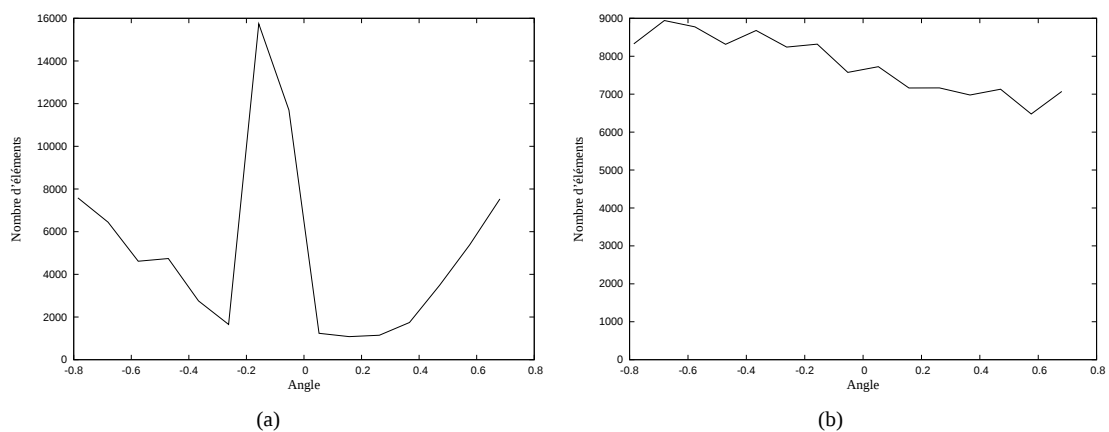


FIG. 7.16 – Exemple d'histogramme d'angles pour le faux positif (a) et pour la lésion (b) présentés à la figure 7.15(b).

7.7 Conclusion

Dans ce chapitre, nous avons abordé le problème de la détection de zones de convergence en tomosynthèse du sein. Ces zones pouvant être un signe de présence de cancer, il s'agit d'un nouvel outil dans le cadre d'un système d'aide à la détection. Seules les définitions théoriques, les propriétés et la question de leur application en pratique ont été considérées. L'évaluation des performances sera étudiée dans le chapitre 8.

Les travaux proposés reposent sur la modélisation a contrario (Desolneux *et al.*, 2000) de critères de convergence similaires à ceux qu'ont proposés Karssemeijer et te Brake (1996). La méthode proposée peut se résumer ainsi : on définit ce qu'est le contenu d'un sein ne contenant pas de lésion, puis on en déduit les valeurs peu probables de notre mesure de convergence. Cela nous permet d'obtenir un ensemble de seuils qui seront utilisés pour la détection de zones suspectes.

Une mise en œuvre rapide a aussi été proposée permettant ainsi l'utilisation de l'approche en pratique sur des volumes de tomosynthèse du sein. Cette dernière repose sur une formulation récursive du problème qui est inspirée des techniques de programmation dynamique. Brièvement l'algorithme proposé calcule les différences de convergence ou non convergence lorsque la taille supposée de la lésion augmente. Cela permet de réduire le temps de calcul de toutes les mesures de convergence en passant d'une complexité de l'ordre de $O(NR^3)$ à une complexité de l'ordre de $O(NR^2)$, avec N le nombre de pixels à considérer et R le rayon de la lésion de plus grande taille que l'on souhaite détecter.

Enfin un nouveau critère a été proposé pour la réduction de faux positifs. Ce critère repose sur l'analyse de l'histogramme des orientations dans les zones détectées par l'approche a contrario. L'idée est de différencier les croisements de fibres, qui présentent des histogrammes avec quelques pics, des lésions dont les histogrammes sont plus homogènes.

Des pistes d'amélioration peuvent être envisagées comme le raffinement du modèle naïf proposé. En effet, on pourrait considérer la non uniformité exacte de la distribution des angles induite par la quantification des niveaux de gris ou encore la corrélation entre différents niveaux de gris de pixels voisins dans des coupes de tomosynthèse. Cette corrélation peut provenir de la méthode de reconstruction utilisée ou d'une potentielle étape de filtrage avant extraction d'orientation. Dans le même esprit, on pourrait adapter les seuils λ_r en fonction du contenu de l'image, c'est-à-dire en définissant un modèle naïf dépendant de cette dernière.

Chapitre 8

Détection automatique de lésions malignes en mammographie 3D

La tomosynthèse du sein peut potentiellement pallier les limites de la mammographie conventionnelle, notamment en atténuant l'effet de masquage de lésions ou de construction de signes suspects par superposition de tissus. La contrepartie à ces potentielles améliorations est l'accroissement de la quantité de données à revoir pour le radiologue. Dans ce contexte, un système de détection automatique est plus que jamais d'actualité. En effet il pourrait permettre de préserver, voire d'améliorer, les performances cliniques du radiologue tout en diminuant le temps de lecture des images malgré l'accroissement du volume de données.

Dans ce chapitre, nous détaillerons une chaîne d'aide à la détection automatique d'opacités. Nous commencerons par une vue d'ensemble, puis nous détaillerons chacun des deux canaux composant l'approche. Nous finirons par une conclusion sur les performances globales de l'approche.

8.1 Description globale de la méthode

Nous allons commencer par une description assez haut niveau de l'approche que nous proposons dans ce chapitre. Nous détaillerons aussi la base de données utilisée pour l'évaluation des performances.

8.1.1 Approche

Pour traiter un volume, nous proposons une chaîne de traitement travaillant directement dans le volume de tomosynthèse reconstruit. La chaîne de traitement se décompose en deux sous-parties qui sont indépendantes. Le premier canal est destiné à la détection des opacités caractérisées par une surdensité alors que le deuxième canal est voué à détecter les zones de convergence. En y ajoutant un canal de détection de foyers de microcalcifications (Bernard *et al.*, 2008), le système obtenu permettrait de détecter une grande partie des signes de cancer que l'on peut retrouver en tomosynthèse : seules quelques lésions marginales ne sont pas prises en compte comme on le verra plus tard.

Le canal de détection de densités repose sur les travaux relatifs aux filtres connexes flous présentés dans les premiers chapitres de ce manuscrit. De manière générale le processus se décompose de manière séquentielle comme suit : le volume est tout d'abord sous-échantillonné, puis un filtre connexe flou est appliqué pour obtenir une carte floue des zones potentiellement suspectes. Chaque zone suspecte est ensuite segmentée dans le but d'extraire des caractéristiques qui serviront à confirmer ou infirmer la suspicion. Cette étape de segmentation/classification vise donc à réduire le nombre de fausses détections.

Le canal de détection de convergences est assez similaire. Le détecteur a contrario décrit au chapitre 7 est appliqué sur les coupes du volume en résolution native. Puis les zones suspectes sont regroupées en 3D pour ensuite être classifiées de manière similaire au précédent canal.

Les deux canaux ayant des populations cibles distinctes, l'étape de fusion est faite de manière directe : une zone dans le volume est considérée comme suspecte si elle a été détectée par au moins l'un des deux canaux. La figure 8.1 illustre la chaîne complète de détection.

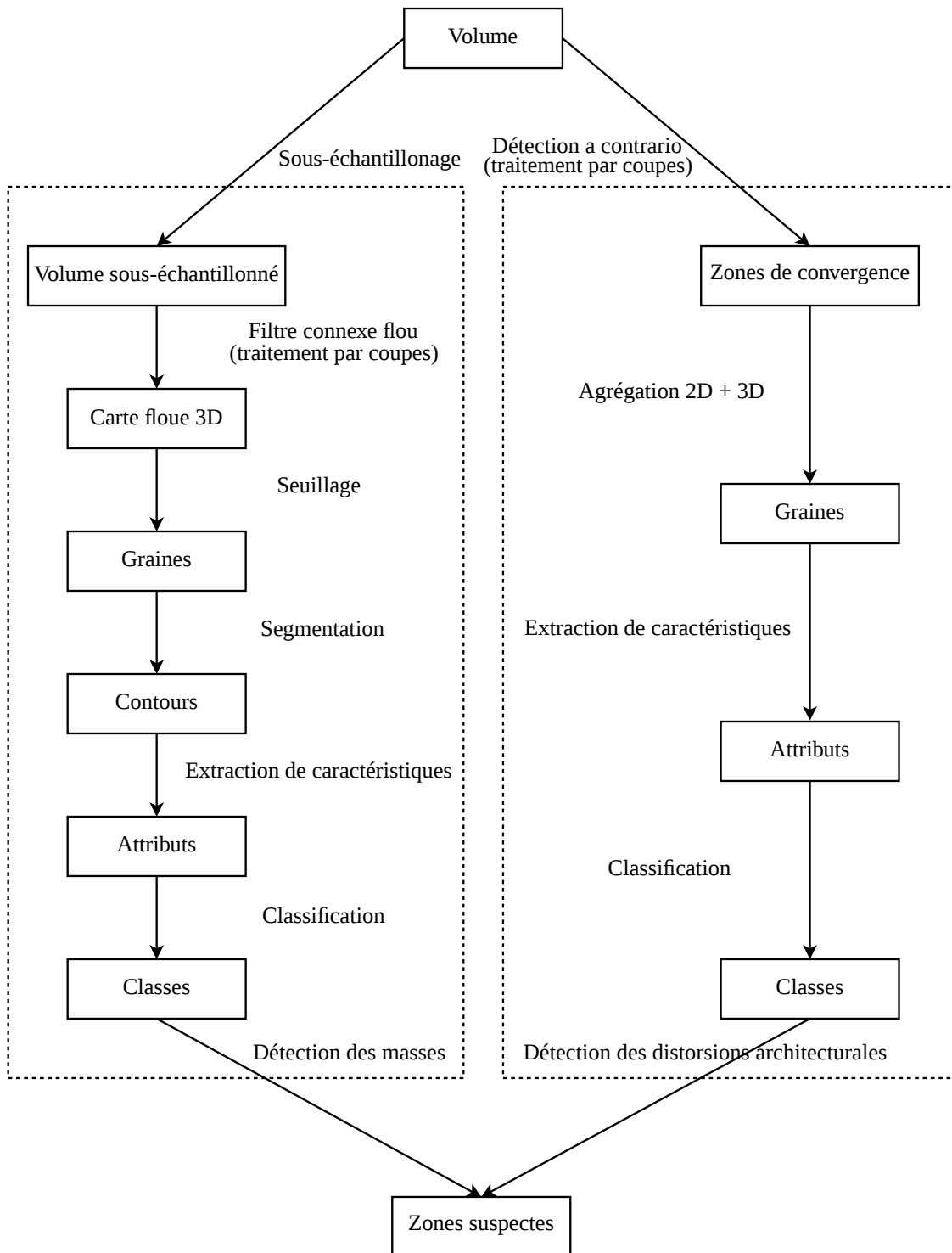


FIG. 8.1 – Schéma global de la chaîne de détection pour les opacités et les distorsions architecturales.

8.1.2 Base de données

La base de données utilisée dans ce chapitre pour la validation des différents canaux est composée de 53 examens unilatéraux de tomosynthèse contenant une ou plusieurs lésions malignes (56 au total) ainsi que

48 examens unilatéraux de tomosynthèse ne contenant aucune lésion. Les images utilisées pour les calculs de spécificité ne contiennent aucune zone suspecte repérée par les radiologues alors que les images utilisées pour les calculs de sensibilité contiennent au moins une lésion vue par des radiologues et dont la malignité est prouvée par biopsie.

Parmi les lésions malignes, on trouve plusieurs types de lésions. Ainsi, on dispose de sept lésions irrégulières, quatre lobulées, trente neuf spiculées, trois distorsions architecturales. Parmi les lésions spiculées, dix présentaient un motif stellaire prédominant par rapport au noyau dense.

Dans le but de tester les deux canaux proposés précédemment, la base de données a été scindée en deux. Une première partie contient les distorsions architecturales ainsi que les dix lésions spiculées à forte convergence, soit un total de treize lésions. Ce premier ensemble a été utilisé pour la validation du canal de détection de convergences. Les lésions irrégulières, lobulées et spiculées restantes, soit quarante lésions, ont été utilisées pour valider le canal utilisant la détection de sur-densités. Remarquons enfin que trois lésions non représentatives n'ont pas été utilisées pour l'étape de validation. En effet, ces dernières ne présentant pas vraiment de caractéristiques communes avec le reste de la base de données, il aurait fallu les traiter à l'aide d'un troisième canal, dont la validation aurait été impossible de par le peu de données disponibles. Ces lésions n'étant pas détectables par les travaux présentés, l'impact se traduira par une baisse de sensibilité d'environ 0,05 à déduire des performances présentées. Un exemple d'une telle lésion est présenté à la figure 8.2.

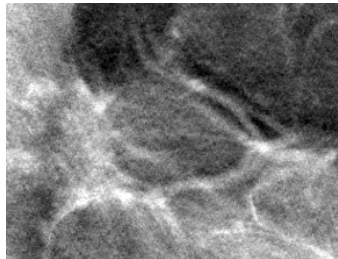


FIG. 8.2 – Exemple de lésion non incluse dans l'évaluation des performances.

8.2 Détection des lésions à noyaux denses

Le premier canal de la chaîne de détection proposée repose sur la détection de sur-densités dans l'image et de leur classification. Nous allons décrire ici les différentes étapes composant ce module, à savoir l'utilisation des filtres connexes flous pour la détection des sur-densités, la segmentation des lésions potentielles et leur classification.

8.2.1 Détection des densités

La détection des sur-densités s'effectue sur une version sous-échantillonnée du volume original. Cette étape permet de réduire le temps de calcul n'impacte pas la qualité de la détection, dans la mesure où les opacités sont généralement assez grosses en comparaison par exemple à des microcalcifications où une haute résolution est nécessaire pour leur visualisation.

Filtre de détection

Pour effectuer cette étape de détection, des filtres connexes flous tels qu'introduits au chapitre 2 ont été utilisés. L'avantage de ce type de filtre est double. Tout d'abord, ils permettent de manipuler l'incertitude sur la forme des composantes connexes floues ainsi que l'incertitude sur la réponse aux critères de sélection. L'utilisation de tels filtres durant l'étape de marquage se justifie par le fait qu'il est raisonnable de s'attendre à ce que les sur-densités dans une image de tomosynthèse puissent grossièrement être récupérées par différents seuillages. En effet ce type d'images permet en théorie d'atténuer, si ce n'est supprimer, les

superpositions de tissus que l'on observe en mammographie standard. On se rapproche ainsi du cas idéal d'un objet à détecter sur un fond uniforme.

Dans notre mise en œuvre, nous avons utilisé une version simplifiée du formalisme des images d'ombres floues : les images utilisées sont en fait des images d'ombres classiques. Ainsi, le filtre proposé ne tire que l'avantage de réponse graduée induite par le flou en n'utilisant pas l'éventuelle imprécision provenant des niveaux de gris. Pour justifier ce choix, on peut remarquer deux choses. Premièrement, dans notre application, l'étape de marquage est une étape grossière, où l'on souhaite obtenir une forte sensibilité même si la spécificité n'est pas très bonne. Dans ce contexte, les critères utilisés sont assez lâches et une analyse précise prenant en compte l'imprécision sur la forme n'est donc pas forcément nécessaire : l'étape de caractérisation prendra place plus tard dans la chaîne. Le deuxième argument pour justifier cette simplification se situe au niveau du temps de calcul. En effet, même si l'on a proposé au chapitre 3 un algorithme générique permettant de filtrer rapidement une image, le fait de travailler avec une image d'ombres qui est nette en entrée permet d'obtenir un traitement encore plus rapide : une représentation sous forme d'un unique arbre permet de filtrer l'image. En pratique on s'aperçoit que l'utilisation d'images d'ombres est assez robuste. On a d'ailleurs proposé un exemple n'introduisant aucune imprécision dans une telle image au chapitre 2.

Plus concrètement, le filtre utilisé sur les différents niveaux de gris considère toutes les composantes connexes que l'on peut extraire à partir de l'image pour chacun de ces derniers et conserve uniquement celles qui vérifient simultanément un critère de taille, de compacité et de contraste. Ce dernier se définit mathématiquement de la manière suivante :

$$\forall F \in \mathcal{F}, \forall f \in \mathcal{S}, \forall p \in \Omega \quad \psi^F(f)(p) = \begin{aligned} & \max_{h \in \mathcal{H}(f)} (\min(h(p), t_{u_1}(fcard(h)))) \\ & \top \max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_2}(fcomp(h)))) \\ & \top \max_{h \in \mathcal{H}(f)} (\min(h(p), r_{u_3}(frose(F, h)))) \end{aligned}$$

avec t_u et r_u des fonctions rampes de paramètres u , $fcard$, $fcomp$ les opérateurs définis au chapitre 2 et $frose$ l'opérateur suivant :

$$\forall F \in \mathcal{F}, \forall f \in \mathcal{S}, \forall h \in \mathcal{S} \quad frose_h(F, f) = \frac{fctrast_h(F, f)}{fmean(F, D_h(f) \cap \bar{f})} \sqrt{fcard(f)fmean(F, D_h(f) \cap \bar{f})}$$

Le filtre ψ^F peut s'interpréter de la manière suivante : pour chaque point du domaine de définition, le degré associé correspond à la t-norme $\top$ entre les degrés issus des différents critères calculés sur les différentes composantes connexes extraites.

Le filtre de détection appliqué aux coupes I du volume s'exprime alors $\phi(I) = agg(\delta(um(I)))$, avec $\forall p \in \Omega, \forall g \in \mathcal{G}, \delta(um(I))(p, g) = \psi^{um(I)}(um(I)(*, g))(p)$. Comme au chapitre 2, l'opérateur agg permet l'agrégation des résultats de filtrage pour les différents niveaux de gris. Ainsi la détection en un pixel de l'image sera l'agrégation des degrés d'appartenance obtenus pour chaque niveau de gris en ce même pixel.

On remarquera que l'opérateur ψ^F est similaire au filtre ψ_{mass}^F introduit au chapitre 2. La seule différence se situe au niveau de la mesure $fctrast$ qui est remplacée par $frose$. Cette dernière mesure est en fait une version floue du rapport signal à bruit exprimé dans le modèle de Rose, c'est-à-dire une formalisation où la taille de l'objet est prise en compte (Beutel *et al.*, 2000). Originellement ce critère a été défini pour juger de la détectabilité d'une structure circulaire de taille connue dans une image corrompue avec un bruit de Poisson. Par extension on peut donc l'utiliser pour justement détecter des structures d'intérêt. Bien que le contenu des coupes de tomosynthèse ne soit pas corrompu avec un bruit de Poisson, il est apparu expérimentalement que cette mesure donnait de meilleures performances que le contraste flou.

Illustration

La figure 8.3 illustre le type de résultats que l'on obtient en appliquant le filtre sur un volume de tomosynthèse du sein.

8.2.2 Extraction de marqueurs à partir de la carte de détection floue

Une fois toutes les coupes filtrées, une étape d'extraction de marqueurs est appliquée dans le but de mieux caractériser les zones détectées et ainsi réduire le nombre de faux positifs.

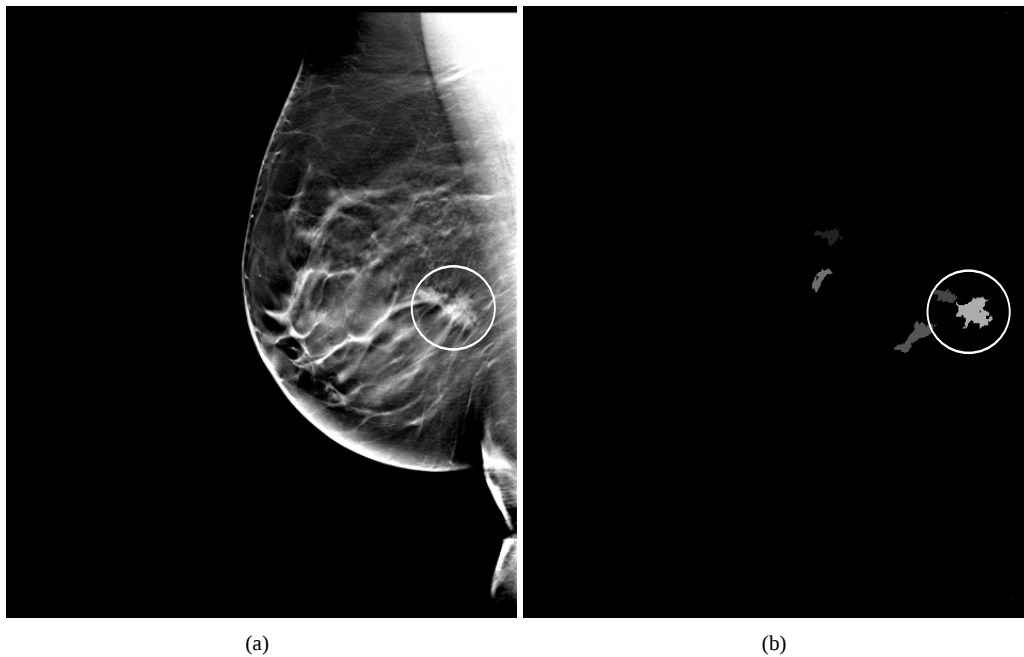


FIG. 8.3 – Illustration de l'étape de marquage de noyaux dense sur un examen de tomosynthèse contenant une lésion spiculée. (a) Coupe d'un volume contenant une lésion spiculée (entourée en blanc). (b) Carte de détection floue.

Seuillage adaptatif

Une approche par seuillage adaptatif a été mise en œuvre pour l'extraction de marqueurs à partir de la carte de détection floue obtenue à l'étape précédente. En effet, la texture pouvant varier d'un sein à un autre, un seuillage unique ne permet pas de traiter de manière satisfaisante, c'est-à-dire avec une spécificité raisonnable, un ensemble assez large de d'examen : certains volumes ne produisent que un ou deux marqueurs alors que d'autres peuvent en produire vingt fois plus.

La stratégie utilisée consiste donc à durcir le seuil tant que le nombre de détections dans le volume est supérieur à ce que l'on souhaite accepter. Ce nouveau paramètre va permettre de régler de manière plus intuitive et plus efficace les performances de l'étape de marquage. Le nombre de détections correspond au nombre de zones agrégées par la méthode présentée plus loin.

Pour illustrer la dépendance entre nombre de détection et composition du sein, on peut se reporter à la figure 8.4 qui présente les distributions estimées sur 122 cas du nombre de détections non limitées, c'est-à-dire en utilisant un seuil fixe assez lâche, pour différentes densités de sein. La classification sur une échelle allant de 1 à 4 proposée par l'*American College of Radiology* (ACR, 2003) a été utilisée pour rendre compte de ces densités. Les densités ont été évaluées par des radiologues sur les volumes de tomosynthèse. Il faut remarquer que les histogrammes présentés sont fortement régularisés. De plus les densités 1 et 4 étant plus rares que les autres, leurs histogrammes, qui ne sont calculés que sur quelques cas, ne sont que peu représentatifs. Dans cet illustration on s'aperçoit que les seins les plus denses semblent produire un plus grand nombre de détections en moyenne. Il semble donc raisonnable de s'adapter aux données lors de l'étape de marquage.

Agrégation 3D

Pour tirer parti de l'information 3D, les composantes connexes issues de l'étape de marquage réalisée par le biais du seuillage adaptatif ont été regroupées en 3D. Ce regroupement est effectué en utilisant la connexité entre la détection pour différentes coupes consécutives couplée à une notion de similarité de marqueurs. Cette notion de similarité est exprimée de manière pratique en introduisant une notion de

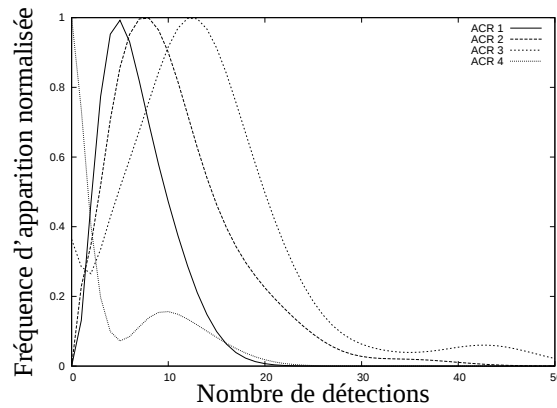


FIG. 8.4 – Histogrammes normalisés du nombre de détections pour des seins de densités différentes.

variation entre différents marqueurs ($\{C_z\}$) :

$$\text{variation}(\{C_z\}) = \frac{\max_z |C_z|}{|\bigcup_z C_z|}$$

avec $|\cdot|$ le cardinal d'un marqueur.

De manière pratique, pour agréger les zones marquées dans le volume, on considère tous les couples de marqueurs 2D, extraits à partir des différentes coupes, qui sont connexes entre eux selon l'axe des Z . On agrège chacun de ces couples si le marqueur 3D formé par l'agrégation des deux marqueurs 2D considérés a une variation inférieure à un certain seuil. De cette manière on s'assure de ne pas considérer par la suite des marqueurs qui ont une forte dissimilarité d'une coupe à l'autre. En pratique le seuil est suffisamment sélectif pour éviter toute agrégation abusive. Cela peut cependant résulter en une non agrégation de marqueurs qui devraient l'être. Ce comportement n'est toutefois pas gênant puisque, s'il s'agit d'une unique lésion, cette dernière n'est pas perdue et sera dans le pire des cas marquée plusieurs fois.

Performances

Cette étape de marquage a été évaluée par validation croisée sur la base de données contenant les lésions lobulées, irrégulières et les lésions spiculées non utilisées pour l'évaluation de la détection de convergence. Au total, ce sont 40 lésions malignes qui ont été utilisées. La figure 8.5 illustre les différentes sensibilités associées aux spécificités obtenues en faisant varier le nombre maximal de marqueurs autorisés comme expliqué précédemment.

On voit que l'on obtient, à ce stade, une sensibilité de 82%, pour un nombre de faux positifs égal à 3,83 ou bien une sensibilité de 90% pour 5,12 faux positifs par volume. Cependant, on utilisera une sensibilité de 100% à cette étape du traitement afin de conserver tous les signes radiologiques potentiellement malins, en acceptant un nombre moyen de faux positifs de 6,8 par volume de tomosynthèse. Les étapes de segmentation, de mesure des caractéristiques et de classification vont permettre la diminution de ce taux.

8.2.3 Segmentation

Chaque marqueur issu de l'étape précédente prend la forme d'un ensemble de voxels dans le volume. Dans la mesure où les objets dans les volumes de tomosynthèse présentent généralement une forte distorsion selon l'axe des Z , les lésions potentielles ont été segmentées en 2D dans le plan (X, Y) . La coupe considérée est celle qui est la plus représentative. Pour déterminer cette dernière, nous avons procédé comme suit : une mesure de contraste a été calculée pour chaque coupe intersectée par le marqueur en faisant la moyenne entre l'intérieur et l'extérieur de ce dernier, et la coupe de plus fort contraste a été considérée comme la plus représentative.

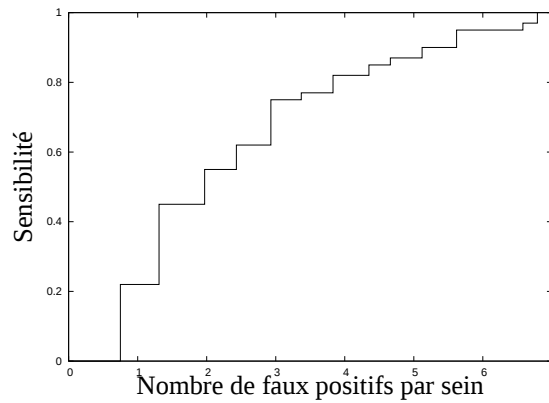


FIG. 8.5 – Courbe de performances de l'étape de marquage pour la détection de noyaux denses suspects.

Dans les évaluations de performances conduites sur cette approche, deux méthodes de segmentation ont été utilisées en remplacement l'une de l'autre. La première approche est fondée sur l'extraction d'un contour à partir de différents seuillages de l'image et la seconde repose sur la modélisation du contour par un chemin minimum dans un paysage énergétique déduit de l'image. Ces deux méthodes sont celles présentées au chapitre 5.

Remarquons néanmoins que les méthodes de segmentation utilisées ici ne sont pas les versions floues introduites dans le chapitre précédemment cité. En effet, pour des contraintes de temps de calcul, l'application présentée ici après étape de marquage ne manipule plus l'imprécision sur les détections.

Dans le cas où la segmentation provient de multi-seuillages, le contour considéré est le contour de plus fort degré d'appartenance. Dans le cas de l'approche par programmation dynamique, le premier chemin trouvé est le contour considéré.

8.2.4 Caractéristiques

Dans l'optique de caractériser les contours obtenus à l'étape précédente, plusieurs mesures ont été calculées sur ces derniers dans le but d'utiliser le classifieur qui sera détaillé dans la section suivante.

Parmi ces mesures, certaines n'utilisent que la morphologie du contour considéré, d'autres son interaction avec le contenu de l'image. Les mesures utilisées sont listées dans le tableau 8.1.

Le calcul de ces mesures sur les différents contours extraits à partir des marqueurs permet de constituer une base d'éléments à classifier. Dans la suite, nous les noterons ces vecteurs x_i , chaque dimension correspondant à une des mesures présentées au tableau 8.1.

8.2.5 Séparateur à Vaste Marge

Pour l'étape de classification, un classifieur classique de type séparateur à vaste marge (SVM) a été utilisé (Cristianini et Shawe-Taylor, 2000). Ce choix est motivé par le fait qu'un tel classifieur est assez standard et qu'il a déjà été utilisé avec succès dans des domaines assez variés dont, notamment, la classification de lésions en mammographie standard. Dans l'optique de comprendre son utilisation, nous allons le décrire de manière générale.

Séparation linéaire des données

Le principe des SVM est de trouver une séparation linéaire entre deux classes. De manière plus formelle, cette dernière peut s'exprimer à l'aide de la fonction $h(x)$ suivante qui prend en entrée un élément à classifier (vecteur contenant les différentes mesures pour la zone détectée à l'étape de marquage) et dont le signe identifie la classe d'appartenance (e.g. + pour suspect et - pour non suspect) :

$$h(x) = w^t x + w_0 \quad (8.1)$$

Mesures	Description
Homogénéité	Variance dans la région d'intérêt contenant le marqueur
Entropie	Entropie dans la région d'intérêt
Contraste	Contraste de la structure (différence entre moyenne à l'intérieur et à l'extérieur du contour)
RGI	Indice de radialité sur le contour (Kupinski et Giger, 1998)
RGI à l'intérieur du contour	Idem mais dans le contour
RGI dans la région d'intérêt	Idem mais dans la région d'intérêt
Gradient moyen sur le contour	Moyenne du gradient de l'image calculée sur le contour
Entropie des orientations dans le contour	Mesure statistique sur l'histogramme (pondéré ou non) des orientations dans le contour
Entropie des orientations sur le contour	Idem mais l'histogramme est calculé sur le contour
Entropie des orientations dans la région d'intérêt	Idem mais l'histogramme est calculé sur tout le voisinage contenant la lésion
Kurtosis, variance, moyenne, entropie, skewness des orientations relatives dans le contour	Mesures statistiques calculées sur un histogramme (pondéré ou non) d'orientations relatives, par rapport au centre de gravité du contour, dans le contour
Kurtosis, variance, moyenne, entropie, skewness des orientations relatives sur le contour	Idem mais l'histogramme est calculé avec les éléments présents sur le contour
Kurtosis, variance, moyenne, entropie, skewness des orientations relatives dans la région d'intérêt	Idem mais l'histogramme est calculé avec toutes les valeurs présentes dans la région contenant la lésion
Probabilité minimale de convergence	Mesure dérivée du formalisme décrit au chapitre 7 ($\min\{P[Z_{c,r} \geq z_{c,r}], r \leq R_{\max}, c \in \text{contour}\}$, avec $z_{c,r}$ la réalisation de $Z_{c,r}$ pour l'image considérée)

TAB. 8.1 – Mesures utilisées pour caractériser les lésions à noyaux denses.

avec w et w_0 les paramètres de la séparation linéaire.

Dans le cas où les données sont bien linéairement séparables, il peut exister une infinité de séparations possibles. Dans le cas des SVM, la séparation retenue est celle qui maximise la marge entre les éléments des deux classes proches de la frontière. La figure 8.6 illustre cette notion de marge. L'expression formelle de cette notion de marge maximale peut se reformuler de la manière suivante :

$$\arg \max_{w, w_0} \min_{j \in \{1; J\}} \{ \|x - x_j\| : x \in \mathbb{R}^N, w^t x + w_0 = 0 \}$$

avec N le nombre total de mesures et J le nombre total d'éléments dans la base d'apprentissage.

De manière concrète, la résolution de ce problème se fait par une optimisation par la méthode des multiplicateurs de Lagrange. En effet, on peut montrer que la marge peut s'exprimer par $\|w\|^{-1}$ à une normalisation près. Ainsi on souhaite minimiser $\|w\|$ sous la contrainte de bonne classification des éléments d'apprentissage ($\forall j, l_j(w^t x_j + w_0) \geq 1$, l_j étant la classe (+1 ou -1) du $j^{\text{ème}}$ élément de la base d'apprentissage) :

$$L(w, w_0, \alpha) = \frac{1}{2} \|w\|^2 - \sum_{j=1}^J \alpha_k (l_j(w^t x_j + w_0) - 1)$$

où le premier terme correspond à la marge et le second aux contraintes de classification.

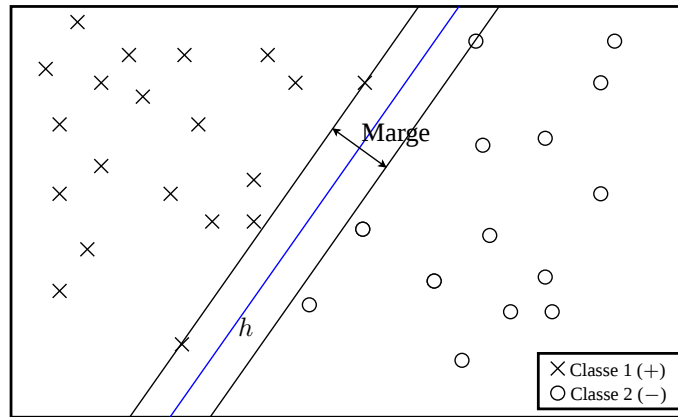


FIG. 8.6 – Illustration des SVM dans le cas de données linéairement séparables

Ainsi, pour les coefficients de Lagrange optimaux α_k^* , on obtient l'expression de la séparation linéaire :

$$h(x) = \sum_{j=1}^J \alpha_j^* l_k(x \cdot x_k) + w_0 \quad (8.2)$$

Assouplissement de la contrainte d'erreur

Dans le but d'une meilleure généralisation, il est parfois préférable d'autoriser une erreur sur les données d'apprentissage. Dans cette optique, une amélioration des SVM a été proposée (Cortes et Vapnik, 1995). L'idée est d'introduire des coefficients ξ_i qui vont autoriser une mauvaise classification. La nouvelle version du problème se traduit par la minimisation pour w , w_0 et ξ_i de :

$$\frac{1}{2} \|w\|^2 - C \sum_{j=1}^J \xi_j$$

sous la contrainte suivante :

$$l_j(w^t x_j + w_0) \geq 1 - \xi_j \quad \xi_j \geq 0, \quad 1 \leq j \leq J$$

Le nouveau paramètre C est dépendant des données du problème de classification considéré. Son optimisation sera abordée plus tard.

Cas de données non linéairement séparables

Dans la plupart des problèmes réels, les données ne sont généralement pas linéairement séparables. Pour pallier cette limitation, on peut exprimer les vecteurs à classer dans un espace de dimension supérieur à celui d'origine qui pourra potentiellement permettre l'expression d'une séparation linéaire. En appliquant une transformation non linéaire ϕ aux éléments à classer, on peut réécrire l'équation 8.2 comme suit :

$$h(x) = \sum_{j=1}^J \alpha_j^* l_j(\phi(x) \cdot \phi(x_k)) + w_0 \quad (8.3)$$

L'astuce citée précédemment (Vapnik, 1995; Cristianini et Shawe-Taylor, 2000) repose sur l'utilisation d'un noyau K vérifiant $K(x_i, x_j) = \phi(x_i) \cdot \phi(x_j)$ qui permet de ne pas calculer le produit scalaire dans l'espace de dimension supérieure.

Dans les expérimentations présentées dans ce chapitre, le noyau utilisé est une fonction à base radiale ou noyau gaussien :

$$K(x_i, x_j) = e^{-\gamma \|x_i - x_j\|^2}$$

On remarquera l'apparition d'un nouveau paramètre γ . Tout comme le paramètre C introduit dans l'assouplissement des marges, son optimisation sera abordée ultérieurement.

Utilisation des SVM

Les différents points de fonctionnement présentés dans les courbes de performances qui seront détaillées plus loin ont été obtenues grâce à l'estimation de probabilités d'appartenance aux différentes classes comme décrit dans Wu *et al.* (2004a).

8.2.6 Résultats

Processus de validation

Pour valider les performances du canal de détection de densités, une approche de type validation croisée a été utilisée. Pour chaque volume de tomosynthèse, les autres volumes de la base ont été utilisés pour apprendre un SVM qui est ensuite utilisé pour tester le volume mis à l'écart. Pour l'étape d'apprentissage, une validation croisée a aussi été utilisée pour optimiser les paramètres C et γ du SVM et du noyau. Le schéma de la figure 8.7 illustre cette validation.

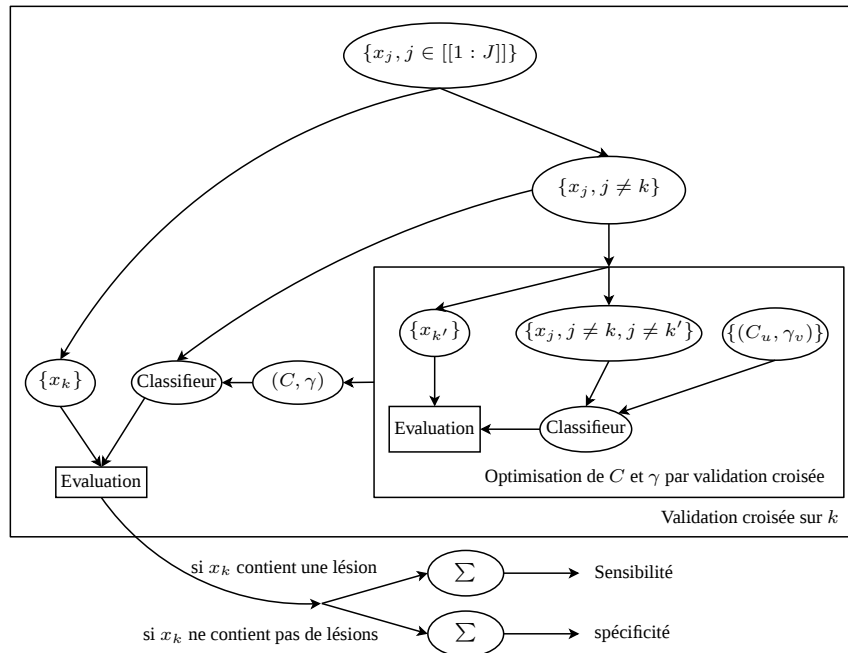


FIG. 8.7 – Validation croisée du canal de détection de densités.

Performances

La figure 8.8 présente les performances de l'étape de classification. Dans la mesure où on a fixé le point de fonctionnement à l'étape de marquage à une sensibilité de 1, cette figure représente la performance de toute la chaîne de détection d'opacités. Dans cette figure, on distingue deux courbes. La première correspond aux performances obtenues lorsque l'étape de segmentation est réalisée à l'aide de multi-seuillage des lésions potentielles. La seconde courbe correspond à l'approche de segmentation par programmation dynamique. On remarquera que les performances obtenues dans le second cas sont meilleures que dans le premier pour des sensibilités élevées. Cela peut s'expliquer de manière générale par une meilleure qualité des contours fournis par l'approche programmation dynamique. En effet, comme on a pu le voir, l'hypothèse de segmentation par multi-seuillage est qu'il n'y a plus de superposition de structures dans l'image.

Malheureusement, les volumes de tomosynthèse n'étant que le résultat d'une approximation grossière du théorème de Radon (1917), une certaine superposition des tissus persiste.

Sur cette même figure, on voit qu'après segmentation, extraction des caractéristiques et classification, le nombre moyen de faux positifs par volume de tomosynthèse est ramené à 1,23 pour une sensibilité de 90%.

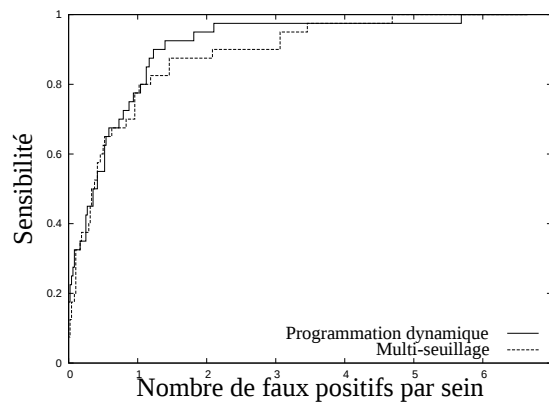


FIG. 8.8 – Courbe de performance du canal de détection de noyaux denses suspects.

La figure 8.9 illustre le canal de détection sur un sein contenant une lésions spiculée. On voit dans cet exemple que l'étape de marquage pointe plusieurs zones, dont la lésion, et que ces zones sont en grande partie supprimées par l'étape de réduction de faux positifs. A la fin de la chaîne il ne reste qu'un faux positif qui apparaît en partie sur la même coupe que la lésion.

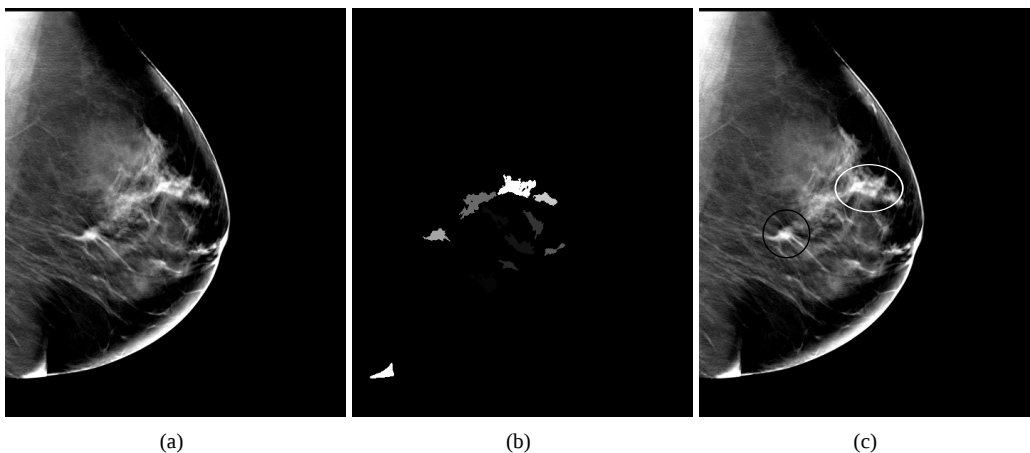


FIG. 8.9 – Exemple d'un traitement complet du canal de détection d'opacités. (a) Coupe d'un volume de tomosynthèse contenant une lésion spiculée. (b) Carte de détection floue issue de l'étape de marquage. (c) résultat après classification : la lésion (en noir) est détectée ainsi qu'un faux positif (en blanc).

8.3 Détection des lésions sans noyau dense

Certaines lésions ne se traduisent pas par l'apparition d'un noyau dense, mais seulement par un motif fort de convergence. Pour détecter ce type de lésions, nous avons besoin du deuxième canal évoqué à la section 8.1. Nous allons décrire maintenant la constitution de ce dernier, ainsi que sa validation sur une partie de la base de données décrite précédemment.

8.3.1 Détection de convergence

Pour détecter les convergences, nous utilisons l'approche a contrario décrite au chapitre 7 comme étape de marquage. Cette méthode nous permet de récupérer les distorsions architecturales de motif stellaire ainsi que des lésions spiculées à forte convergence. De par leur morphologie, les lésions irrégulières et lobulées ont très peu de chance d'être détectées. De plus, dans l'optique d'avoir une forte sensibilité associée à une spécificité raisonnable, on ne peut essayer de détecter toutes les lésions spiculées par ce critère de convergence. Pour illustrer cela on se référera à la courbe de détection de la figure 8.10. On remarquera que dans cette figure, l'échelle du nombre de faux positifs par sein a été contractée pour permettre le traçage de la courbe traduisant ainsi le grand nombre de faux positifs pour des sensibilités modestes.

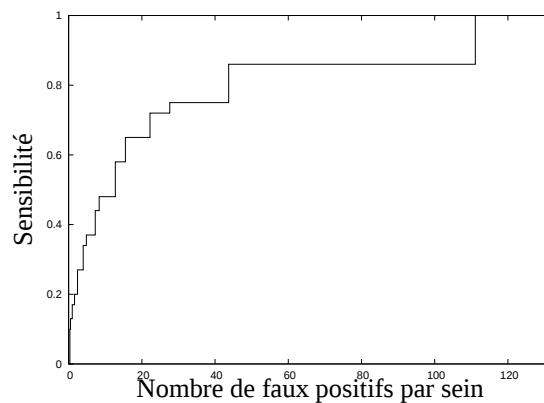


FIG. 8.10 – Courbe de performance du modèle a contrario pour des lésions uniquement spiculées.

Si maintenant on se restreint aux seules lésions présentant un motif stellaire, c'est-à-dire les distorsions architecturales en étoile et les lésions très fortement spiculées (13 lésions au total), on peut obtenir une première étape de marquage viable comme on peut le voir sur la figure 8.11. On note une sensibilité de 92% pour un taux de faux positifs par sein égal à 0,9.

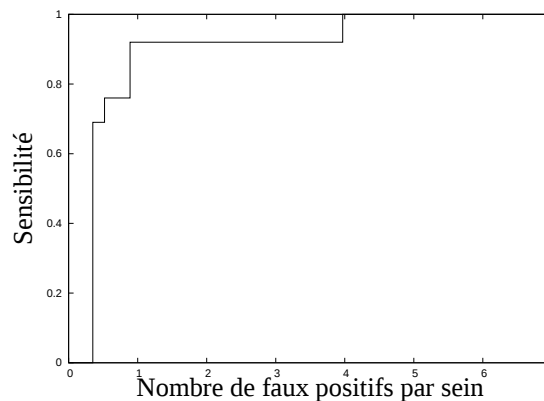


FIG. 8.11 – Courbe de performance du modèle a contrario pour les distorsions architecturales, et des lésions très fortement spiculées.

Les courbes des figures 8.10 et 8.11 ont été obtenues en faisant varier le paramètre ϵ représentant le nombre de fausses détections que l'on est prêt à accepter dans le modèle naïf (c.f. chapitre 7).

8.3.2 Réduction de faux positifs

Une méthode similaire à l'étape de classification utilisée dans le canal de détection de noyaux denses a été choisie pour réduire le taux de faux positifs. La différence principale réside dans le fait que pour les lésions sans noyau dense toutes les coupes sont traitées pour l'extraction de caractéristiques. En effet, les contractions de tissu laissant paraître un motif stellaire étant assez subtiles, une analyse par coupe fine permet une meilleure caractérisation.

Agrégation 3D

Dans le but d'utiliser l'information 3D non encore exploitée, les détections sur chacune des coupes sont regroupées en 3D. Pour ce faire, une carte de détection A est construite en agrégeant les disques de centre c et de rayon αr pour tous les couples (c, r) marqués comme suspects lors de la première étape. Les composantes connexes A^j de cette carte 3D sont ensuite étiquetées. Pour tout j , on notera $\{A_i^j\}$ l'ensemble des composantes connexes contenues dans A^j pour chacune des coupes que ce dernier traverse. Chaque composante A^j est ensuite gardée si l'étape de classification considère qu'au moins un A_i^j correspond à une lésion. De plus, le fait de garder la composante A^j en entier permet de ne pas couper la lésion s'il y a une mauvaise classification sur une coupe intermédiaire. Cette étape d'agrégation est illustrée à la figure 8.12.

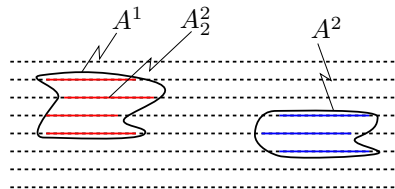


FIG. 8.12 – Illustration de l'étape d'agrégation des marqueurs pour le canal de détection de convergences.

Sélection de caractéristiques

Comme on l'a vu au chapitre 7, les faux positifs générés par l'approche a contrario correspondent souvent à des croisements de fibres normales. Pour différencier ces structures des potentielles lésions, des caractéristiques reposant sur l'analyse statistique (entropie) des orientations dans le voisinage du marqueurs ont été utilisées comme entrée pour le classifieur. Plus précisément, les mesures ont été réalisées dans trois zones différentes : dans la zone définie par l'agrégation des centres de convergence (disque de rayon αr), sur la frontière définie par cette même zone ainsi que dans la zone définie par l'agrégation des zones de convergence (disques de rayon r). A ces mesures ont été ajoutés un indice de radialité (Kupinski et Giger, 1998) sur la frontière de la zone de convergence ainsi que des mesures sur les rapports de tailles de volumes d'intérêt englobant les différents A^j (e.g. $\min(\text{longueur}, \text{largeur})/\text{profondeur}$).

8.3.3 Performance

L'évaluation des performances de ce canal a été conduite de la même manière que pour la chaîne de détection de densités (c.f. section 8.2.6). La validation a été faite sur les treize cas présentés au début de ce chapitre limitant ainsi les conclusions sur les performances réelles de ce canal. Néanmoins, cela permet d'avoir une idée sur la validité de la démarche.

Les performances obtenues sont synthétisées par la courbe de la figure 8.13. On peut remarquer que la spécificité est généralement meilleure que pour la détection de noyaux denses pour des sensibilités équivalentes. Ainsi, on atteint une sensibilité de 92% pour un taux de faux positifs égal à 0,48.

La figure 8.14 illustre sur un exemple la chaîne complète de détection de convergences suspectes. L'image de la figure 8.14(a) est une coupe de tomosynthèse de sein contenant une distorsion architecturale. La détection a contrario (c.f. figure 8.14(b)) met en évidence quatre zones de convergences suspectes dans cette coupe, l'une d'entre elles étant la lésion. La réduction de faux positifs (c.f. figure 8.14(c)) permet d'en éliminer deux. Ainsi au final, deux zones sont marquées : la lésion et un faux positif.

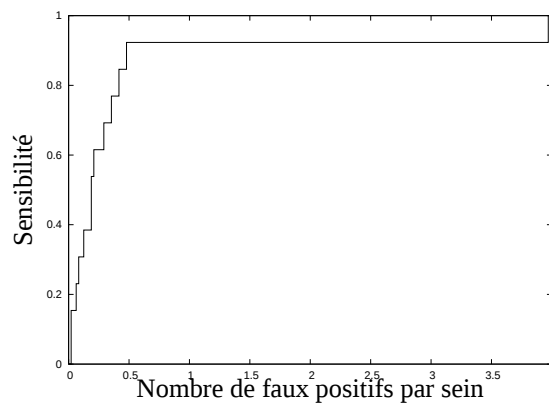


FIG. 8.13 – Performance du canal de détection de convergence suspectes.

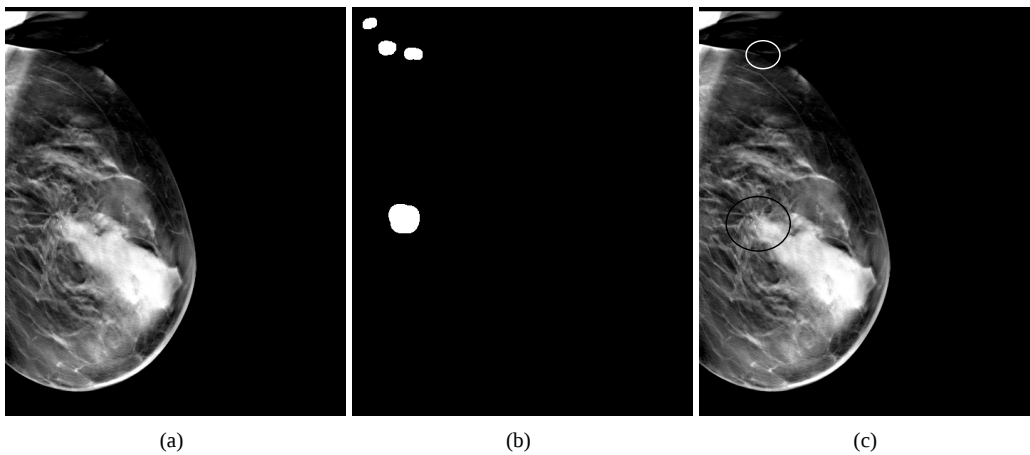


FIG. 8.14 – Exemple d'un traitement complet du canal de détection de convergences. (a) Coupe d'un volume de tomosynthèse contenant une distorsion architecturale. (b) Résultats de la détection a contrario. (c) résultat après classification : la lésion (en noir) est détectée ainsi qu'un faux positif (en blanc).

8.4 Performances globales de la chaîne

Nous allons maintenant nous intéresser à la fusion des résultats des deux canaux précédents. Après avoir présenté l'approche retenue, nous exposerons les performances complètes de notre chaîne de détection.

8.4.1 Fusion des deux canaux

Nous avons présenté précédemment les performances individuelles de chaque canal de notre approche de détection. Dans la mesure où ces canaux sont destinés à la détection de structures différentes, la méthode de fusion des résultats retenue est de type disjonctive, c'est-à-dire que l'on considère une zone suspecte si au moins l'un des deux canaux la considère comme suspecte.

Cette approche permet d'avoir rapidement des bornes en termes de sensibilité et de spécificité en combinant les courbes des figures 8.8 et 8.13. En effet, en regroupant les deux bases de données utilisées précédemment, pour une spécificité on est assuré d'atteindre au moins la somme pondérée respectivement par la proportion des cas malins dans chaque base, des sensibilités des deux canaux. De même, pour une sensibilité donnée on est assuré d'avoir au plus la somme des faux positifs obtenus pour les deux canaux.

Performances et discussion des résultats

Les performances présentées au tableau 8.2 ont été obtenues en utilisant les seuils pour les probabilités des classifieurs correspondant à des sensibilités équivalentes dans les courbes des figures 8.8 et 8.13. Les résultats des deux chaînes ont été fusionnés par l'application d'un *ou logique* entre les masques de détection obtenus par chacun des canaux. Le même processus de validation croisée qu'utilisé précédemment a été employé.

Sensibilité (%)	Spécificité (faux positifs par sein)
81,13	1,31
90,57	1,60
96,23	1,81

TAB. 8.2 – Performances de la chaîne complète après fusion des deux canaux.

On remarquera que l'utilisation simultanée des deux chaînes permet d'avoir des meilleures performances que les bornes que l'on peut obtenir en combinant les courbes des figures 8.8 et 8.13 (e.g. sensibilité de 90% pour 1,7 faux positifs par volume ou 80% pour 1,45 faux positifs par volume). Cela peut s'expliquer par le fait que même si les deux canaux ne sont pas destinés à la détection des mêmes structures, la séparation entre les deux populations cibles n'est pas facilement descriptible. Ainsi, il existe une proportion de signes radiologiques que les deux méthodes sont susceptibles de détecter. C'est pourquoi les sensibilités du tableau 8.2 sont meilleures qu'une simple combinaison de celles qu'on peut trouver à spécificité équivalente pour les deux canaux. La figure 8.15 illustre cela avec une lésion très fortement spiculée utilisée pour la validation du canal de détection de convergence. Bien que cette dernière soit considérée comme non suspecte par l'étape de réduction de faux positifs de ce canal, le processus de détection d'opacités permet rattraper cette erreur lors de l'étape de fusion des deux canaux.



FIG. 8.15 – Coupe d'un volume de tomosynthèse contenant une lésion fortement spiculée ayant servi à l'évaluation du canal de détection de convergences. Cette lésion n'est pas détectée par ce canal mais est détectée par le canal de détection d'opacités.

Les résultats présentés sont comparables avec ce que l'on peut trouver dans la littérature. Parmi les premiers travaux citant des performances pour l'application considérée, on peut citer ceux de Chan *et al.* (2005) qui proposent une méthode de détection en 3D où une étude préliminaire sur 26 volumes contenant 13 opacités malignes, 2 distorsions architecturales malignes, 1 bénigne et 10 opacités bénignes, a permis d'obtenir une sensibilité de 80% (respectivement 85%) pour 2 (respectivement 2,2) faux positifs par cas. Ces premiers résultats sont difficilement comparables à ce que nous présentons dans la mesure où la population cible n'est pas la même : dans ce que nous avons présenté, nous ne cherchons à détecter que les lésions malignes. Peu après, Chan *et al.* (2007) ont proposé de combiner l'approche précédente à une détection dans les images de projections. Sur une base de données contenant 41 opacités malignes et 11 bénignes, ils

ont montré qu'en rajoutant l'information 2D, sur leurs données ils pouvaient, pour une sensibilité de 80%, passer d'une spécificité de 1,6 faux positif par sein à une spécificité de 1,19 ou pour une sensibilité de 90% passer d'une spécificité de 3,04 à une spécificité de 2,27. Plus récemment encore, Chan *et al.* (2008a) ont étudié plus en détails les différences entre détection en 3D et en 2D et par combinaison des deux méthodes. Ces travaux sont les plus pertinents pour une comparaison avec les performances que nous obtenons car la taille de la base de données utilisée est du même ordre de grandeur (100 volumes contenant 69 lésions malignes contre 53 lésions malignes pour notre validation). De plus les résultats proposés dans ces travaux concernent la détection de lésions malignes ou bénignes ainsi que la détection de lésions malignes seules. Les points de fonctionnement de leurs trois systèmes de détection pour les lésions malignes sont présentés à la table 8.3.

Sensibilité (%)	Spécificité (faux positifs par sein)		
	Approche 2D	Approche 3D	2D + 3D
80	2,43	1,46	0,84
85	3,16	2,06	1,06
90	3,65	2,52	1,61

TAB. 8.3 – Performances des chaînes de détection proposées par Chan *et al.* (2008a).

On remarquera que même si les points de fonctionnement ne sont pas exactement les mêmes entre les tableaux 8.2 et 8.3 on peut faire une brève comparaison entre les spécificités obtenues pour des sensibilités de 80 et 90%. On remarquera que pour une sensibilité de 80% l'approche combinée proposée par Chan *et al.* (2008a) est meilleure en termes de spécificité que notre chaîne globale, les deux autres approches étant un peu moins bonnes. Cependant, pour une sensibilité de 90%, notre chaîne donne des résultats comparables à l'approche combinée. On remarquera aussi que notre approche permet de manière raisonnable d'augmenter la sensibilité au delà des 90%. Au delà de cette comparaison assez simpliste, on peut prendre un peu de recul et faire plusieurs remarques. La première est le fait que les bases de données dans les deux cas sont relativement petites en regard de la variabilité des opacités en général. Deuxièmement, les deux bases ne sont pas les mêmes, rendant difficile toute comparaison. Enfin, on remarquera que, dans leurs travaux, Chan *et al.* (2008a) ne semblent se focaliser que sur les opacités et délaissent les distorsions architecturales. Ainsi on pourrait comparer les résultats du tableau 8.3 aux résultats du premier canal présentés à la figure 8.8. En faisant cela on s'aperçoit que la différence de spécificité se réduit pour une sensibilité de 80% pour atteindre 1,04 faux positif par volume. D'autre part la spécificité du canal de détection de densité est meilleure pour une sensibilité de 90% avec 1,23 faux positifs par volume. Pour ces raisons, même si encore une fois toute comparaison reste difficile, on conclura que l'approche proposée est comparable aux algorithmes de référence.

On peut aussi citer les travaux de Singh (2008) qui proposent une approche autre pour la détection d'opacités en tomosynthèse du sein reposant sur la théorie de l'information. Pour une base de données contenant 28 lésions dont 10 malignes, une spécificité de 2,4 faux positifs par sein est obtenue pour une sensibilité de 90% (Singh *et al.*, 2008a).

Enfin, on peut citer une dernière référence qui expose des résultats. Chan *et al.* (2008b) étudient ainsi l'influence du nombre de projections sur les performances de leur CAD travaillant uniquement en 3D (Chan *et al.*, 2005). Ainsi pour sur une base de données contenant 14 lésions malignes pour une sensibilité de 80% les taux de faux positifs pour 11 et 21 projections sont 0,87 et 1,11, respectivement.

8.5 Conclusion

Dans ce chapitre nous avons proposé un schéma à deux canaux pour la détection d'opacités et de distorsions architecturales souvent signes de la présence de cancer. La procédure proposée permet d'apporter un exemple concret pour la mise en pratique des développements théoriques introduits dans les chapitres précédents.

Nous avons montré que notre approche, bien que n'étant pas encore parfaite, est comparable à ce qui existe dans la littérature. Cela permet de considérer l'approche proposée comme une solution viable qui

peut être le point de départ d'un système de détection voué à évoluer et à être amélioré.

En perspective d'évolution, on pourrait envisager une modélisation de l'imprécision induite par les artefacts de la tomosynthèse du sein dans le cadre des images d'ombres floues utilisées dans le premier canal. En effet, le formalisme employé laisse la porte ouverte à ce type d'amélioration sous la contrainte unique d'allonger le temps de traitement. On pourrait aussi envisager de modéliser l'imprécision ou l'incertitude dans l'étape de segmentation/classification en s'inspirant des travaux présentés au chapitre 6. L'apport d'une telle technique par rapport à la chaîne présentée serait intéressant à étudier.

En ce qui concerne le canal de détection de convergences qui, comme on l'a vu au chapitre 7, est plus une méthode développée pour la détection de lésions en tomosynthèse, contrairement au formalisme générique des filtres connexes flous, il serait intéressant d'évaluer l'apport en termes de performances d'une modélisation plus réaliste du contenu des coupes d'un volume. Ainsi la prise en compte de la corrélation entre voxels serait un bon point de départ.

Enfin, comme piste future d'investigation, on peut citer le travail qui reste à faire sur la constitution d'une plus grande base de données pour la validation. En effet, même si l'on peut se comparer à la littérature en ce qui concerne la taille de la base utilisée pour l'évaluation de notre approche, la quantité de données à disposition est assez faible compte tenu la variabilité des opacités. Ce dernier point est tout particulièrement vrai pour le canal de détection de convergences. En effet l'évaluation conduite sur treize cas peut juste nous permettre d'avoir une idée sur la tendance de la validité de l'approche. De manière similaire, la faible quantité de données nous a aussi conduit à écarter quelques lésions de notre validation. Ainsi en augmentant le nombre de lésions, on pourrait définir par exemple un troisième canal pour les prendre en compte.

Chapitre 9

Conclusions et perspectives

Filtrage flou d'images

Dans ce manuscrit, nous avons proposé un ensemble d'outils génériques de traitement d'image utilisant la théorie des ensembles flous. Nous avons proposé une extension dans ce formalisme des filtres connexes dédiés aux images de niveaux de gris. Pour cela nous avons développé une chaîne générique permettant de traiter complètement une image à niveaux de gris. Ainsi nous avons abordé le problème de la représentation des défauts dans une image à niveaux de gris, la définition de filtres dédiés à ces dernières ainsi que leur mise en œuvre de manière efficace.

Nous avons proposé une manière originale de représenter l'imprécision dans les images à niveaux de gris en introduisant le concept d'images floues. Ces dernières sont des images où les valeurs d'intensité ne sont plus représentées par des nombres classiques, mais par des sous-ensembles flous. Nous avons étudié deux grandes classes d'images floues : les images d'intervalles flous et les images d'ombres floues. La première famille est l'extension naturelle des images classiques alors que la deuxième est l'extension des images d'ombres parfois utilisées en morphologie mathématique. Ces deux représentations disposent de propriétés différentes qui sont souhaitables ou non en fonction du filtre que l'on veut leur appliquer. Par exemple, les images d'ombres floues permettent d'exprimer des extensions de filtres de nivellement, alors que les images d'intervalles flous sont plutôt utilisées pour du filtrage par zone plates.

La construction de ces images a été abordée. Ainsi nous avons étudié comment convertir une image classique en image floue. Si cela ne pose pas de problème dans le cas d'imprécision, ce point est assez délicat dans le cas de dégradations d'ordre statistique qui ne sont pas directement représentables par une image floue. Un résultat intéressant résultant de l'étude de ce problème est la possibilité de transformer ces défauts de nature statique contenus dans une image par l'introduction d'imprécision de débruitage. Ainsi le formalisme proposé permet de manipuler directement ou indirectement différentes sources de défauts, ce qu'une approche statistique par exemple ne pourrait pas forcément faire. Parmi les approches proposées reposant sur le concept d'imprécision de débruitage, une adaptation du filtrage en ondelettes a été proposée. Ces considérations sur les techniques de construction d'images floues ont été étudiées sur des images synthétiques et réelles dans le cadre de la détectabilité de structures contenues dans ces dernières.

L'utilisation des précédentes images floues a été proposée par l'introduction de nouveaux types de filtres qui ont pour but de généraliser la famille des filtres connexes. Dans cette optique de généralisation nous avons proposé un formalisme générique permettant d'exprimer les généralisations floues des filtres classiques comme par exemple les filtres de nivellement, ou d'amincissement. Ce formalisme permet d'introduire différentes familles de filtres ayant différentes propriétés. Ces travaux théoriques ont donc l'intérêt d'étudier de manière abstraite les filtres proposés dans ce manuscrit ou d'autres qui pourraient être développés dans des applications futures.

Pour finir notre chaîne générique de traitement, nous avons développé un algorithme permettant la mise en œuvre des filtres précédemment discutés. Notre approche repose sur la représentation des ensembles flous contenus dans les images floues à l'aide d'une structure arborescente. Grâce à cette dernière, nous avons introduit des mécanismes permettant de tirer parti des redondances entre les ensembles flous extraits

pour deux niveaux de gris successifs. La conception de ces derniers a été validée de manière théorique par une étude approfondie des algorithmes résultants.

Parmi les perspectives concernant cette chaîne de traitement, on pourrait envisager de la valider sur d'autres types d'images que celles proposées. On pourrait aussi étudier les optimisations possibles de certains filtres précis : l'écriture générique de l'algorithme rapide de filtrage peut avoir un surcoût dans le calcul de certaines mesures. Une stratégie de mises à jour des attributs en parallèle des mises à jour d'arbres est une piste à explorer. L'introduction d'autres sources d'imprécision lors de la construction des images floues pourrait aussi être étudiée. Néanmoins, cela peut être dépendant du type d'image traitée : on pourrait imaginer par exemple en tomosynthèse une estimation des artefacts de reconstruction qui serait introduite dans les images floues.

Segmentation

Le problème de segmentation de signes radiologiques en tomosynthèse a aussi été abordé. La problématique qui apparaît avec ce type de structures provient essentiellement de l'ambiguïté provenant des images. En effet, certains signes radiologiques sont difficilement discernables et d'autres sont difficilement délimitables. Ces deux points sont respectivement liés à l'incertitude et à l'imprécision provenant des données à traiter. En s'appuyant sur des travaux existants (Peters, 2007) nous avons choisi d'utiliser un formalisme par contours flous qui est bien adapté à la représentation de ces deux notions. Dans ce contexte, nous avons étudié trois méthodes de segmentation floues. La première est une adaptation de l'approche par seuillages multiples utilisée par Peters (2007). Les limitations en tomosynthèse d'une telle approche ainsi que le lien fort qui existe avec le formalisme de filtrage d'images d'ombres floues a été étudié. Ce dernier point met en évidence l'équivalence entre une extraction de contours flous par seuillages multiples et le filtrage d'une image d'ombres non floue par un filtre connexe flou reposant sur les mêmes attributs discriminants que ceux utilisés pour le calcul des degrés d'appartenance dans le contour flou. La seconde approche étudiée est une extension d'une méthode de segmentation classique exprimée à l'aide du formalisme des ensembles de niveaux (Peters *et al.*, 2007). Cette approche a été utilisée en pratique dans une chaîne d'aide à la décision visant à montrer la pertinence de l'utilisation des contours flous en tomosynthèse. Enfin nous avons proposé une troisième méthode de segmentation floue qui repose sur la modélisation d'un contour par un chemin dans un graphe de coût lié à l'image à segmenter (Timp et Karssemeijer, 2004). Dans l'extension proposée nous avons présenté des techniques permettant l'extraction de plusieurs contours à partir d'une image. Nous avons étudié la sémantique portée par les contours extraits par ces différentes techniques, à savoir l'incertitude et l'imprécision. Nous avons enfin proposé une méthode originale pour calculer des degrés d'appartenance pour les différents contours extraits reposant sur une notion de carte de densité de contours.

Cette étude de méthodes de segmentation floue constitue en fait un travail préliminaire pour l'utilisation robuste de telles techniques pour la détection et la caractérisation de signes radiologiques suspects en tomosynthèse du sein. Ainsi parmi les prochaines étapes, on peut citer l'étude nécessaire de techniques de classification adaptées à des données floues. Des travaux préliminaires ont déjà été réalisés (Bothorel, 1996) et ont été présentés dans ce manuscrit pour illustrer l'intérêt des contours flous dans une application pratique de caractérisation de signes radiologiques. Cependant ceux-ci restent marginaux en comparaison des techniques de classification conçues pour des données non floues. On pourrait ainsi envisager de considérer des techniques standard de classification couramment utilisées pour les étendre afin qu'elles puissent gérer des données floues. Pour faire cela, il faudrait s'assurer de la validité de l'extension proposée pour le traitement de ces données, aussi bien d'un point de vue théorique (sens des traitements effectués) que d'un point de vue expérimental (apport de méthode floue par rapport la méthode originale).

Détection de convergence

Un signe caractéristique des cancers en mammographie par tomosynthèse, est l'apparition de structures en étoile. Dans ce contexte, une autre contribution proposée dans ce manuscrit fut l'élaboration d'une méthode de détection de convergences dans des images. L'approche proposée combine une variation de

l'approche de Karssemeijer et te Brake (1996) avec la détection de structures par modélisation a contrario proposée par Desolneux *et al.* (2000).

Nous avons étudié différentes variations possibles de la mesure de convergence originellement proposée par Karssemeijer et te Brake (1996). En définissant un modèle statistique du contenu des coupes d'un volume de tomosynthèse du sein, nous avons déduit les valeurs peu probables d'occurrence de notre mesure de convergence. Ces travaux théoriques nous ont permis de définir une approche simple et facile à paramétrer pour détecter les motifs stellaires.

Dans le même domaine, une deuxième contribution a été de développer un algorithme efficace pour la mise en œuvre de la précédente approche. Pour cela nous avons reformulé le problème sous forme réursive. Cela a permis de passer d'une complexité de l'ordre de $O(NR^3)$ pour la mise en œuvre naïve à une complexité de $O(NR^2)$, avec N le nombre de pixels à traiter et R le rayon maximal d'une zone de convergence que l'on s'autorise à détecter.

Enfin une troisième contribution a été l'introduction d'une mesure pour différencier les convergences malignes des convergences dues à des croisements de structures normales telles que des fibres. Cette dernière repose sur l'analyse des orientations des structures linéaires présentes dans la zone de détection.

Parmi les améliorations possibles de ces travaux, nous pouvons citer l'amélioration de la modélisation du contenu d'un sein sans pathologie qui constitue le modèle naïf. En effet, dans notre modélisation, nous considérons que les voxels sont indépendants. On pourrait envisager de considérer la texture présente dans le sein, ainsi que la corrélation induite par les algorithmes de reconstruction. Cela permettrait d'étudier leur impact sur le modèle naïf et sur les réalisations de la mesure de convergence. Cela pourrait se faire de manière théorique ou encore expérimentalement en estimant les distributions de probabilité mises en œuvre dans la méthode sur des données réelles.

Chaîne de détection de cancer

Les outils précédemment introduits ont non seulement été étudiés de manière théorique, mais aussi utilisés en pratique. Ainsi avons nous proposé une chaîne de détection automatique de signes radiologiques suspects pour des volumes de tomosynthèse du sein. Nous nous sommes focalisés sur la détection de deux types de signes radiologiques, à savoir les opacités et les distorsions architecturales. L'approche que nous avons proposée se décompose en deux canaux qui travaillent en parallèle sur le volume à traiter pour détecter les deux types de signes radiologiques qui nous intéressent

Le canal de détection d'opacités se décompose en une étape de marquage, une étape de segmentation des zones suspectes et une étape de prise de décision au travers d'une classification utilisant des attributs extraits à partir des contours obtenus. L'étape de marquage est réalisée à l'aide d'un filtre connexe flou exprimé dans le formalisme introduit dans ce manuscrit. Ce filtre est appliqué sur les différentes coupes composant le volume pour limiter l'impact des artéfacts de reconstruction, l'information 3D étant prise en compte après filtrage par des techniques de regroupement. Il utilise des critères simples pour détecter les composantes connexes qui sont contenues dans le volume et qui peuvent être suspectes. Ainsi la discrimination se fait sur une contrainte dérivée d'une mesure de contraste couplée à des critères de taille et de compacité. Le formalisme flou permet d'être assez robuste en ce qui concerne les lésions qui sont à la limite de ces critères de détection. Pour la segmentation, des approches non floues ont été utilisées. Deux méthodes de segmentation ont été employées. La première repose sur des seuillages multiples de l'image et la seconde sur la modélisation du contour par un chemin de coût minimal dans l'image. Les performances de ce canal ont été évaluées sur une base de données contenant 40 cancers prouvés par biopsie et 53 seins sans pathologie. La méthode de segmentation a un impact sur les performances obtenues, ainsi la modélisation par chemins minimaux offre de meilleures performances. Pour cette dernière, des sensibilités de 90% et 80% ont été obtenues pour respectivement 1,23 et 1,04 faux positifs par sein.

Le canal de détection de distorsions architecturales a été mis en œuvre grâce aux travaux proposés sur la détection de convergences suspectes. Ce dernier est aussi décomposé en plusieurs étapes, à savoir détection de zones suspectes, caractérisation des zones puis classification. L'étape de détection met en œuvre la modélisation a contrario utilisant comme critère discriminant une forte convergence de certaines structures. Les autres contributions, comme l'analyse plus fine des orientations, sont quant à elle utilisées dans l'étape de caractérisation des zones suspectes. Enfin la classification se fait, comme pour le premier

canal, à l'aide de techniques standard reposant sur les SVM. Ce canal a été validé à l'aide de 65 volumes dont 13 contenaient une lésion maligne prouvée par biopsie. Ainsi des sensibilités de 92% et 76% ont été obtenues pour respectivement 0,48 et 0,35 faux positifs par sein.

Une méthode simple d'agrégation de type disjonctive a aussi été proposée pour la fusion des résultats provenant des deux canaux. L'idée est de tirer parti de la superposition des populations ciblées par les deux canaux. En pratique cela se traduit par une augmentation de la sensibilité pour une spécificité donnée par rapport à une combinaison directe des performances des deux canaux. La chaîne ainsi obtenue est comparable aux travaux de référence que l'on peut trouver dans la littérature. Ajoutons enfin que notre chaîne a un fort potentiel d'amélioration puisque l'utilisation de méthodes de segmentation floue et de classification floue reste à être mis en œuvre.

La chaîne de détection proposée constitue une base pour la détection de certains signes radiologiques. On peut envisager plusieurs pistes pour la faire évoluer. Ainsi on pourrait modéliser la fonction d'étalement de point associée à la reconstruction pour déduire une approche tirant directement parti de l'information 3D lors de la détection. Une comparaison de l'apport d'une telle approche par rapport à notre modélisation serait intéressante. Une seconde piste d'investigation pourrait être l'étude de l'apport du formalisme de segmentation floue s'il était intégré dans cette chaîne. Enfin, un des points les plus importants dont dépend toute évaluation est l'agrandissement de la base de données. En effet, de par la variabilité des structures recherchées un nombre conséquent de données devrait être utilisé pour pouvoir être précis quant aux performances réelles de notre approche et de ses potentielles évolutions. Cette étape permettrait aussi de définir une méthode de détection pour les structures cancéreuses non prises en compte dans nos travaux pour cause de non représentativité.

Annexe A

Quelques notions sur les filtres en traitement d'images

A.1 Image d'ombres

Une image d'ombres est un concept originellement introduit pour étendre les outils de morphologie mathématique aux images à niveaux de gris. En effet, ces outils ont été historiquement définis pour des ensembles. L'idée est de transformer une image à niveaux de gris en un ensemble binaire de dimension supérieure pour pouvoir utiliser les traitements classiques (érosion, dilatation, etc.).

En pratique l'ensemble binaire est obtenu à partir de l'image en considérant que tous les éléments constitués d'un pixel p et d'un niveau de gris g (appartenant donc à $\Omega \times \mathcal{G}$, avec Ω le domaine de définition de l'image et $\mathcal{G}$ l'ensemble des niveaux de gris) est inclus dans l'ensemble si l'image originale a un niveau de gris supérieur ou égal à g au point p . Plus formellement cela peut s'écrire comme :

$$\{(p, g) \in \Omega \times \mathcal{G} / I(p) \geq g\}$$

avec I l'image à niveaux de gris.

Plus concrètement, l'ensemble obtenu correspond à l'ombre sous le signal de l'image considéré. La figure A.1 illustre ce concept sur une image 1D.

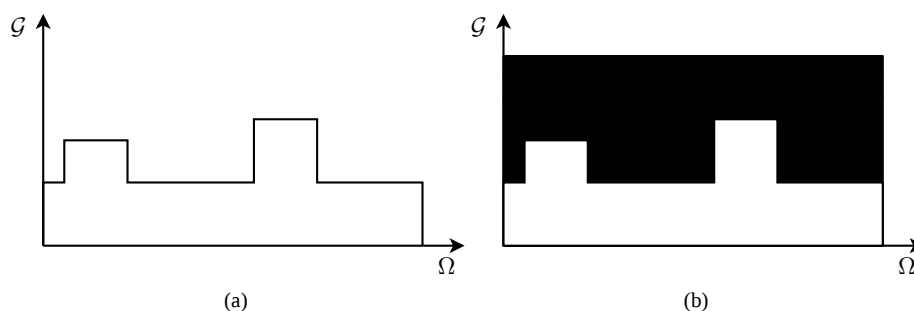


FIG. A.1 – Exemple de construction d'une image d'ombres. (a) Image 1D originale. (b) Image d'ombres obtenue (les éléments à l'intérieur de l'ensemble apparaissent en blanc).

A.2 Connexité

Nous allons présenter ici quelques notions sur la connexité dans les images discrètes. Pour une description plus formelle et plus complète, on pourra se référer aux ouvrages de Chassery et Montanvert (1991) et Klette et Rosenfeld (2004).

Dans la plupart des problèmes de traitement d'image, les données à traiter sont discrètes. Ainsi on définit souvent un pavage composé de cellules auxquelles on associe un niveau de gris. On peut également associer un point central à chacune de ces cellules pour définir un maillage et ainsi raisonner de manière topologique sur ces cellules. Si on travaille sur des ensembles définis sur des images discrètes, on peut se retrouver face à un problème de définition lorsque l'on cherche à savoir si deux éléments d'un ensemble sont reliés entre eux (connexes, voisins). On peut de manière arbitraire fixer une convention, ainsi par exemple on peut, en 2D, considérer qu'un pixel est seulement connexe à ses voisins du dessus, du dessous, de droite et de gauche. Cette convention est connue sous le nom de 4-connexité. Si on décide d'ajouter les voisins accessibles en diagonale, on parle alors de 8-connexité. Ces deux conventions sont illustrées à la figure A.2.

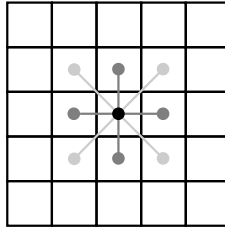


FIG. A.2 – Exemple de 4-connexité (gris foncé) et 8-connexité (gris clair + gris foncé).

En utilisant cette notion de connexité, on peut définir un ensemble connexe comme étant un ensemble tel que pour tout couple de points de l'ensemble, il existe un chemin (suite de points connexes selon la connexité choisie) qui les relie et qui soit inclus dans l'ensemble.

En faisant cela, on s'aperçoit que le formalisme n'est pas bien posé, en tout cas pour les grilles de type rectangulaire. En effet, on peut remarquer que considérer une même connexité pour le fond et l'objet, peut mener à des paradoxes d'ordre topologique comme illustré à la figure A.3. Pour pallier ces problèmes, on utilise généralement des connexités différentes pour le fond et l'objet. Ainsi en 2D, on peut utiliser une 4-connexité pour l'un et une 8-connexité pour l'autre, et inversement. Formellement, c'est l'équivalent discret du théorème de Jordan qui permet de montrer ce résultat.

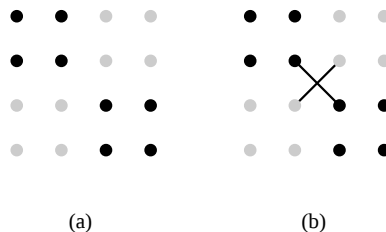


FIG. A.3 – Problème de connexité entre fond et objet. (a) En 4-connexité, il peut exister des zones qui ne font ni partie de l'objet ni de son complémentaire. (b) En 8-connexité, deux composantes connexes peuvent se croiser.

A.3 Différence filtre/opérateur

En morphologie mathématique, on définit des outils de traitement d'ensembles ou de fonctions. Ces derniers ont été étudiés de manière théorique, notamment au niveau des propriétés qu'ils peuvent avoir. On parlera ainsi d'opérateurs pour définir les outils les plus simples, et en fonction des propriétés que ces derniers peuvent avoir, on emploiera des terminologies plus spécialisées. Parmi les propriétés évoquées dans ce manuscrit qui comptent parmi les plus courantes, on citera, pour un opérateur ϕ donné :

- la croissance : pour tout couple d'ensembles X et Y $X \subseteq Y \Rightarrow \phi(X) \subseteq \phi(Y)$,
- l'idempotence : pour tout ensemble X $\phi(\phi(X)) = \phi(X)$,
- et l'anti-extensivité : pour tout ensemble X $\phi(X) \subseteq X$.

D'un point de vue de la terminologie, un opérateur cumulant les deux premières propriétés sera communément appelé un filtre morphologique. De même si on ajoute la troisième propriété, on parlera d'ouverture algébrique.

A.4 Rappels sur les filtres connexes

Les filtres connexes sont au cœur des développements proposés dans les chapitres 2, 3 et 4. Nous allons présenter quelles sont les différences entre les principales familles de filtres évoquées.

A.4.1 Ouverture d'attribut, amincissements

Parmi les filtres connexes pour images à niveaux de gris les plus utilisés, on trouve les ouvertures d'attributs ainsi que les filtres d'amincissement. Ces deux familles de filtres sont assez semblables dans leur expression que nous rappelons dans le cas du filtrage d'une image I pour comprendre les différences entre ces dernières :

$$\forall p \in \Omega \quad \delta(I)(p) = \max \left\{ t/C \left(\hat{f}_{X_t^+(I)}^p \right) = 1 \right\}$$

avec $X_t^+(I)$ l'opérateur de seuillage qui retourne l'ensemble des points dont l'intensité est supérieure ou égale à t et C un critère à valeur dans $\{0, 1\}$ que les composantes connexes doivent respecter.

Cette expression peut s'interpréter comme l'extraction de toutes les composantes connexes contenues dans tous les ensembles issus de tous les seuillages possibles de l'image. Certaines de ces dernières sont ensuite supprimées et d'autres conservées. Le niveau de gris de l'image filtrée est ensuite déduit en prenant le niveau maximal associé aux composantes connexes conservées pour un point donné.

En pratique les ouvertures d'attribut se comportent comme des filtres qui suppriment complètement ou partiellement les maxima de l'image. Cette propriété est en fait induite par le fait que le critère C doit être croissant pour ce type de filtre. Ainsi, si on considère un maximum de l'image, et que l'on s'intéresse aux différentes composantes connexes constituant ce maximum et issues des différents seuillages de l'image, on s'aperçoit que le fait qu'elles soient emboîtées assure que quand la composante connexe atteindra une certaine taille (pour un niveau de gris donné), le critère C sera vérifié pour toutes les composantes issues de niveaux de gris inférieurs. Un exemple d'un tel filtre est proposé à la figure A.4(a).

Dans le cas des filtres d'amincissement, la contrainte de croissance sur l'opérateur C est relâchée. Ainsi, on n'est plus assuré de ne changer que les maxima de l'image. Cela a pour impact de potentiellement créer des bords fictifs dans l'image filtrée. Un exemple d'un tel filtre est proposé à la figure A.4(b).

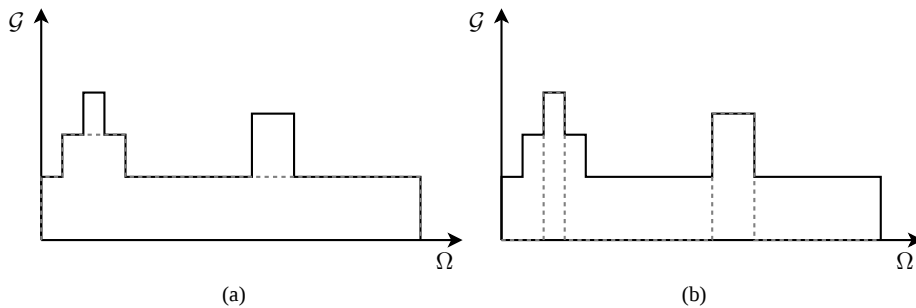


FIG. A.4 – Exemple d'une ouverture d'aire avec $C(X) = (\text{aire}(X) \geq 3)$ (a) et d'un filtre d'amincissement avec $C(X) = (\text{aire}(X) \leq 2)$ (b). Les données non filtrées (resp. filtrées) apparaissent en trait plein (resp. pointillé).

A.4.2 Opérateur d'extinction

Un opérateur d'extinction se définit à l'aire d'une famille décroissante $\Psi = \{\psi_\lambda, \lambda \in \mathbb{R}^+\}$ d'opérateurs connexes anti-extensifs (Vachier, 1995) :

- pour tout $\lambda \geq 0$ ψ_λ est connexe,
- pour tout ensemble X , pour tout $\mu \geq \lambda \geq 0$ $\psi_\mu(X) \subseteq \psi_\lambda(X)$,
- pour tout ensemble X , pour tout $\lambda \geq 0$ $\psi_\lambda(X) \subseteq X$ et $\psi_0 = id$

Grâce à une telle famille d'opérateurs, on peut définir les valeurs maximales de λ pour lesquelles les composantes connexes incluses dans un ensemble donné disparaissent. Ces valeurs sont connues sous le nom de valeur d'extinction. Les quantités floues représentant la persistance d'un ensemble introduites au chapitre 2 sont en fait l'extension de ces valeurs d'extinction.

Annexe B

Quelques notions sur les ensembles flous

B.1 Opérateurs d'agrégation

Dans ce manuscrit, les notions de t-norme et t-conorme sont souvent utilisées. On peut voir ces deux opérateurs d'agrégation comme des équivalents flous du *et* et du *ou* binaire.

Définition B.1.1. Une t-norme $\top$ est une fonction de $[0, 1] \times [0, 1]$ dans $[0, 1]$ telle que :

- $\top$ est commutative : $\forall (x, y) \in [0, 1]^2 \quad x \top y = y \top x$,
- $\top$ est associative : $\forall (x, y, z) \in [0, 1]^3 \quad (x \top y) \top z = x \top (y \top z)$,
- 1 est élément neutre : $\forall x \in [0, 1] \quad x \top 1 = 1 \top x = x$,
- $\top$ est croissante par rapport aux deux variables : $\forall (x, x', y, y') \in [0, 1]^4 \quad ((x \leq x') \wedge (y \leq y')) \Rightarrow (x \top y) \leq (x' \top y')$

Définition B.1.2. Une t-conorme $\perp$ est une fonction de $[0, 1] \times [0, 1]$ dans $[0, 1]$ telle que :

- $\perp$ est commutative : $\forall (x, y) \in [0, 1]^2 \quad x \perp y = y \perp x$,
- $\perp$ est associative : $\forall (x, y, z) \in [0, 1]^3 \quad (x \perp y) \perp z = x \perp (y \perp z)$,
- 0 est élément neutre : $\forall x \in [0, 1] \quad x \perp 0 = 0 \perp x = x$,
- $\perp$ est croissante par rapport aux deux variables : $\forall (x, x', y, y') \in [0, 1]^4 \quad ((x \leq x') \wedge (y \leq y')) \Rightarrow (x \perp y) \leq (x' \perp y')$

En pratique les t-normes et t-conormes peuvent prendre plusieurs formes. Le tableau B.1 présente les couples les plus courants.

Description	t-norme ($x \top y$)	t-conorme ($x \perp y$)
Minimum/maximum	$\min(x, y)$	$\max(x, y)$
Probabiliste	xy	$x + y - xy$
Lukasiewicz	$\max(0, x + y - 1)$	$\min(1, x + y)$

FIG. B.1 – Exemples de t-normes/t-conormes.

B.2 Quantités floues

Dans cette section nous considérerons des ensembles flous définis sur un ensemble X muni d'une relation d'ordre totale (typiquement $\mathbb{R}$).

Définition B.2.1. Une quantité floue est un ensemble flou défini sur X

Définition B.2.2. Un intervalle flou est une quantité floue convexe, c'est-à-dire que toutes ses α -coupes sont des intervalles.

Définition B.2.3. Un nombre flou est un intervalle flou dont les α -coupes sont des intervalles fermés et dont le noyau se réduit à un unique élément.

Annexe C

Preuves sur les mises à jour d'arbres

C.1 Différents cas de l'opérateur $\overline{M}$

Nous présentons ici la justification des différents cas évoqués pour la construction de l'opérateur $\overline{M}$ de la définition 3.2.1. Ces cas sont en fait déduits des différentes inégalités possibles entre les degrés associés aux racines des arbres à fusionner. Ils sont aussi séparés en fonction de l'inclusion ou non des marqueurs dans un des sous-arbres issus des racines. Le tableau C.1 montre que toutes les configurations possibles sont bien prises en compte par les différents cas décrits dans la définition de l'opérateur.

$\mathcal{A}(\mathcal{R}(t_1)) - \mathcal{A}(\mathcal{R}(t_2))$	$ Q_{p_1}(sa(t_1, \mathcal{R}(t_1))) $	$ Q_{p_2}(sa(t_2, \mathcal{R}(t_2))) $	Numéro de cas
> 0	0	0	cas 5
> 0	0	1	cas 5
> 0	1	0	cas 5
> 0	1	1	cas 5
< 0	0	0	cas 4
< 0	0	1	cas 4
< 0	1	0	cas 3
< 0	1	1	cas 3
$= 0$	1	1	cas 2
$= 0$	0	0	cas 1
$= 0$	0	1	cas 1
$= 0$	1	0	cas 1

TAB. C.1 – Énumération de toutes les configurations possibles pour les entrées de l'opérateur $\overline{M}$ associés avec les numéros de cas de la définition 3.2.1.

C.2 Preuve du théorème 3.2.1

Dans le but de prouver le théorème 3.2.1, nous considérons toutes les configurations possibles de manière récursive comme illustré à la figure C.1. L'idée est de montrer que le théorème 3.2.1 est vrai pour les cas terminaux (c'est-à-dire les cas qui ne nécessitent pas d'appels supplémentaires à $\overline{M}$), ainsi que pour les autres configurations sous la seule hypothèse que leurs appels récursifs à $\overline{M}$ vérifient le théorème.

Démonstration. Prouvons ce théorème récursivement :

Cas terminaux : Soit $(t_1, t_2) \in \mathcal{T}^{e^2} / \forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset$ et $(p_1, p_2) \in \Omega$ et $t_3 = \overline{M}(t_1, t_2, p_1, p_2)$. Deux cas terminaux existent (pas d'appel récursif à $\overline{M}$) : cas 1 et 4 :

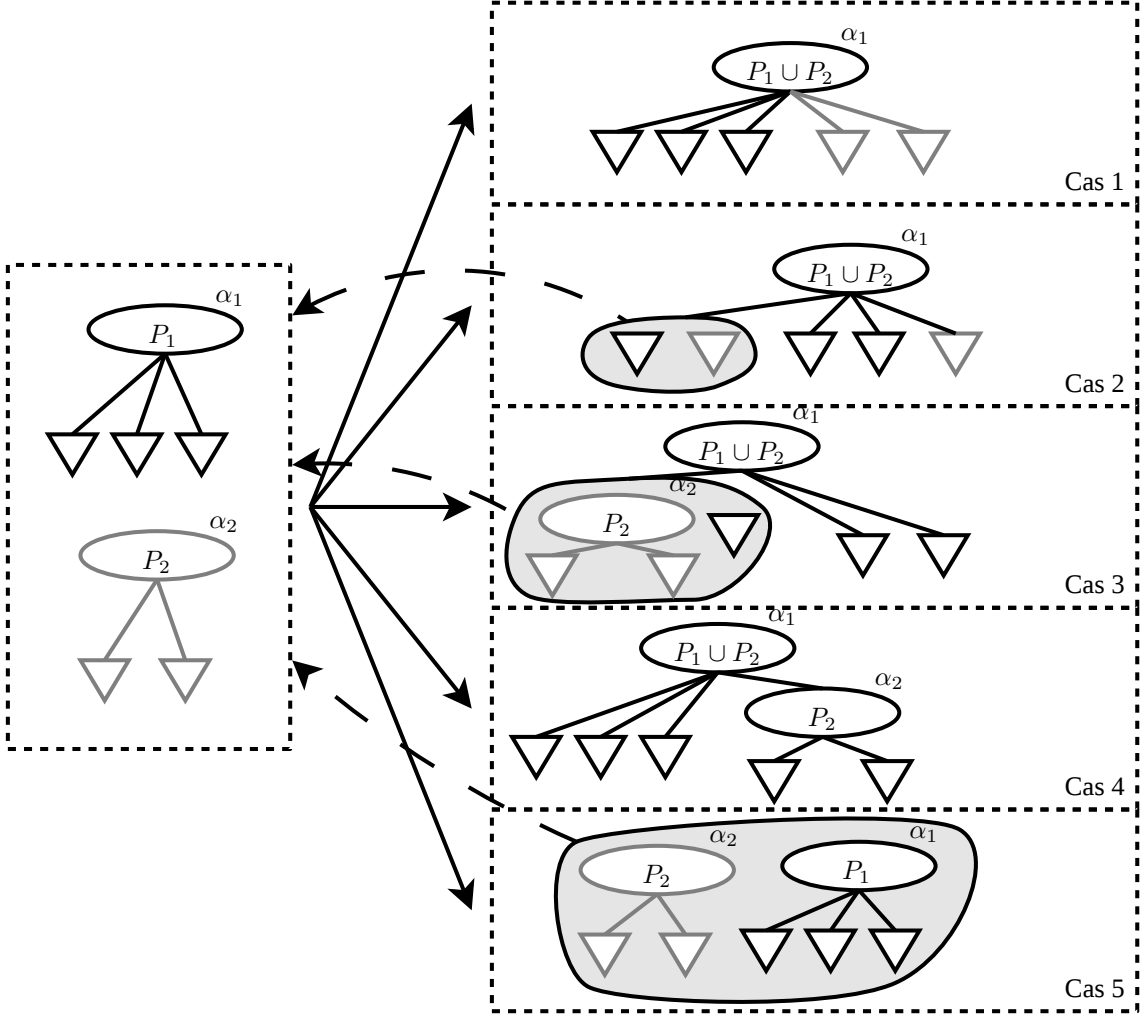


FIG. C.1 – Preuve par récurrence du théorème 3.2.1. Les cas 1 et 4 sont terminaux (la preuve ne fait aucune hypothèse à propos des deux sous-arbres d'entrée) et ils vérifient le théorème. Les cas 2, 3, 5 sont vrais si le théorème est vrai pour la fusion des sous-arbres dans la zone grise. On prouve de manière récursive cela (flèches en pointillés) en utilisant toutes les nouvelles configurations possibles (les 5 cas) jusqu'à ce qu'on arrive sur un cas terminal.

cas 1 $(\mathcal{A}(\mathcal{R}(t_1)) = \mathcal{A}(\mathcal{R}(t_2))) \wedge ((|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 0) \vee (|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))| = 0))$

Dans ce cas, on a $\forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(n) \quad \mathcal{A}(s) > \mathcal{A}(n)$ (équation 3.10) parce que :

- $\forall t \in sa(t_1, \mathcal{R}(t_1)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$ parce que $t_1 \in \mathcal{T}^e$
- $\forall t \in sa(t_2, \mathcal{R}(t_2)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$ parce que $t_2 \in \mathcal{T}^e$
- $\forall n \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2)) \quad \mathcal{A}(\mathcal{R}(t_3)) < \mathcal{A}(n)$ parce que $\mathcal{A}(\mathcal{R}(t_3)) = \mathcal{A}(\mathcal{R}(t_1)) = \mathcal{A}(\mathcal{R}(t_2))$

D'autre part, on a aussi $\forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(t_3, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$ (équation 3.11) parce que :

- $\forall t \in sa(t_1, \mathcal{R}(t_1)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$ parce que $t_1 \in \mathcal{T}^e$
- $\forall t \in sa(t_2, \mathcal{R}(t_2)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$ parce que $t_2 \in \mathcal{T}^e$
- $\forall n \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2)) \quad \mathcal{P}(n) \subseteq \mathcal{P}(\mathcal{R}(t_3))$ parce que $\mathcal{P}(\mathcal{R}(t_3)) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2))$

Finalement, on a $\forall n \in \mathcal{N}(t_3) \quad \bigcap_{s \in \mathcal{W}(t_3, n)} \mathcal{P}(s) = \emptyset$ (équation 3.12) parce que :

- $\forall t \in sa(t_1, \mathcal{R}(t_1)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \bigcap_{s \in \mathcal{W}(t, n)} \mathcal{P}(s) = \emptyset$ parce que $t_1 \in \mathcal{T}^e$

- $\forall t \in sa(t_2, \mathcal{R}(t_2)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \bigcap_{s \in \mathcal{W}(t, n)} \mathcal{P}(s) = \emptyset$ parce que $t_2 \in \mathcal{T}^e$
- $\bigcap_{s \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2))} \mathcal{P}(s) = \emptyset$ parce que $\forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \quad \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset$

Donc $t_3 \in \mathcal{T}^e$. De plus, en utilisant la construction de t_3 , on obtient $\mathcal{A}(\mathcal{R}(\overline{M}(t_1, t_2, p_1, p_2))) = \min(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{A}(\mathcal{R}(t_2))) = \mathcal{A}(\mathcal{R}(t_1)) = \mathcal{A}(\mathcal{R}(t_2))$ et $\mathcal{P}(\mathcal{R}(\overline{M}(t_1, t_2, p_1, p_2))) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2))$.

cas 4 $(\mathcal{A}(\mathcal{R}(t_1)) < \mathcal{A}(\mathcal{R}(t_2))) \wedge (|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 0)$

Dans ce cas, on peut faire le même raisonnement que précédemment, en substituant $\mathcal{W}(t_2)$ par t_2 .

Dans les deux cas terminaux, le théorème 3.2.1 est vérifié.

Cas non terminaux Soit $(t_1, t_2) \in \mathcal{T}^{e^2} / \forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \quad \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset$ et $(p_1, p_2) \in \Omega$ et $t_3 = \overline{M}(t_1, t_2, p_1, p_2)$.

cas 2 $(\mathcal{A}(\mathcal{R}(t_1)) = \mathcal{A}(\mathcal{R}(t_2))) \wedge (|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| = 1) \wedge (|Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))| = 1)$

En supposant que les hypothèses du théorème 3.2.1 sont vérifiées pour les éléments $Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))$, $Q_{p_2}(sa(t_2, \mathcal{R}(t_2)))$, p_1 et p_2 , on obtient (équation 3.11) :

$$\forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$$

parce que :

- $\forall t \in sa(t_1, \mathcal{R}(t_1)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$
- $\forall t \in sa(t_2, \mathcal{R}(t_2)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$
- $\forall n \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2)) \quad \mathcal{P}(n) \subseteq \mathcal{P}(\mathcal{R}(t_3))$ car $\mathcal{P}(\mathcal{R}(t_3)) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2))$
- par hypothèse $\mathcal{P}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2)))$ est égal à $\mathcal{P}(\mathcal{R}((Q_{p_1}(sa(t_1, \mathcal{R}(t_1))) \cup \mathcal{P}(Q_{p_2}(sa(t_2, \mathcal{R}(t_2))))))$

De manière similaire (équation 3.10) :

$$\forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(n) \quad \mathcal{A}(s) > \mathcal{A}(n)$$

car :

- $\forall t \in sa(t_1, \mathcal{R}(t_1)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$
- $\forall t \in sa(t_2, \mathcal{R}(t_2)), \forall n \in \mathcal{N}(t), \forall s \in \mathcal{W}(t, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$
- $\forall n \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2)) \quad \mathcal{A}(n) > \mathcal{A}(\mathcal{R}(t_3))$ parce que $\mathcal{A}(\mathcal{R}(t_3)) = \min(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{A}(\mathcal{R}(t_2)))$
- par hypothèse $\mathcal{A}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2)))$ est égal à $\min(\mathcal{A}(\mathcal{R}((Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), \mathcal{A}(Q_{p_2}(sa(t_2, \mathcal{R}(t_2))))))$

Finalement, on a (équation 3.12) :

$$\forall n \in \mathcal{N}(t_3) \quad \bigcap_{s \in \mathcal{W}(t_3, n)} \mathcal{P}(s) = \emptyset$$

parce que :

- $\forall t \in sa(t_1, \mathcal{R}(t_1)), \forall n \in \mathcal{N}(t) \quad \bigcap_{s \in \mathcal{W}(t, n)} \mathcal{P}(s) = \emptyset$
- $\forall t \in sa(t_2, \mathcal{R}(t_2)), \forall n \in \mathcal{N}(t) \quad \bigcap_{s \in \mathcal{W}(t, n)} \mathcal{P}(s) = \emptyset$
- $\bigcap_{s \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2))} \mathcal{P}(s) = \emptyset$ car $s \in \mathcal{W}(t_1, \mathcal{R}(t_1)) \cup \mathcal{W}(t_2, \mathcal{R}(t_2)) \implies s \in \mathcal{W}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2))$
- $\forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \quad \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset$
- et $\mathcal{P}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), Q_{p_2}(sa(t_2, \mathcal{R}(t_2))), p_1, p_2))) = \mathcal{P}(\mathcal{R}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))) \cup \mathcal{P}(\mathcal{R}(Q_{p_2}(sa(t_2, \mathcal{R}(t_2))))))$
- et $\bigcap_{s \in \mathcal{W}(t_1, \mathcal{R}(t_1))} \mathcal{P}(s) = \emptyset$
- et $\bigcap_{s \in \mathcal{W}(t_2, \mathcal{R}(t_2))} \mathcal{P}(s) = \emptyset$

Sous les hypothèses précédentes, on obtient que $t_3 \in \mathcal{T}^e$, et en remarquant que $\mathcal{A}(\mathcal{R}(\overline{M}(t_1, t_2, p_1, p_2))) = \min(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{A}(\mathcal{R}(t_2)))$ et que $\mathcal{P}(\mathcal{R}(\overline{M}(t_1, t_2, p_1, p_2))) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2))$ on arrive à la conclusion que sous les mêmes hypothèses le théorème est vérifié.

cas 3 : $(\mathcal{A}(\mathcal{R}(t_1)) < \mathcal{A}(\mathcal{R}(t_2))) \wedge (|Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))| \neq 0)$

En supposant que le théorème 3.2.1 est vérifié pour $\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2)$ on obtient :

- $\forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(t_3, n) \quad \mathcal{P}(s) \subseteq \mathcal{P}(n)$ (équation 3.11) car :
 - $\mathcal{P}(\mathcal{R}(t_3)) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2)))$
 - et $t_1 \in \mathcal{T}^e$
 - et par hypothèse, on a $\mathcal{P}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2))) = \mathcal{P}(\mathcal{R}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))))) \cup \mathcal{P}(\mathcal{R}(t_2))$
- $\forall n \in \mathcal{N}(t_3), \forall s \in \mathcal{W}(t_3, n) \quad \mathcal{A}(s) > \mathcal{A}(n)$ (équation 3.10) car :
 - $\mathcal{A}(\mathcal{R}(t_3)) = \min(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{A}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2))))$
 - et $t_1 \in \mathcal{T}^e$
 - et par hypothèse $\mathcal{A}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2)))$ est égal à $\min(\mathcal{A}(\mathcal{R}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))))) , \mathcal{A}(\mathcal{R}(t_2)))$
- $\forall n \in \mathcal{N}(t_3) \quad \bigcap_{s \in \mathcal{W}(t_3, n)} \mathcal{P}(s) = \emptyset$ (équation 3.12) car :
 - par hypothèse $\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2) \in \mathcal{T}^e$
 - $\forall t \in sa(t_3, \mathcal{R}(t_3)) \quad t \in \mathcal{T}^e$
 - $\bigcap_{s \in \mathcal{W}(t_3, \mathcal{R}(t_3))} \mathcal{P}(s) = \emptyset$ car $\forall n_1 \in \mathcal{N}(t_1), \forall n_2 \in \mathcal{N}(t_2) \quad \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset$ et $\mathcal{P}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2))) = \mathcal{P}(\mathcal{R}(t_2)) \cup \mathcal{P}(\mathcal{R}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1)))))$

Ainsi sous l'hypothèse précédente, $t_3 \in \mathcal{T}^e$. De plus, en remarquant que $\mathcal{A}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2))) = \min(\mathcal{A}(\mathcal{R}(t_1)), \mathcal{A}(\mathcal{R}(t_2)))$ et que $\mathcal{P}(\mathcal{R}(\overline{M}(Q_{p_1}(sa(t_1, \mathcal{R}(t_1))), t_2, p_1, p_2))) = \mathcal{P}(\mathcal{R}(t_1)) \cup \mathcal{P}(\mathcal{R}(t_2))$, on peut conclure que le théorème 3.2.1 est vérifié car $t_3 \in \mathcal{T}^e$.

cas 5 Dans les autres configurations, on peut se ramener aux cas précédents $\overline{M}(t_1, t_2, p_1, p_2) = \overline{M}(t_2, t_1, p_2, p_1)$.

Conclusion Puisque tous les cas terminaux et toutes les étapes de récurrence vérifient le théorème 3.2.1, ce dernier est vrai $\forall (t_1, t_2) \in \mathcal{T}^{e^2}, \forall (p_1, p_2) \in \Omega$. $\square$

C.3 Preuve du théorème 3.2.4

Démonstration. Soit $\tilde{p} \in \Omega, P = \{p_1..p_k\} \in 2^\Omega, \tilde{t} \in \mathcal{T}^e$. Soit l'hypothèse de récurrence $H_i, i \in [[0..k]]$ définie comme : $M(\tilde{p}, \{p_1..p_i\}, \tilde{t}) \in \mathcal{T}^e$

Initialisation : Par définition, $M(\tilde{p}, \{\}, \tilde{t}) = \tilde{t} \in \mathcal{T}^e$ donc H_0 est vraie.

Pas de récurrence : Supposons que H_i soit vraie

- $H_i \Rightarrow M(\tilde{p}, \{p_1..p_i\}, \tilde{t}) \in \mathcal{T}^e$
- $M(\tilde{p}, \{p_1..p_{i+1}\}, \tilde{t}) = \text{subst}(M(\tilde{p}, \{p_1..p_i\}, t), t', t'')$ avec, en utilisant la définition de M , $\mathcal{R}(t') = \mathcal{R}(t'')$
- Montrons que $t'' \in \mathcal{T}^e$
 - $Q_{\tilde{p}}(T)$ et $Q_{p_i}(T)$ sont des sous-arbres d'un même emboîtement car H_i est vraie.
 - De plus, par définition de M , on a $\mathcal{R}(Q_{\tilde{p}}(T)) \in \mathcal{W}(t', \mathcal{R}(t'))$ et $\mathcal{R}(Q_{p_i}(T)) \in \mathcal{W}(t', \mathcal{R}(t'))$ ainsi $\mathcal{P}(\mathcal{R}(Q_{\tilde{p}})) \cap \mathcal{P}(\mathcal{R}(Q_{p_i})) = \emptyset$.
 - En utilisant l'équation 3.11 sur $Q_{\tilde{p}}(T)$ et $Q_{p_i}(T)$, on obtient $\forall n_1 \in \mathcal{N}(Q_{\tilde{p}}(T)), \forall n_2 \in \mathcal{N}(Q_{p_i}(T)) \quad \mathcal{P}(n_1) \cap \mathcal{P}(n_2) = \emptyset$.
 - Ainsi les conditions du théorème 3.2.1 sont vérifiées. On a donc :
 - $\overline{M}(Q_{\tilde{p}}(T), Q_{p_i}(T), \tilde{p}, p_i) \in \mathcal{T}^e$
 - $\mathcal{A}(\mathcal{R}(\overline{M}(Q_{\tilde{p}}(T), Q_{p_i}(T), \tilde{p}, p_i))) = \min(\mathcal{A}(\mathcal{R}(Q_{\tilde{p}}(T))), \mathcal{A}(\mathcal{R}(Q_{p_i}(T))))$
 - $\mathcal{P}(\mathcal{R}(\overline{M}(Q_{\tilde{p}}(T), Q_{p_i}(T), \tilde{p}, p_i))) = \mathcal{P}(\mathcal{R}(Q_{\tilde{p}}(T))) \cup \mathcal{P}(\mathcal{R}(Q_{p_i}(T)))$
 - Par définition de t'' , les seuls nœuds et arcs de t' qui peuvent être modifiés sont ceux de $Q_{\tilde{p}}(T)$ et $Q_{p_i}(T)$. Or comme par hypothèse les deux derniers arbres sont remplacés par $\overline{M}(Q_{\tilde{p}}(T), Q_{p_i}(T), \tilde{p}, p_i)$, qui est un emboîtement dont le degré d'appartenance associé au nœud racine est le minimum

- de ceux de $Q_{\tilde{p}}(T)$ et $Q_{p_i}(T)$, et comme l'ensemble de points de cette nouvelle racine est l'union de $\mathcal{P}(\mathcal{R}(Q_{\tilde{p}}(T)))$ et de $\mathcal{P}(\mathcal{R}(Q_{p_i}(T)))$, on a $t'' \in \mathcal{T}^e$.
- Puisque $M(\tilde{p}, \{p_1..p_i\}, \tilde{t}) \in \mathcal{T}^e$ et $t' \in \mathcal{T}^e$ et $t'' \in \mathcal{T}^e$, on peut utiliser le théorème 3.2.3 et conclure que H_{i+1} est vraie.

Conclusion Puisque H_0 est vrai et puisque $\forall i \in [[0..k-1]] \quad H_i \Rightarrow H_{i+1}$, on peut conclure que H_i est vraie pour tout i . De plus, la démonstration de H_i est valide pour tout P , $\tilde{p}$ et $\tilde{t}$ sans hypothèses aucune sur ces derniers, on peut donc conclure que le théorème 3.2.4 est vérifié. $\square$

C.4 Preuve du théorème 3.2.5

Démonstration. Montrons les équations 3.13, 3.14 et 3.15.

Equation 3.13 $M(\tilde{p}, \text{voisinage}(\tilde{p}), \tilde{t}) \in \mathcal{T}^e$

L'arbre $\tilde{t}$ est un emboîtement car $\tilde{p}$ appartient au père de $S_{\tilde{p}}^{f'(\tilde{p})}$ dont le degré d'appartenance est plus grand que celui de son père (par définition de $f'(\tilde{p}) > f(\tilde{p})$), et car aucun frère de $S_{\tilde{p}}^{f'(\tilde{p})}$ ne contient $\tilde{p}$. En utilisant le théorème 3.2.4, on obtient que $M(\tilde{p}, \text{voisinage}(\tilde{p}), \tilde{t}) \in \mathcal{T}^e$.

Equation 3.14 $\forall p \in \Omega, \forall \alpha \in [0, 1] \quad p \in f_\alpha \Rightarrow \exists n \in \mathcal{N}(M(\tilde{p}, \text{voisinage}(\tilde{p}), \tilde{t}))/\mathcal{P}(n) = \hat{\Gamma}_{f_\alpha}^p$

Soient $p \in \Omega$ et $\alpha \in]f(\tilde{p}), 1]$. Soit $\{p_1..p_k\} = \text{voisinage}(\tilde{p})$ l'ensemble des k voisins de $\tilde{p}$. Soit $H_i, i \leq k$ l'hypothèse de récurrence suivante :

$$p \in \hat{\Gamma}_{f'_\alpha}^{\tilde{p}} \Rightarrow \exists n \in \mathcal{N}(M(\tilde{p}, \{p_1..p_i\}, \tilde{t}))/\mathcal{P}(n) = \bigcup_{p' \in \{p_1..p_i\}/p' \in f_\alpha} \hat{\Gamma}_{f_\alpha}^{p'} \cup \{\tilde{p}\}$$

Initialisation $M(\tilde{p}, \{\}, \tilde{t}) = \tilde{t}$, or $S_{\tilde{p}}^{f'(\tilde{p})} \in \mathcal{N}(\tilde{t})$ ainsi H_0 est vraie.

Pas de récurrence Supposons H_i vraie

- $M(\tilde{p}, \{p_1..p_{i+1}\}, \tilde{t}) = \text{subst}(M(\tilde{p}, \{p_1..p_i\}, \tilde{t}), t', t'')$
- par hypothèse, il existe $n_1 \in \mathcal{N}(M(\tilde{p}, \{p_1..p_i\}, \tilde{t}))$ tel que $\mathcal{P}(n_1) = \bigcup_{p' \in \{p_1..p_i\}/p' \in f_\alpha} \hat{\Gamma}_{f_\alpha}^{p'} \cup \{\tilde{p}\}$
- si $p_{i+1} \in f_\alpha$ (ou encore $p_{i+1} \in f'_\alpha$) alors $\exists n_2 \in \mathcal{N}(t')/\mathcal{P}(n_2) = \hat{\Gamma}_{f_\alpha}^{p_{i+1}}$
- puisque p_{i+1} et $\tilde{p}$ sont voisins, nous sommes assurés de l'existence d'un couple (n_1, n_2) de nœuds compatibles qui sont par conséquent fusionnés par l'opérateur $\overline{M}$ (c.f. théorème 3.2.2).
- Ainsi, on a l'existence d'un nœud $\tilde{n}$ dans t'' et par conséquent dans $M(\tilde{p}, \{p_1..p_{i+1}\}, \tilde{t})$ qui vérifie $\mathcal{P}(\tilde{n}) = \bigcup_{p' \in \{p_1..p_{i+1}\}/p' \in f_\alpha} \hat{\Gamma}_{f_\alpha}^{p'} \cup \{\tilde{p}\}$. Cela signifie que H_{i+1} est vraie.

Conclusion Par récursion, H_i est vraie pour tout $i \in [[0..k]]$.

Si on utilise ce résultat, avec $\{p_1..p_k\}$ l'ensemble $\text{voisinage}(\tilde{p})$ de voisins de $\tilde{p}$, et si on remarque que les configurations $\alpha \in [0, 1], p \in \Omega$ telles que $p \notin \hat{\Gamma}_{f'_\alpha}^{\tilde{p}} \hat{\Gamma}_{f_\alpha}^p = \hat{\Gamma}_{f'_\alpha}^p$ et que les nœuds concernés se sont pas modifiés par M (ligne en trait plein dans les figures 3.9(c), 3.9(e) et 3.9(g)), on peut en déduire que l'équation 3.14 est vérifiée.

Equation 3.15 $\forall n \in \mathcal{N}(M(\tilde{p}, \text{voisinage}(\tilde{p}), \tilde{t})), \forall p \in \Omega \quad p \in \mathcal{P}(n) \Rightarrow \hat{\Gamma}_{f_{\mathcal{A}(n)}}^p = \mathcal{P}(n)$

En remarquant que les nœuds $n \in \mathcal{N}(t)$ contenant un voisin de $\tilde{p}$ qui vérifient $f(\tilde{p}) < \mathcal{A}(n) \leq f'(\tilde{p})$ (i.e. les nœuds qui ne représentent pas une composante connexe dans f') sont supprimés par les différents appels à $\overline{M}$ (arbres en pointillés rouges sur les figures 3.9(b), 3.9(d) et 3.9(f)), et que comme on vient de le voir les seuls nœuds introduits correspondent aux composantes connexes de f' qui n'étaient pas déjà présentes dans f (arbres en pointillés bleus sur les figures 3.9(c), 3.9(e) et 3.9(g)), on montre que l'équation 3.15 est vérifiée.

Conclusion Comme les équations 3.13, 3.14 et 3.15 sont vérifiées par $M(\tilde{p}, \text{voisinage}(\tilde{p}), \tilde{t})$, ce dernier est une représentation de f' . $\square$

Remarque C.4.1. Dans l'étape de pas de récursion de la démonstration de l'équation 3.14, il peut y avoir plusieurs couples de nœuds (n_1, n_2) qui vérifient $\mathcal{P}(n_1) = \bigcup_{p' \in \{p_1 \dots p_i\} / p' \in f_\alpha} \hat{\Gamma}_{f_\alpha}^{p'} \cup \{\tilde{p}\}$ et $\mathcal{P}(n_2) = \hat{\Gamma}_{f_\alpha}^{p_{i+1}}$.

Néanmoins, seul un d'entre eux est composé de nœuds compatibles. Dans l'exemple de la figure 3.9(d), pour $\alpha_4 < \alpha \leq \alpha_5$, il existe seulement un candidat pour n_1 ($\mathcal{P}(n_1) = E \cup \{\tilde{p}\}$ et $\mathcal{A}(n_1) = \alpha_5$) et deux candidats pour n_2 ($\mathcal{P}(n_2) = D$ et $\mathcal{A}(n_2) = \alpha_5$ ou $\mathcal{A}(n_2) = \alpha_6$) mais seulement un nœud n_2 qui vérifie que $\mathcal{A}(n_2) = \alpha_5$ est compatible avec n_1 . Dans le même exemple (c.f. figure 3.9(e)), ces nœuds produisent un nœud $\tilde{n}$ tel que $\mathcal{P}(\tilde{n}) = E \cup D \cup \{\tilde{p}\}$ et $\mathcal{A}(\tilde{n}) = \alpha_5$.

C.5 Preuve du théorème 3.3.2

Démonstration. Soit $t \in \mathcal{T}^e$, $\tilde{p} \in \Omega$ et $\alpha \in [0, 1]$.

Montrons tout d'abord que $\mathcal{Y}(t, \tilde{p}, \alpha)$ est bien un emboîtement ($\mathcal{Y}(t, \tilde{p}, \alpha) \in \mathcal{T}^e$) :

Montrons les équations 3.10 et 3.11 $\forall n \in \mathcal{N}(\mathcal{Y}(t, \tilde{p}, \alpha)), \forall s \in \mathcal{W}(\mathcal{Y}(t, \tilde{p}, \alpha), n)$, deux cas sont possibles :

- soit $s \in \mathcal{W}(t, n)$ ce qui implique que les équations 3.10 et 3.11 sont vérifiées car t est un emboîtement,
- ou $s \notin \mathcal{W}(t, n)$, dans ce cas on a, en utilisant l'équation 3.31, on a que $s \in \mathcal{D}(t, n)$ et donc, en utilisant le théorème 3.3.1, que les équations 3.10 et 3.11 sont vérifiées.

Montrons l'équation 3.12 $\forall n \in \mathcal{N}(\mathcal{Y}(t, \tilde{p}, \alpha)), \forall s \in \mathcal{W}(\mathcal{Y}(t, \tilde{p}, \alpha), n)$, on a de manière exclusive :

- soit $s \in \mathcal{W}(t, n)$,
- ou $\exists! s' \in \mathcal{W}(t, n) / (s \in \mathcal{D}(t, s')) \wedge (\forall s'' \neq s \in \mathcal{W}(\mathcal{Y}(t, \tilde{p}, \alpha), n) s'' \notin \mathcal{D}(t, s'))$

c'est-à-dire que pour un nœud donné, chacun de ses fils est soit remplacé de manière unique par lui-même, ou par un de ses descendants (alternativement il peut disparaître).

En utilisant le théorème 3.3.1, on a donc l'équation 3.12 qui est vérifiée. On a donc bien $\mathcal{Y}(t, \tilde{p}, \alpha)$ qui est un emboîtement.

Montrons maintenant que $\mathcal{Y}(t, \tilde{p}, \alpha)$ est une représentation de $f' \subseteq f^t$ définie comme :

$$\forall p \in \Omega \quad f'(p) = \max \left\{ \beta \in [0, 1] / (p \in f_\beta^t) \wedge \left((\beta \leq f^t(\tilde{p})) \vee (\beta > f^t(\tilde{p})) \vee (p \notin \Gamma_{f_\beta^t}^{\tilde{p}}) \right) \right\} \quad (\text{C.1})$$

Montrons que l'équation 3.14 est vérifiée Soit $\beta \in [0, 1]$ et $p \in f'_\beta$.

- si $\beta > f^t(\tilde{p})$ $\Gamma_{f'_\beta}^p = \Gamma_{f_\beta^t}^p$ et les nœuds correspondants sont conservés dans $\mathcal{Y}(t, \tilde{p}, \alpha)$ (c.f. nœuds dont le degré d'appartenance est strictement supérieur à α_4 dans les figures 3.10(c) et 3.10(d))
- si $\beta \leq \max_{n \in \mathcal{N}(t) / (\tilde{p} \in \mathcal{P}(n)) \wedge \mathcal{A}(n) \leq \alpha} \mathcal{A}(n)$, on est dans le même cas (c.f. les nœuds dont le degré d'appartenance est inférieur ou égal à α_1 dans les figures 3.10(c) et 3.10(d)),
- sinon $\max_{n \in \mathcal{N}(t) / (\tilde{p} \in \mathcal{P}(n)) \wedge \mathcal{A}(n) \leq \alpha} \mathcal{A}(n) < \beta \leq f^t(\tilde{p})$ et pour un tel degré d'appartenance, seul les nœuds de t

qui ne contiennent pas $\tilde{p}$ sont conservés dans $\mathcal{Y}(t, \tilde{p}, \alpha)$ (c.f. équation 3.30). De plus si $p \notin \Gamma_{f_\beta^t}^{\tilde{p}}$, on a $\Gamma_{f'_\beta}^p = \Gamma_{f_\beta^t}^p$ (c.f. nœuds dont le degré d'appartenance est α_3 ou α_4 dans la figure 3.10(d)). Enfin si

$p \in \Gamma_{f_\beta^t}^{\tilde{p}}$ on a l'existence d'un $\beta' \leq \max_{n \in \mathcal{N}(t) / (\tilde{p} \in \mathcal{P}(n)) \wedge \mathcal{A}(n) \leq \alpha} \mathcal{A}(n)$ ou $\beta' > f^t(p)$ qui vérifie $\Gamma_{f'_\beta}^p = \Gamma_{f'_\beta}^p$

(c.f. équation C.1), et on peut donc se reporter à l'un des deux premiers cas.

L'arbre $\mathcal{Y}(t, \tilde{p}, \alpha)$ vérifie donc l'équation 3.14.

Montrons que l'équation 3.15 est vérifiée Par construction de $\mathcal{Y}(t, \tilde{p}, \alpha)$ (équation 3.30) et par définition de f' (équation C.1), l'équation 3.15 est vérifiée.

En conclusion, le théorème 3.3.2 est bien vérifié. $\square$

Remarque C.5.1. L'expression de f' donnée à l'équation C.1 s'obtient à partir des équations 3.30 et 3.31 où l'opérateur $\mathcal{Y}$ ne fait que supprimer des nœuds (c.f. équation 3.30), et ces derniers ne correspondent qu'aux nœuds de $\mathcal{N}(t)$ qui ont un degré d'appartenance compris entre α et $f^t(\tilde{p})$ et qui contiennent $\tilde{p}$ (c.f. définition de $Z_{f^t(\tilde{p}),\alpha}^{\tilde{p}}$). Les figures 3.10(b) et 3.10(d) illustrent l'équation C.1.

Annexe D

Nombre d'éléments dans B_γ et $\overline{B_\gamma}$

Dans cette annexe, nous allons partiellement justifier les équations 7.19 et 7.20 correspondant au cardinal des éléments structurants utilisés pour le calcul de la mesure de convergence introduite au chapitre 7.

Pour des raisons de clarté et dans la mesure où nous souhaitons juste justifier la complexité de l'approche de détection de convergence, nous ne montrerons que les inégalités de ces équations, c'est-à-dire :

$$|B_\gamma| \leq 4 \left(\int_0^{\frac{\alpha R_{\max}}{\sqrt{\alpha^2+1}}} \sqrt{R_{\max}^2 - x^2} dx - \int_0^{\alpha R_{\max} \cos(\frac{\pi}{4})} x dx - \int_{\alpha R_{\max} \cos(\frac{\pi}{4})}^{\frac{\alpha R_{\max}}{\sqrt{\alpha^2+1}}} \sqrt{\frac{x^4}{(\alpha R_{\max})^2 - x^2}} dx \right) \quad (D.1)$$

et

$$|\overline{B_\gamma}| \leq 4 \left(\int_0^{\frac{\alpha R_{\max}}{\sqrt{2}}} \sqrt{(\alpha R_{\max})^2 - x^2} dx - \int_0^{\frac{\alpha R_{\max}}{\sqrt{2}}} x dx \right) \quad (D.2)$$

Cela nous permettra de justifier que les complexités exposées au chapitre 7 sont des bornes supérieures.

D.1 Cardinal de B_γ

Pour montrer l'équation D.1, nous allons introduire trois théorèmes correspondant aux trois parties de cette équation.

Théorème D.1.1.

$$\forall c \in \Omega, \forall q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0 \quad \exists r \in [[R_{\min}; R_{\max}]]/K_{c,q,r} = 1 \Rightarrow x < y$$

Démonstration. Soit $c \in \Omega, q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0$. Soit $r \in [[R_{\min}; R_{\max}]]/K_{c,q,r} = 1$.

On a par définition de $K_{c,q,r}$:

$$\begin{aligned} (\alpha r < \|\vec{c\bar{q}}\|) \wedge (\tan(\theta)\|\vec{c\bar{q}}\| \leq \alpha r) &\Rightarrow \tan(\theta)\|\vec{c\bar{q}}\| \leq \alpha r < \|\vec{c\bar{q}}\| \\ &\Rightarrow \tan(\theta)\|\vec{c\bar{q}}\| < \|\vec{c\bar{q}}\| \\ &\Rightarrow \tan(\theta) < 1 \\ &\Rightarrow \frac{x}{y} < 1 \\ &\Rightarrow x < y \end{aligned}$$

Le théorème est donc bien vérifié. □

Théorème D.1.2.

$$\begin{aligned} \forall c \in \Omega, \forall q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0 \\ \exists r \in [[R_{\min}; R_{\max}]]/K_{c,q,r} = 1 \Rightarrow y < \sqrt{R_{\max}^2 - x^2} \end{aligned}$$

Démonstration. Soit $c \in \Omega, q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0$. Soit $r \in [[R_{\min}; R_{\max}]]/K_{c,q,r} = 1$. Par définition, on s'interdit de considérer des $r > R_{\max}$, donc $r \leq R_{\max}$. En utilisant le fait que $\|\vec{cq}\| < r$ (définition de $K_{c,q,r}$), on a :

$$\begin{aligned} \|\vec{cq}\| < r \leq R_{\max} &\Rightarrow \sqrt{x^2 + y^2} < R_{\max} \\ &\Rightarrow x^2 + y^2 < R_{\max}^2 \\ &\Rightarrow y < \sqrt{R_{\max}^2 - x^2} \end{aligned}$$

Le théorème est donc bien vérifié. $\square$

Théorème D.1.3.

$$\begin{aligned} \forall c \in \Omega, \forall q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0 \\ \exists r \in [[R_{\min}; R_{\max}]]/K_{c,q,r} = 1 \Rightarrow y > \sqrt{\frac{x^4}{\sqrt{(\alpha R_{\max})^2 - x^2}}} \end{aligned}$$

Démonstration. Soit $c \in \Omega, q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0$. Soit $r \in [[R_{\min}; R_{\max}]]/K_{c,q,r} = 1$. Par définition, on a :

$$K_{cqr} = 1 \Rightarrow \tan(\theta) \|cq\| \leq \alpha r$$

Comme $r \leq R_{\max}$, on a aussi :

$$\begin{aligned} \tan(\theta) \|cq\| \leq \alpha R_{\max} &\Rightarrow \frac{x}{y} \sqrt{x^2 + y^2} \leq \alpha R_{\max} \\ &\Rightarrow x^2(x^2 + y^2) \leq y^2 \alpha^2 R_{\max}^2 \\ &\Rightarrow x^4 + x^2 y^2 - y^2 \alpha^2 R_{\max}^2 \leq 0 \\ &\Rightarrow y^2 \leq \frac{-x^4}{x^2 - \alpha^2 R_{\max}^2} \Rightarrow y \leq \sqrt{\frac{x^4}{\alpha^2 R_{\max}^2 - x^2}} \text{ car } \alpha^2 R_{\max}^2 - x^2 \text{ est positif} \end{aligned}$$

On a donc bien que le théorème D.1.3 est vérifié. $\square$

En utilisant les théorèmes D.1.1, D.1.2 et D.1.3, on a que les éléments contenus dans le quart supérieur droit de l'élément structurant B sont dans la zone définie par l'équation D.1. Les bornes des intégrales s'obtiennent facilement en posant les égalités entre les différentes équations de y en fonction de x . La symétrie de l'élément structurant permet de compter les individus dans les autres quadrants (coefficient 4 au début de l'équation).

D.2 Cardinal de $\overline{B_\gamma}$

Pour démontrer l'équation D.2, nous avons besoin d'un dernier théorème.

Théorème D.2.1.

$$\begin{aligned} \forall c \in \Omega, \forall q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0 \\ \exists r \in [[R_{\min}; R_{\max}]]/(K_{c,q,r} = 0) \wedge (K_{c,q,r-1} = 1) \Rightarrow y \leq \sqrt{(\alpha R_{\max})^2 - x^2} \end{aligned}$$

Démonstration. Soit $c \in \Omega, q \in \Omega/c - q = (x, y), x \geq 0, y \geq 0$. Soit $r \in [[R_{\min}; R_{\max}]]/(K_{c,q,r} = 0) \wedge (K_{c,q,r-1} = 1)$.

Trois cas sont possibles :

– cas 1 : $K_{c,q,r} = 0 \Rightarrow \alpha r \geq \|cq\|$, or $r \leq R_{\max}$, donc :

$$\begin{aligned} \alpha R_{\max} \geq \alpha r \geq \|cq\| &\Rightarrow \alpha R_{\max} \geq \|cq\| \\ &\Rightarrow (\alpha R_{\max})^2 \geq x^2 + y^2 \\ &\Rightarrow y \leq \sqrt{(\alpha R_{\max})^2 - x^2} \end{aligned}$$

– cas 2 : $K_{c,q,r} = 0 \Rightarrow \|cq\| \geq r$
or $\|cq\| < r - 1$ car $K_{c,q,r-1} = 1$, on a donc $r \leq \|cq\| < r - 1$ alors que $r > r - 1$. Ce cas est donc impossible.

- cas 3 : $K_{c,q,r} = 0 \Rightarrow \tan(\theta)||cq|| > \alpha r$
 or $\tan(\theta)||cq|| \leq \alpha(r-1)$ car $K_{c,q,r-1} = 1$
 on a donc $\alpha(r-1) \geq \tan(\theta)||cq|| > \alpha r$ ce qui implique que $r-1 \geq r \Rightarrow$ alors que $r > r-1$. Ce cas est donc impossible.

Les cas 2 et 3 étant impossibles et le cas 1 vérifiant l'implication, on a que le théorème D.2.1 est vrai. $\square$

En utilisant ce théorème ainsi que le théorème D.1.1, on obtient directement l'équation D.2. Comme pour le cardinal de B , les bornes de l'intégrale dans l'équation D.2 s'obtiennent en posant l'égalité entre les équations de y en fonction de x provenant des deux théorèmes utilisés.

Liste des publications

- Palma, G., Bloch, I. et Muller, S. (2008a). Fuzzy connected filters for fuzzy gray scale images. *International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU)*, pages 667–674, Malaga, Spain
- Palma, G., Nempont, O., Bloch, I. et Muller, S. (2008b). Extraction de "zones plates floues" dans des images de quantités floues. *Rencontre Francophones sur la Logique Floue et ses Applications*, pages 364–371, Lens, France
- Palma, G., Peters, G., Muller, S. et Bloch, I. (2008c). Masses classification using fuzzy active contours and fuzzy decision trees. *SPIE Symposium on Medical Imaging*, numéro 6915, San Diego, CA, USA
- Palma, G., Bloch, I., Muller, S. et Iordache, R. (2009a). Fuzzifying images using fuzzy wavelet denoising. *FUZZ-IEEE*, pages 135–140, Jeju, Korea
- Palma, G., Muller, S., Bloch, I. et Iordache, R. (2009b). Convergence areas detection in digital breast tomosynthesis volumes using a contrario modeling. *SPIE Symposium on Medical Imaging*, Lake Buena Vista, FL, USA
- Palma, G., Muller, S., Bloch, I. et Iordache, R. (2009c). Fast detection of convergence areas in digital breast tomosynthesis. *IEEE International Symposium on Biomedical Imaging (ISBI)*, pages 847–850, Boston, MA, USA
- Apffel, L., Palma, G., Bloch, I. et Muller, S. (2010). Fuzzy segmentation of masses in digital breast tomosynthesis images based on dynamic programming. *International Conference on Imaging Theory and Applications (IMAGAPP)*, Angers, France
- Palma, G., Bloch, I. et Muller, S. (2010a). Fast fuzzy connected filter implementation using max-tree updates. *Fuzzy Systems and Systems*, 161(1):118–146
- Palma, G., Bloch, I. et Muller, S. (2010b). Spiculated lesions and architectural distortions detection in digital breast tomosynthesis datasets. *International Workshop on Digital Mammography (IWDM)*, Girona, Spain

Bibliographie

- Abdel-Mottaleb, M., Carman, C., Hill, C. et Vafai, S. (1996). Locating the boundary between the breast skin edge and the background in digitized mammograms. *International Workshop on Digital Mammography (IWDM)*, pages 221–228.
- ACR (2003). *BI-RADS Breast Imaging Reporting and Data System*. American College of Radiology.
- Andersen, A. H. et Kak, A. C. (1984). Simultaneous algebraic reconstruction technique (SART) : A superior implementation of the ART algorithm. *Ultrasonic Imaging*, 6(1):81–94.
- Anscombe, F. J. (1948). The transformation of poisson, binomial and negative-binomial data. *Biometrika*, 35(3-4):246–254.
- Apffel, L., Palma, G., Bloch, I. et Muller, S. (2010). Fuzzy segmentation of masses in digital breast tomosynthesis images based on dynamic programming. *International Conference on Imaging Theory and Applications (IMAGAPP)*, Angers, France.
- Arbach, L., Reinhardt, J., Bennett, D. et Fallouh, G. (2003). Mammographic masses classification : comparison between backpropagation neural network (BNN), K nearest neighbors (KNN), and human readers. *IEEE Canadian Conference on Electrical and Computer Engineering (CCECE)*, 3:1441–1444.
- Avinash, G., Israni, K. et Li, B. (2006). Characterization of point spread function in linear digital tomosynthesis : a simulation study. *SPIE Symposium on Medical Imaging*, volume 6142, page 614258.
- Ayres, F. et Rangayyan, R. (2005). Characterization of architectural distortion in mammograms. *IEEE Engineering in Medicine and Biology Magazine*, 24(1):59–67.
- Ayres, F. J. et Rangayyan, R. M. (2007). Reduction of false positives in the detection of architectural distortion in mammograms by using a geometrically constrained phase portrait model. *International Journal of Computer Assisted Radiology and Surgery*, pages 361–369.
- Ball, J. et Bruce, L. (2007a). Digital mammogram spiculated mass detection and spicule segmentation using level sets. *IEEE Engineering in Medicine and Biology Society*, pages 4979–4984.
- Ball, J. et Bruce, L. (2007b). Digital mammographic computer aided diagnosis (CAD) using adaptive level set segmentation. *IEEE Engineering in Medicine and Biology Society*, pages 4973–4978.
- Ball, J. et Bruce, L. (2007c). Level set-based core segmentation of mammographic masses facilitating three stage (core, periphery, spiculation) analysis. *IEEE Engineering in Medicine and Biology Society*, pages 819–824.
- Bankman, I., Tsai, J., Kim, D., Gatewood, O. et Brody, W. (1994). Detection of microcalcification clusters using neural networks. *IEEE Engineering in Medicine and Biology Society*, pages 590–591.
- Berger, C., Géraud, T., Levillain, R., Widynski, N., Baillard, A. et Bertin, E. (2007). Effective component tree computation with application to pattern recognition in astronomical imaging. *IEEE International Conference on Image Processing*, pages 41–44.

- Bernard, S., Muller, S. et Onativia, J. (2008). Computer-aided microcalcification detection on digital breast tomosynthesis data : A preliminary evaluation. *International Workshop on Digital Mammography (IWDM)*, pages 151–157.
- Bernard, S., Muller, S., Peters, G. et Iordache, R. (2006). Microcalcification detection on simple back-projection reconstructed slices using wavelets. *Computer Assisted Radiology and Surgery (CARS)*, 1:84–86.
- Bernard, S., Muller, S., Peters, G. et Iordache, R. (2007). Fast microcalcification detection on digital breast tomosynthesis datasets. *SPIE Symposium on Medical Imaging*, volume 6514, page 65141X.
- Beutel, J., Kundel, H. L. et Van Metter, R. L. (2000). *Handbook of medical imaging. Physics and psychophysics*, volume 1. SPIE Press, Washington, USA.
- Blessing, M., Godfrey, D., Lohr, F. et Yin, F. (2006). Analysis of the point spread function of isocentric digital tomosynthesis (DTS). *Medical Physics*, 33(6):2266–2266.
- Bloch, I. et Maitre, H. (1995). Fuzzy mathematical morphologies : a comparative study. *Pattern Recognition*, 28:1341–1387.
- Boone, J., Nelson, T., Lindfors, K. et Seibert, J. (2001). Dedicated breast CT : radiation dose and image quality evaluation. *Radiology*, 221:657–667.
- Bornefalk, H. (2005). Use of quadrature filters for detection of stellate lesions in mammograms. *Scandinavian Conference on Image Analysis (SCIA)*, pages 649–658.
- Bothorel, S. (1996). *Analyse d'images par arbre de décision flou*. Thèse de doctorat, Université Paris VI.
- Bourrelly, C. et Muller, S. (1989). Detection of microcalcifications in mammographic images. *NATO ASI Series in Computer and System Sciences*, volume 68, pages 325–328.
- Braccialarghe, D. et Kaufmann, G. H. (1996). Contrast enhancement of mammographic features : a comparison of four methods. *Optical Engineering*, 35(1):76–80.
- Braga-Neto, U. et Goutsias, J. (2003a). A multiscale approach to connectivity. *Computer Vision and Image Understanding*, 89(1):70–107.
- Braga-Neto, U. et Goutsias, J. (2003b). A Theoretical Tour of Connectivity in Image Processing and Analysis. *Journal of Mathematical Imaging and Vision*, 19(1):5–31.
- Breen, E. et Jones, R. (1996). Attribute openings, thinnings, and granulometries. *Computer Vision and Image Understanding*, 64(3):377–389.
- Bruynooghe, M. (2006). Mammographic mass detection using unsupervised clustering in synergy with a parcimonious supervised rule-based classifier. *International Workshop on Digital Mammography (IWDM)*, pages 68–75.
- Brzakovic, D., Luo, X. et Brzakovic, P. (1990). An approach to automated detection of tumors in mammograms. *IEEE Transactions on Medical Imaging*, 9(3):233–241.
- Bustince, H., Barrenechea, E., Pagola, M. et Orduna, R. (2007). Construction of interval type 2 fuzzy images to represent images in grayscale. False edges. *FUZZ-IEEE*, pages 1–6.
- Campanini, R., Dongiovanni, D., Iampieri, E., Lanconelli, N., Masotti, M., Palermo, G., Riccardi, A. et Roffilli, M. (2004). A novel featureless approach to mass detection in digital mammograms based on support vector machines. *Physics in Medicine and Biology*, 49(6):961–975.
- Canny, J. (1986). A computational approach to edge detection. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 8(6):679–698.

- Cao, A., Song, Q., Yang, X., Liu, S. et Guo, C. (2004). Mammographic mass detection by vicinal support vector machine. *IEEE International Joint Conference on Neural Networks*, 3:1953–1958.
- Chan, H.-P., Vyborny, C. J., MacMahon, H., Metz, C. E., Doi, K. et Sickles, E. A. (1987). Digital mammography. ROC studies of the effects of pixel size and unsharp-mask filtering on the detection of subtle microcalcifications. *Investigative Radiology*, 22(7):581–589.
- Chan, H.-P., Wei, J., Sahiner, B., Rafferty, E. A., Wu, T., Roubidoux, M. A., Moore, R. H., Kopans, D. B., Hadjiiski, L. M. et Helvie, M. A. (2004). Computerized detection of masses on digital tomosynthesis mammograms - a preliminary study. *International Workshop on Digital Mammography (IWDM)*.
- Chan, H.-P., Wei, J., Sahiner, B., Rafferty, E. A., Wu, T., Roubidoux, M. A., Moore, R. H., Kopans, D. B., Hadjiiski, L. M. et Helvie, M. A. (2005). Computer-aided detection system for breast masses on digital tomosynthesis mammograms : Preliminary experience. *Radiology*, 237(3):1075–1080.
- Chan, H.-P., Wei, J., Zhang, Y., Helvie, M. A., Moore, R. H., Sahiner, B., Hadjiiski, L. et Kopans, D. B. (2008a). Computer-aided detection of masses in digital tomosynthesis mammography : Comparison of three approaches. *Medical Physics*, 35(9):4087–4095.
- Chan, H.-P., Wei, J., Zhang, Y., Moore, R. H., Kopans, D. B., Hadjiiski, L., Sahiner, B., Roubidoux, M. A. et Helvie, M. A. (2007). Computer-aided detection of masses in digital tomosynthesis mammography : combination of 3D and 2D detection information. *SPIE Symposium on Medical Imaging*, volume 6514, page 651416.
- Chan, H.-P., Wei, J., Zhang, Y., Sahiner, B., Hadjiiski, L. et Helvie, M. A. (2008b). Detection of masses in digital breast tomosynthesis mammography : Effects of the number of projection views and dose. *International Workshop on Digital Mammography (IWDM)*, pages 279–285, Berlin, Heidelberg. Springer-Verlag.
- Chan, T. et Vese, L. (2001). Active contours without edges. *IEEE Transactions on Image Processing*, 10(2):266–277.
- Chang, C., Nesbit, D., Fisher, D., Fritz, S., 3rd S., D., Templeton, A., Lin, F. et Jewell, W. (1982). Computed tomographic mammography using a conventional body scanner. *American Journal of Roentgenology*, 138:553–558.
- Chang, C., Sibala, J., Fritz, S., Dwyer, S. et Templeton, A. (1979). Specific value of computed tomographic breast scanner (CT / M) in diagnosis of breast diseases. *Radiology*, 132:647–652.
- Chang, C.-M. et Laine, A. (1997). Enhancement of mammograms from oriented information. *IEEE International Conference on Image Processing*, volume 3, pages 524–527.
- Chassery, J.-M. et Montanvert, A. (1991). *Géométrie Discrètes en Analyse d'Images*. Hermes, Paris.
- Chen, C. H. et Lee, G. G. (1997). On digital mammogram segmentation and microcalcification detection using multiresolution wavelet analysis. *Graphical Models and Image Processing*, 59(5):349–364.
- Cheng, H., Shi, X., Min, R., Hu, L., Cai, X. et Du, H. (2006). Approaches for automated detection and classification of masses in mammograms. *Pattern Recognition*, 39(4):646 – 668. Graph-based Representations.
- Cheng, H. D., Cai, X., Chen, X., Hu, L. et Lou, X. (2003). Computer-aided detection and classification of microcalcifications in mammograms : a survey. *Pattern Recognition*, 36(12):2967 – 2991.
- Cheng, H.-D., Lui, Y. M. et Freimanis, R. (1998). A novel approach to microcalcification detection using fuzzy logic technique. *IEEE Transactions on Medical Imaging*, 17(3):442–450.
- Cheng, S., Chan, H., Helvie, M., Goodsitt, M., Adler, D. et St. Clair, D. (1994). Classification of mass and non-mass regions on mammograms using artificial neural networks. *Journal of Imaging Science and Technology*, 38:598–603.

- Chitre, Y., Dhawan, A. et Moskowitz, M. (1994). Artificial neural network based classification of mammographic microcalcifications using image structure and cluster features. *IEEE Engineering in Medicine and Biology Society*, 1:592–593.
- Claus, B. et Eberhard, J. (2002). A new method for 3D reconstruction in digital tomosynthesis. *SPIE Symposium on Medical Imaging*, volume 4684.
- Claus, B., Eberhard, J., Thomas, J., Galbo, C., Pakenas, W. et Muller, S. (2002). Preference study of reconstructed image quality in mammographic tomosynthesis. *International Workshop on Digital Mammography (IWDM)*.
- Cortes, C. et Vapnik, V. (1995). Support-vector networks. *Machine Learning*, pages 273–297.
- Couto, P., Pagola, M., Bustince, H., Barrenechea, E. et Melo-Pinto, P. (2008). Image segmentation using A-IFSs. *International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU)*, pages 1620–1627, Malaga, Spain.
- Cover, T. et Hart, P. (1967). Nearest neighbor pattern classification. *IEEE Transactions on Information Theory*, 13(1):21–27.
- Crespo, J., Schafer, R. W., Serra, J., Gratin, C. et Meyer, F. (1997). The flat zone approach : a general low-level region merging segmentation method. *Signal Processing*, 62(1):37–60.
- Cristianini, N. et Shawe-Taylor, J. (2000). *An Introduction to Support Vector Machines and other kernel-based learning methods*. Cambridge University Press.
- Daubechies, I. (1988). Orthonormal bases of compactly supported wavelets. *Communications on Pure and Applied Mathematics*, 41:909–996.
- De Luca, A. et Termini, S. (1972). A definition of non-probabilistic entropy in the setting of fuzzy sets theory. *Information and Control*, 20:301–312.
- Dengler, J., Behrens, S. et Desaga, J. F. (1993). Segmentation of microcalcification in mammograms. *IEEE Transactions on Medical Imaging*, 12(4):634–642.
- Desolneux, A., Moisan, L. et Morel, J. (2000). Meaningful alignments. *International Journal of Computer Vision*, 40(1):7–23.
- Desolneux, A., Moisan, L. et Morel, J. (2003). A grouping principle and four applications. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 25(4):508–513.
- Desolneux, A., Moisan, L. et Morel, J.-M. (2001). Edge detection by helmholtz principle. *Journal of Mathematical Imaging and Vision*, 14(3):271–284.
- Dhawan, A., Chitre, Y. et Kaiser Bonasso, C. (1996). Analysis of mammographic microcalcifications using gray-level image structure features. *IEEE Transactions on Medical Imaging*, 15(3):246–259.
- Dobbins III, J. T. et Godfrey, D. J. (2003). Digital x-ray tomosynthesis : current state of the art and clinical potential. *Physics in Medicine and Biology*, 48(19):R65–R106.
- Domínguez, A. R. et Nandi, A. K. (2007). Improved dynamic-programming-based algorithms for segmentation of masses in mammograms. *Medical Physics*, 34(11):4256–4269.
- Domínguez, A. R. et Nandi, A. K. (2009). Toward breast cancer diagnosis based on automated segmentation of masses in mammograms. *Pattern Recognition*, 42(6):1138 – 1148. Digital Image Processing and Pattern Recognition Techniques for the Detection of Cancer.
- Donoho, D. L. (1995). De-noising by soft-thresholding. *IEEE Transactions on Information Theory*, 41(3):613–627.

- Donoho, D. L. et Johnstone, I. M. (1995). Adapting to unknown smoothness via wavelet shrinkage. *Journal of the American Statistical Association*, 90(432):1200–1224.
- Donoho, D. L., Johnstone, I. M., Kerkycharian, G. et Picard, D. (1995). Wavelet shrinkage : asymptopia ? (with discussion). *Journal of the Royal Statistical Society : Series B (Statistical Methodology)*, 57:301–369.
- Dubois, D. et Prade, H. (1980). *Fuzzy Sets and Systems : Theory and Applications*. Academic Press, New York.
- Duffy, S., Tabar, I. et Vitak, B. (2003). The relative contribution of screen-detected in situ and invasive breast carcinomas in reducing mortality from the disease. *European Journal of Cancer*, 39.
- Eltonsy, N., Tourassi, G. et Elmaghraby, A. (2007). A concentric morphology model for the detection of masses in mammography. *IEEE Transactions on Medical Imaging*, 26(6):880–889.
- Floyd Jr., C. E., Lo, J. Y., Yun, A. J., Sullivan, D. C. et Kornguth, P. J. (1994). Prediction of breast cancer malignancy using an artificial neural network. *Cancer*, 74(11):2944–2948.
- Fogel, D. B., III, E. C. W., Boughton, E. M. et Porto, V. W. (1998). Evolving artificial neural networks for screening features from mammograms. *Artificial Intelligence in Medicine*, 14(3):317–326.
- Freer, T. W. et Ullissey, M. J. (2001). Screening Mammography with Computer-aided Detection : Prospective Study of 12,860 Patients in a Community Breast Center. *Radiology*, 220(3):781–786.
- Galloway, M. M. (1975). Texture analysis using gray level run lengths. *Computer Graphics and Image Processing*, 4(2):172–179.
- Gavrielides, M. A., Lo, J. Y. et Carey E. Floyd, J. (2002). Parameter optimization of a computer-aided diagnosis scheme for the segmentation of microcalcification clusters in mammograms. *Medical Physics*, 29(4):475–483.
- Gennaro, G., Baldan, E., Bezzon, E., Grassa, M. L., Pescarini, L. et di Maggio, C. (2008). Clinical performance of digital breast tomosynthesis versus full-field digital mammography : Preliminary results. *International Workshop on Digital Mammography (IWDM)*, pages 477–482, Berlin, Heidelberg. Springer-Verlag.
- Géraud, T., Palma, G. et Van-Vliet, N. (2004). Fast Color Image Segmentation Based on Levellings in Feature Space. *International Conference on Computer Vision and Graphics (ICCVG)*, volume 32, pages 800–813.
- Gilbert, F. J., Astley, S. M., McGee, M. A., Gillan, M. G. C., Boggis, C. R. M., Griffiths, P. M. et Duffy, S. W. (2006). Single reading with computer-aided detection and double reading of screening mammograms in the united kingdom national breast screening program. *Radiology*, 241(1):47–53.
- Gilbert, P. (1972). Iterative methods for the three-dimensional reconstruction of an object from projections. *Journal of Theoretical Biology*, 36(1):105 – 117.
- Gisvold, J., Reese, D. et Karsell, P. (1979). Computed tomographic mammography (CTM). *American Journal of Roentgenology*, 133:1143–1149.
- Gordon, R., Bender, R. et Herman, G. T. (1970). Algebraic reconstruction techniques (ART) for three-dimensional electron microscopy and X-ray photography. *Journal of Theoretical Biology*, 29(3):471–81.
- Gordon, R. et Herman, G. T. (1974). Three-dimensional reconstruction from projections : a review of algorithms. *International Review of Cytology*, 38:111–151.
- Gout, C., Guyader, C. L. et Vese, L. (2005). Segmentation under geometrical conditions using geodesic active contours and interpolation using level set methods. *Numerical Algorithms*, 39:155–173(19).

- Grimaud, M. (1991). *La géodésie numérique en morphologie mathématique. Application à la détection automatique des microcalcifications*. Thèse de doctorat, Ecole nationale supérieure des mines de Paris.
- Grosjean, B. (2007). *Lesion detectability in Digital Mammography : Impact of Texture*. Thèse de doctorat, Ecole centrale des arts et manufactures "Ecole Centrale Paris", France.
- Gulato, D., Rangayyan, R., Carnielli, W., Zuffo, J. et Desautels, J. (1998). Segmentation of breast tumors in mammograms by fuzzy region growing. *IEEE Engineering in Medicine and Biology Society*, 2:1002–1005.
- Hara, T., Makita, T., Matsubara, T., Fujita, H., Inenaga, Y., Endo, T. et Iwase, T. (2006). Automated detection method for architectural distortion with spiculation based on distribution assessment of mammary gland on mammogram. *International Workshop on Digital Mammography (IWDM)*, pages 370–375.
- Haralick, R. M., Shanmugam, K. et Dinstein, I. (1973). Textural features for image classification. *IEEE Transactions on Systems, Man and Cybernetics*, 3(6):610–621.
- Hassanien, A. et Ali, J. (2004). Digital mammogram segmentation algorithm using pulse coupled neural networks. *International Conference on Image and Graphics*, pages 92–95.
- Hatanaka, Y., Hara, T., Fujita, H., Kasai, S., Endo, T. et Iwase, T. (2001). Development of an automated method for detecting mammographic masses with a partial loss of region. *IEEE Transactions on Medical Imaging*, 20(12):1209–1214.
- Heath, M. D. et Bowyer, K. W. (2000). Mass detection by relative image intensity. *International Workshop on Digital Mammography (IWDM)*. Medical Physics Publishing (Madison, WI).
- Heijmans, H. (1997). Connected morphological operators and filters for binary images. *IEEE International Conference on Image Processing*, volume 2, pages 211–214.
- Holland, J. H. (1992). *Adaptation in Natural and Artificial Systems : An Introductory Analysis with Applications to Biology, Control, and Artificial Intelligence*. The MIT Press.
- Holland, R., Mravunac, M., Hendriks, J. H. C. L. et Bekker, B. V. (1982). So-called interval cancers of the breast : Pathologic and radiologic analysis of sixty-four cases. *Cancer*, 49:2527–2533.
- Huo, Z., Giger, M. et Vyborny, C. (2001). Computerized analysis of multiple-mammographic views : potential usefulness of special view mammograms in computer-aided diagnosis. *IEEE Transactions on Medical Imaging*, 20(12):1285–1292.
- Huo, Z., Giger, M. L., Vyborny, C. J., Wolverton, D. E., Schmidt, R. A. et Doi, K. (1998). Automated computerized classification of malignant and benign masses on digitized mammograms. *Academic Radiology*, 5(3):155–168.
- Janikow, C. (1998). Fuzzy decision trees : issues and methods. *IEEE Transactions on Systems, Man and Cybernetics*, 28(1):1–14.
- Jansen, M. (2001). *Noise Reduction by Wavelet Thresholding*. Springer.
- Jansen, M., Malfait, M. et Bultheel, A. (1997). Generalized cross validation for wavelet thresholding. *Signal Processing*, 56:33–44.
- Jeunehomme, F. (2005). *Mammographie numérique avec injection de produit de contraste*. Thèse de doctorat, université Paris XI Orsay.
- Jiang, M. et Wang, G. (2003). Convergence of the simultaneous algebraic reconstruction technique (sart). *IEEE Transactions on Image Processing*, 12(8):957–961.
- Jiang, Y., Nishikawa, R., Wolverton, D., Metz, C., Schmidt, R. et Doi, K. (1997). Computerized classification of malignant and benign clustered microcalcifications in mammograms. *IEEE Engineering in Medicine and Biology Society*, 2:521–523.

- John, V. et Ewen, K. (1989). CT scanning of the breast in problem cases. *Strahlentherapie und Onkologie*, 165:657–662.
- Kallergi, M., Woods, K., Clarke, L., Qian, W. et Clark, R. (1992). Image segmentation in digital mammography : comparison of local thresholding and region growing algorithms. *Computerized medical imaging and graphics*, 16(5):231–331.
- Kanzaki, K. (1992). *The Use of Morphology and Fuzzy Set Theory in FLIR Target Segmentation and Classification*. Thèse de doctorat, Polytechnic University of New York.
- Karssemeijer, N. (1993). Adaptive noise equalization and image analysis in mammography. *International Conference on Information Processing in Medical Imaging (IPMI)*, pages 472–486, London, UK. Springer-Verlag.
- Karssemeijer, N., Otten, J. D. M., Verbeek, A. L. M., Groenewoud, J. H., de Koning, H. J., Hendriks, J. H. C. L. et Holland, R. (2003). Computer-aided detection versus independent double reading of masses on mammograms. *Radiology*, 227(1):192–200.
- Karssemeijer, N. et te Brake, G. M. (1996). Detection of stellate distortions in mammograms. *IEEE Transactions on Medical Imaging*, 5(5):611–619.
- Karssemeijer, N. et Te Brake, G. M. (1997). Method and apparatus for automated detection of masses in digital mammograms. United States Patent 6301378.
- Kass, M., Witkin, A. et Terzopoulos, D. (1988). Snakes : Active contour models. *V1(4)*:321–331.
- Kegelmeyer, W., Pruneda, J., Bourland, P. et al. (1994). Computer-aided mammographic screening for spiculated lesions. *Radiology*, 191:331–337.
- Kerlikowske, K., Grady, D., Rubin, S., Sandrock, C. et Ernster, V. (1995). Efficacy of screening mammography. a meta-analysis. *Journal of the American Medical Association (JAMA)*, 274(5):381–382.
- Kilday, J., Palmieri, F. et Fox, M. (1993). Classifying mammographic lesions using computerized image analysis. *IEEE Transactions on Medical Imaging*, 12(4):664–669.
- Kim, H. et Kim, W. (2005). Automatic detection of spiculated masses using fractal analysis in digital mammography. *Computer Analysis of Images and Patterns (CAIP)*, pages 256–263.
- Kim, J. K. et Park, H. W. (1999). Statistical textural features for detection of microcalcifications in digitized mammograms. *IEEE Transactions on Medical Imaging*, 18(3):231–238.
- Klein, J. (1985). *Conception et réalisation d'une unité logique pour l'analyse quantitative d'images*. Thèse de doctorat, Nancy University, France.
- Klette, R. et Rosenfeld, A. (2004). *Digital Geometry : Geometric Methods for Digital Picture Analysis*. Morgan Kaufmann.
- Koenderink, J. J. et van Doorn, A. J. (1992). Generic neighborhood operators. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 14(6):597–605.
- Kramer, D. et Aghdasi, F. (1999). Texture analysis techniques for the classification of microcalcifications in digitised mammograms. *IEEE AFRICON*, 1:395–400.
- Kupinski, M. et Giger, M. (1997). Feature selection and classifiers for the computerized detection of mass lesions in digital mammography. *International Conference on Neural Networks*, 4:2460–2463.
- Kupinski, M. A. et Giger, M. L. (1998). Automated seeded lesion segmentation on digital mammograms. *IEEE Transactions on Medical Imaging*, 17(4):510–517.
- Kupinski, M. A. et Giger, M. L. (1999). Feature selection with limited datasets. *Medical Physics*, 26(10): 2176–2182.

- Kwan, A., Shah, N., Burkett, G., Seibert, J., Lindfors, K., Nelson, T. et Boone, J. (2004). Progress in the development of a dedicated breast CT scanner. *SPIE Symposium on Medical Imaging*, volume 5368, pages 304–310.
- Lai, S.-M., Li, X. et Biscof, W. (1989). On techniques for detecting circumscribed masses in mammograms. *IEEE Transactions on Medical Imaging*, 8(4):377–386.
- Lau, B. A., Reiser, I. S. et Nishikawa, R. M. (2008). Microcalcification detectability in tomosynthesis. *SPIE Symposium on Medical Imaging*, volume 6913, page 69134L.
- Lau, C. (1991). *Neural Networks : Theoretical Foundations and Analysis*. IEEE Press, Piscataway, NJ, USA.
- Leclerc, Y. (1989). Constructing simple stable descriptions for image partitioning. *International Journal of Computer Vision*, 3(1):73–102.
- Lee, K.-M., Lee, K.-M., Lee, J.-H. et Lee-Kwang, H. (1999). A fuzzy decision tree induction method for fuzzy data. volume 1 de *Fuzzy Systems Conference*, pages 16–21.
- Lee, Y., Park, J. et Park, H. (2000). Mammographic mass detection by adaptive thresholding and region growing. *International Journal of Imaging Systems and Technology*, 11(5):340–346.
- Li, L., Mao, F., Qian, W. et Clarke, L. (1997). Wavelet transform for directional feature extraction in medical imaging. *IEEE International Conference on Image Processing*, volume 3, page 500, Los Alamitos, CA, USA. IEEE Computer Society.
- Liu, S., Babbs, C. et Delp, E. (2001). Multiresolution detection of spiculated lesions in digital mammograms. *IEEE Transactions on Medical Imaging*, 10(6):874–884.
- Llobet, R., Paredes, R. et Pérez-Cortes, J. C. (2005). Comparison of feature extraction methods for breast cancer detection. *Iberian Conference on Pattern Recognition and Image Analysis (PRIA)*, pages 495–502.
- Lopez-Diaz, M. et Gil, M. A. (1997). Constructive definitions of fuzzy random variables. *Statistics & Probability Letters*, 36(2):135–143.
- Mallat, S. et Zhong, S. (1992). Characterization of signals from multiscale edges. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 14(7):710–732.
- Mallat, S. G. (1989). A theory for multiresolution signal decomposition : The wavelet representation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 11(7).
- Martinez, V. G., Gamo, D. M., Rios, J. et Vilarrasa, A. (1999). Iterative method for automatic detection of masses in digital mammograms for computer-aided diagnosis. *SPIE Symposium on Medical Imaging*, volume 3661, pages 1086–1093.
- Matheron, G. (1974). *Random sets and integral geometry*. Wiley New York,.
- Matsubara, T., Fukuoka, D., Yagi, N., Hara, T., Fujita, H., Inenaga, Y., Kasai, S., Kano, A., Endo, T. et Iwase, T. (2005). Detection method for architectural distortion based on analysis of structure of mammary gland on mammograms. *Computer Assisted Radiology and Surgery (CARS)*, 1281:1036–1040.
- Meijster, A. et Wilkinson, M. H. F. (2002). A comparison of algorithms for connected set openings and closings. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 24(4):484–494.
- Meyer, F. (1998a). From connected operators to levelings. *International Symposium on Mathematical Morphology and its applications to image and signal processing (ISMM)*, pages 191–198, Norwell, MA, USA. Kluwer Academic Publishers.
- Meyer, F. (1998b). The levelings. *International Symposium on Mathematical Morphology and its applications to image and signal processing (ISMM)*, pages 199–206, Norwell, MA, USA. Kluwer Academic Publishers.

- Meyer, F. (2005). *Space, Structure and Randomness*, volume 183 de *Lecture Notes in Statistics*, chapitre Morphological segmentation revisited, pages 315–347.
- Meyer, F. et Maragos, P. (1999). Morphological scale-space representation with levelings. *International Conference on Scale-Space Theories in Computer Vision*, pages 187–198, London, UK. Springer-Verlag.
- Meyer, F. et Maragos, P. (2000). Nonlinear scale-space representation with morphological levelings. *Journal of Visual Communication and Image Representation*, 11:200–0.
- Moisan, L. (2003). Modèles continus, numériques et statistiques pour l’analyse d’ images. Habilitation à Diriger des Recherches, spécialité : Mathématiques.
- Moisan, L. et Stival, B. (2004). A probabilistic criterion to detect rigid point matches between two images and estimate the fundamental matrix. *International Journal of Computer Vision*, 57(3):201–218.
- Mudigonda, N., Rangayyan, R. et Leo Desautels, J. (2001). Detection of breast masses in mammograms by density slicing and texture flow-field analysis. *IEEE Transactions on Medical Imaging*, 20(12):1215–1227.
- Muller, J., Waes, P. V. et Koehler, P. (1983). Computed tomography of breast lesions : comparison with x-ray mammography. *Journal of Computed Assisted Tomography*, 7:650–654.
- Mumford, D., Shah, J. et Center for Intelligent Control Systems (US) (1988). *Optimal Approximations by Piecewise Smooth Functions and Associated Variational Problems*. Center for Intelligent Control Systems.
- Nachtegaal, M., Mélangé, T. et Kerre, E. E. (2007). The possibilities of fuzzy logic in image processing. *International Conference on Pattern Recognition and Machine Intelligence*, volume 4815 de *Lecture Notes in Computer Science*, pages 198–208. Springer.
- Najman, L. et Couprie, M. (2006). Building the Component Tree in Quasi-Linear Time. *IEEE Transactions on Image Processing*, 15(11):3531–3539.
- National Cancer Institute Consensus Development Panel (1998). Screening mammography for women ages 40–49.
- Nempont, O., Atif, J., Angelini, E. et Bloch, I. (2008). A New Fuzzy Connectivity Class. Application to Structural Recognition in Images. *Discrete Geometry for Computer Imagery*, pages 446–557, Lyon.
- Niklason, L., Christian, B., Niklason, L., Kopans, D., Castleberry, D., Opsahl-Ong, B., Landberg, C., Slanetz, P., Giardino, A., Moore, R., Albagli, D., DeJule, M., Fitzgerald, P., Fobare, D., Giambattista, B., Kwasnick, R., Liu, J., Lubowski, S., Possin, G., Richotte, J., Wei, C. et Wirth, R. (1997). Digital tomosynthesis in breast imaging. *Radiology*, 205:399–406.
- Osher, S. et Sethian, J. (1988). Fronts propagating with curvature-dependent speed : Algorithms based on Hamilton-Jacobi formulations. *Journal of Computational Physics*, 79:12–49.
- Osher, S. J. et Fedkiw, R. (2002). *Level Set Methods and Dynamic Implicit Surfaces*. Springer.
- Palma, G., Bloch, I. et Muller, S. (2008a). Fuzzy connected filters for fuzzy gray scale images. *International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU)*, pages 667–674, Malaga, Spain.
- Palma, G., Bloch, I. et Muller, S. (2010a). Fast fuzzy connected filter implementation using max-tree updates. *Fuzzy Systems and Systems*, 161(1):118–146.
- Palma, G., Bloch, I. et Muller, S. (2010b). Spiculated lesions and architectural distortions detection in digital breast tomosynthesis datasets. *International Workshop on Digital Mammography (IWDM)*, Girona, Spain.

- Palma, G., Bloch, I., Muller, S. et Iordache, R. (2009a). Fuzzifying images using fuzzy wavelet denoising. *FUZZ-IEEE*, pages 135–140, Jeju, Korea.
- Palma, G., Muller, S., Bloch, I. et Iordache, R. (2009b). Convergence areas detection in digital breast tomosynthesis volumes using a contrario modeling. *SPIE Symposium on Medical Imaging*, Lake Buena Vista, FL, USA.
- Palma, G., Muller, S., Bloch, I. et Iordache, R. (2009c). Fast detection of convergence areas in digital breast tomosynthesis. *IEEE International Symposium on Biomedical Imaging (ISBI)*, pages 847–850, Boston, MA, USA.
- Palma, G., Nempont, O., Bloch, I. et Muller, S. (2008b). Extraction de "zones plates floues" dans des images de quantités floues. *Rencontre Francophones sur la Logique Floue et ses Applications*, pages 364–371, Lens, France.
- Palma, G., Peters, G., Muller, S. et Bloch, I. (2008c). Masses classification using fuzzy active contours and fuzzy decision trees. *SPIE Symposium on Medical Imaging*, numéro 6915, San Diego, CA, USA.
- Pearl, J. (1988). *Probabilistic Reasoning in Intelligent Systems : Networks of Plausible Inference*. Morgan Kaufmann.
- Peters, G. (2007). *Computer-aided Detection for Digital Breast Tomosynthesis*. Thèse de doctorat, Ecole Nationale Supérieure des Télécommunications.
- Peters, G., Muller, S., Bernard, S. et Bloch, I. (2006a). *Soft Computing and Image Processing : Recent Advances*, chapitre Wavelets and Fuzzy Contours in 3D-CAD for Digital Breast Tomosynthesis, pages 296–323. Springer-Verlag.
- Peters, G., Muller, S., Bernard, S., Iordache, R. et Bloch, I. (2006b). Reconstruction-independent 3D CAD for mass detection in digital breast tomosynthesis using fuzzy particles. *SPIE Symposium on Medical Imaging*, volume 6144, page 61441Z.
- Peters, G., Muller, S., Bernard, S., Iordache, R., Wheeler, F. et Bloch, I. (2005). Reconstruction-independent 3D CAD for calcification detection in digital breast tomosynthesis using fuzzy particles. *Iberoamerican Congress on Pattern Recognition (CIARP)*, pages 400–408.
- Peters, G., Muller, S., Grosjean, B., Bernard, S. et Bloch, I. (2007). A hybrid active contour model for mass detection in digital breast tomosynthesis. *SPIE Symposium on Medical Imaging*, volume 6514-1V, pages 1–11.
- Pfisterer, R. S. et Aghdasi, F. (1999). Hexagonal wavelets for the detection of masses in digitized mammograms. volume 3813, pages 966–977. SPIE.
- Pisano, E. D., Gatsonis, C., Hendrick, E., Yaffe, M., Baum, J. K., Acharyya, S., Conant, E. F., Fajardo, L. L., Bassett, L., D'Orsi, C., Jong, R., Rebner, M. et the Digital Mammographic Imaging Screening Trial (DMIST) Investigators Group (2005). Diagnostic performance of digital versus film mammography for breast-cancer screening. *N Engl J Med*, 353(17):1773–1783.
- Puong, S. (2008). *Imagerie du sein multispectrale avec injection de produit de contraste*. Thèse de doctorat, université Paris XI Orsay.
- Radon, J. (1917). über die bestimmung von funktionen durch ihre integralwerte längs gewisser mannigfaltigkeiten. *Ber. Ver. Sachs. Akad. Wiss. Leipzig, MathPhys. Kl.*, 69:262–277.
- Rangayyan, R., El-Faramawy, N., Desautels, J. et Alim, O. (1997). Measures of acutance and shape for classification of breast tumors. *IEEE Transactions on Medical Imaging*, 16(6):799–810.
- Rankin, S. (2000). MRI of the breast. *British Journal of Radiology*, 73(872):806–818.
- Rao, A. et Schunck, B. (1989). Computing oriented texture fields. *IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 61–68.

- Reiser, I., Nishikawa, R., Giger, M., Kopans, D., Rafferty, E., Wu, T. et Moore, R. (2005). A multi-scale 3D radial gradient filter for computerized mass detection in digital tomosynthesis breast images. *Computer Assisted Radiology and Surgery (CARS)*, 1281:1058–1062.
- Reiser, I., Nishikawa, R. M., Giger, M. L., Wu, T., Rafferty, E. A., Moore, R. et Kopans, D. B. (2006). Computerized mass detection for digital breast tomosynthesis directly from the projection images. *Medical Physics*, 33(2):482–491.
- Remontet, L., Esteve, J., Bouvier, A.-M., Grosclaude, P., Launoy, G., Menegoz, F., Exbrayat, C., Tretare, B., Carli, P.-M., Guizard, A.-V., Troussard, X., Bercelli, P., Colonna, M., Halna, J.-M., Hedelin, G., Mace-Lescâh, J., Peng, J., Buemi, A., Velten, M., Jouglu, E., Arveux, P., Le Bodic, L., Michel, E., Sauvage, M., Schvartz, C. et Faivre, J. (2003). Cancer incidence and mortality in France over the period 1978-2000. *Revue d'épidémiologie et de santé publique*, 51:3–30.
- Ren, B., Ruth, C., Stein, J., Smith, A., Shaw, I. et Jing, Z. (2005). Design and performance of the prototype full field breast tomosynthesis system with selenium based flat panel detector. *SPIE Symposium on Medical Imaging*, volume 5745, pages 550–561.
- Rick, A. (1999). *Représentation de la variabilité dans le traitement d'images flou*. Thèse de doctorat, université Paris VI.
- Rifqi, M. (1996). *Mesures de comparaison, typicalité et classification d'objets flous : théorie et pratique*. Thèse de doctorat, LAFORIA, Institut Blaise Pascal.
- Rosenfeld, A. (1979). Fuzzy Digital Topology. *Information and Control*, 40:76–87.
- Rosenfeld, A. (1984). The fuzzy geometry of image subsets. *Pattern Recognition Letters*, 2:311–317.
- Rosenfeld, A. (1992). Fuzzy geometry : An overview. *IEEE International Conference on Fuzzy Systems*, pages 113–117.
- Rosenfeld, A. (1998). Fuzzy geometry : an updated overview. *Information Sciences*, 110(3-4):127–133.
- Rosenfeld, A. et Haber, S. (1985). The perimeter of a fuzzy subset. *Pattern Recognition*, 18:125–130.
- Rumelhart, D., Hinton, G. et Williams, R. (1986). Learning representation by back-propagating errors. *Nature*, 323(9):533–536.
- S. Bothorel, B. Bouchon Meunier, S. M. (1997). A fuzzy logic based approach for semiological analysis of microcalcifications in mammographic images. *International Journal of Intelligent Systems*, 12(11-12):819–848.
- Safavian, S. et Landgrebe, D. (1991). A survey of decision tree classifier methodology. *IEEE Transactions on Systems, Man and Cybernetics*, 21(3):660–674.
- Sahiner, B., Chan, H.-P., Petrick, N., Helvie, M. A. et Goodsitt, M. M. (1998). Computerized characterization of masses on mammograms : The rubber band straightening transform and texture analysis. *Medical Physics*, 25(4):516–526.
- Sahiner, B., Chan, H.-P., Petrick, N., Helvie, M. A. et Hadjiiski, L. M. (2001). Improvement of mammographic mass characterization using spiculation measures and morphological features. *Medical Physics*, 28:1455–1465.
- Sahiner, B., Chan, H.-P., Petrick, N., Wei, D., Helvie, M., Adler, D. et Goodsitt, M. (1996a). Classification of mass and normal breast tissue : a convolution neural network classifier with spatial domain and texture images. *IEEE Transactions on Medical Imaging*, 15(5):598–610.
- Sahiner, B., Chan, H.-P., Wei, D., Petrick, N., Helvie, M. A., Adler, D. D. et Goodsitt, M. M. (1996b). Image feature selection by a genetic algorithm : Application to classification of mass and normal breast tissue. *Medical Physics*, 23(10):1671–1684.

- Sahiner, B., Petrick, N., Chan, H.-P., Hadjiiski, L., Paramagul, C., Helvie, M. et Gurcan, M. (2001). Computer-aided characterization of mammographic masses : accuracy of mass segmentation and its effects on characterization. *IEEE Transactions on Medical Imaging*, 20(12):1275–1284.
- Salembier, P., Oliveras, A. et Garrido, L. (1998). Antiextensive connected operators for image and sequence processing. *IEEE Transactions on Image Processing*, 7(4):555–570.
- Salembier, P. et Serra, J. (1995). Flat zones filtering, connected operators, and filters by reconstruction. *IEEE Transactions on Image Processing*, 4(8):1153–1160.
- Sameti, M. et Ward, R. (1996). A fuzzy segmentation algorithm for mammogram partitioning. *International Workshop on Digital Mammography (IWDM)*, pages 471–474.
- Sampat, M. et Bovik, A. (2003). Detection of spiculated lesions in mammograms. *IEEE Engineering in Medicine and Biology Society*, 1:810–813.
- Saslow, D., Boetes, C., Burke, W., Harms, S., Leach, M. O., Lehman, C. D., Morris, E., Pisano, E., Schnall, M., Sener, S., Smith, R. A., Warner, E., Yaffe, M., Andrews, K. S. et Russell, C. A. (2007). American Cancer Society guidelines for breast screening with MRI as an adjunct to mammography. *CA : A Cancer Journal for Clinicians*, 57(2):75–89.
- Schulte, S., Huysmans, B., Pizurica, A., Kerre, E. et Philips, W. (2006a). A new fuzzy-based wavelet shrinkage image denoising technique. *Advanced Concepts for Intelligent Vision Systems*, pages 12–23.
- Schulte, S., Nachtgael, M., de Witte, V., van der Weken, D. et Kerre, E. (2006b). A fuzzy impulse noise detection and reduction method. *IEEE Transactions on Image Processing*, 15(5):1153–1162.
- Serra, J. (1982). *Image Analysis and Mathematical Morphology*. Academic Press, Inc., Orlando, FL, USA.
- Serra, J. (1998). Connectivity on complete lattices. *Journal of Mathematical Imaging and Vision*, 9(3):231–251.
- Serra, J. et Salembier, P. (1993). Connected operators and pyramids. *SPIE Image Algebra and Morphological Image Processing IV*, volume 2030, pages 65–76.
- Sibala, J. L., Chang, C. H. J., Lin, F. et Jewell, W. R. (1981). Computed tomographic mammography : Diagnosis of mammographically and clinically occult carcinoma of the breast. *Archives of Surgery*, 116(1):114–117.
- Singh, S. (2008). *Computer Aided Detection of Masses in Breast Tomosynthesis Imaging Using Information Theory Principles*. Thèse de doctorat, Duke University.
- Singh, S., Tourassi, G. D., Baker, J. A., Samei, E. et Lo, J. Y. (2008a). Automated breast mass detection in 3D reconstructed tomosynthesis volumes : A featureless approach. *Medical Physics*, 35(8):3626–3636.
- Singh, S., Tourassi, G. D., Chawla, A. S., Saunders, R. S., Samei, E. et Lo, J. Y. (2008b). Computer-aided detection of breast masses in tomosynthesis reconstructed volumes using information-theoretic similarity measures. *SPIE Symposium on Medical Imaging*, volume 6915, page 691505.
- Soille, P. (2008). Constrained connectivity for hierarchical image partitioning and simplification. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 30(7):1132–1145.
- Strickland, R. et Hahn, H. I. (1997). Wavelet transform methods for object detection and recovery. *IEEE Transactions on Medical Imaging*, 6(5):724–735.
- Tahmoush, D. et Samet, H. (2006a). Image similarity and asymmetry to improve computer-aided detection of breast cancer. *International Workshop on Digital Mammography (IWDM)*, pages 221–228.
- Tahmoush, D. et Samet, H. (2006b). Using image similarity and asymmetry to detect breast cancer. *SPIE Symposium on Medical Imaging*, volume 6144, pages 605–611.

- te Brake, G. M. et Karssemeijer, N. (1999). Single and multiscale detection of masses in digital mammograms. *IEEE Transactions on Medical Imaging*, 18(7):628–639.
- Tehami, S., Bigand, A. et Colot, O. (2007). Color image segmentation based on type-2 fuzzy sets and region merging. *Advanced Concepts for Intelligent Vision Systems*, volume 4678 de *Lecture Notes in Computer Science*, pages 943–954. Springer.
- Timp, S. et Karssemeijer, N. (2004). A new 2D segmentation method based on dynamic programming applied to computer aided detection in mammography. *Medical Physics*, 31(5):958–971.
- Tizhoosh, H. R. (2008). *Fuzzy Sets and Their Extensions : Representation, Aggregation and Models*, chapitre Type II Fuzzy Image Segmentation, pages 607–619. Springer-Verlag.
- Torre, V. et Poggio, T. A. (1986). On edge detection. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, PAMI-8(2):147–163.
- Tourassi, G. D., Vargas-Voracek, R., David M. Catarious, J. et Carey E. Floyd, J. (2003). Computer-assisted detection of mammographic masses : A template matching scheme based on mutual information. *Medical Physics*, 30(8):2123–2130.
- Vachier, C. (1995). *Extraction de caractéristiques, segmentation d'image et morphologie mathématique*. Thèse de doctorat, Ecole nationale supérieure des mines de Paris.
- Vachier, C. (2001). Morphological scale-space analysis and feature extraction. *IEEE International Conference on Image Processing*, pages III : 676–679.
- van Engeland, S. et Karssemeijer, N. (2007). Combining two mammographic projections in a computer aided mass detection method. *Medical Physics*, 34(3):898–905.
- van Engeland, S., Timp, S. et Karssemeijer, N. (2006). Finding corresponding regions of interest in medio-lateral oblique and craniocaudal mammographic views. *Medical Physics*, 33(9):3203–3212.
- Vapnik, V. N. (1995). *The nature of statistical learning theory*. Springer-Verlag New York, Inc.
- Veldkamp, W. J. et Karssemeijer, N. (1999). Improved method for detection of microcalcification clusters in digital mammograms. *SPIE Symposium on Medical Imaging*, volume 3661, pages 512–522.
- Verma, B. et Zakos, J. (2001). A computer-aided diagnosis system for digital mammograms based on fuzzy-neural and feature extraction techniques. *IEEE Transactions on Information Technology in Biomedicine*, 5(1):46–54.
- Vese, L. A. et Chan, T. F. (2002). A multiphase level set framework for image segmentation using the Mumford and Shah model. *International Journal of Computer Vision*, 50(3):271–293.
- Vincent, L. (1993). Grayscale area openings and closings, their efficient implementation and applications. *Workshop on Mathematical Morphology and its Applications to Signal Processing*, pages 22–27, Barcelona, Spain.
- Viton, J.-L., Rasigni, M., Rasigni, G. et Llebaria, A. (1996). Method for characterizing masses in digital mammograms. *Optical Engineering*, 35(12):3453–3459.
- Wei, D., Chan, H.-P., Petrick, N., Sahiner, B., Helvie, M. A., Adler, D. D. et Goodsitt, M. M. (1997). False-positive reduction technique for detection of masses on digital mammograms : Global and local multiresolution texture analysis. *Medical Physics*, 24(6):903–914.
- Weszka, J. S., Dyer, C. R. et Rosenfeld, A. (1976). A comparative study of texture measures for terrain classification. *IEEE Transactions on Systems, Man and Cybernetics*, 6(4):269–285.
- Wheeler, F. W., Perera, A. G. A., Claus, B. E., Muller, S. L., Peters, G. et Kaufhold, J. P. (2006). Microcalcification detection in digital tomosynthesis mammography. *SPIE Symposium on Medical Imaging*, volume 6144, page 614420.

- Wilkinson, M. H., Gao, H., Hesselink, W. H., Jonker, J.-E. et Meijster, A. (2008). Concurrent computation of attribute filters on shared memory parallel machines. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 30(10):1800–1813.
- Wilkinson, M. H. F. (2007). Attribute-space connectivity and connected filters. *Image and Vision Computing*, 25(4):426–435.
- Wilson, M., Mitra, S., Roberson, G. H. et Shieh, Y.-Y. (1997). Automated microcalcification detection in mammograms using statistical variable-box-threshold filter method. volume 3165, pages 195–200. SPIE.
- Woods, K., Doss, C., Bowyer, K., Solka, J., Priebe, C. et Philip, W. (1993). Comparative evaluation of pattern recognition techniques for detection of microcalcifications in mammography. *International Journal of Pattern Recognition and Artificial Intelligence*, 7(6):1417–1435.
- Wu, T., Lin, C. et Weng, R. (2004a). Probability estimates for multi-class classification by pairwise coupling. *Journal of Machine Learning Research*, 5:975–1005.
- Wu, T., Moore, R., Rafferty, E., Kopans, D., Stewart, A., Phillips, W., Stanton, M., Eberhard, J., Opsahl-Ong, B., Niklason, L. et Williams, M. (2003a). Tomosynthesis mammography reconstruction using a maximum likelihood method. *Medical Physics*, 30(6).
- Wu, T., Moore, R. H., Rafferty, E. A. et Kopans, D. B. (2004b). A comparison of reconstruction algorithms for breast tomosynthesis. *Medical Physics*, 31(9):2636–2647.
- Wu, T., Stewart, A., Stanton, M., McCauley, T., Phillips, W., Kopans, D. B., Moore, R. H., Eberhard, J. W., Opsahl-Ong, B., Niklason, L. et Williams, M. B. (2003b). Tomographic mammography using a limited number of low-dose cone-beam projection images. *Medical Physics*, 30(3):365–380.
- Wu, T., Zhang, J., Moore, R., Rafferty, E., Kopans, D., Meleis, W. et Kaeli, D. (2004c). Digital tomosynthesis mammography using a parallel maximum-likelihood reconstruction method. *SPIE Symposium on Medical Imaging*, volume 5368, pages 1–11.
- Xu, C. et Prince, J. (1997). Gradient vector flow : a new external force for snakes. *IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 66–71.
- Yang, W., Carkaci, S., Chen, L., Lai, C.-J., Sahin, A., Whitman, G. et Shaw, C. (2007). Dedicated cone-beam breast CT : Feasibility study with surgical mastectomy specimens. *American Journal of Roentgenology*, 189:1312–1315.
- Yoshida, H., Doi, K., Nishikawa, R. M., Giger, M. L. et Schmidt, R. A. (1996). An improved computer-assisted diagnostic scheme using wavelet transform for detecting clustered microcalcifications in digital mammograms. *Academic Radiology*, 3(8):621–627.
- Zadeh, H., Nezhad, S. et Rad, F. (2001). Shape-based and texture-based feature extraction for classification of microcalcification in mammograms. *SPIE Symposium on Medical Imaging*, volume 4684, pages 301–310.
- Zadeh, L. A. (1975). The concept of a linguistic variable and its application to approximate reasoning-III. *Information Sciences*, 9(1):43–80.
- Zhang, Y., Chan, H.-P., Sahiner, B., Wei, J., Goodsitt, M. M., Hadjiiski, L. M., Ge, J. et Zhou, C. (2006). A comparative study of limited-angle cone-beam reconstruction methods for breast tomosynthesis. *Medical Physics*, 33(10):3781–3795.
- Zhen, L. et Chan, A. (2001). An artificial intelligent algorithm for tumor detection in screening mammogram. *IEEE Transactions on Medical Imaging*, 20(7):559–567.
- Zheng, B., Chang, Y.-H., Wang, X.-H. et Good, W. (1999). Comparison of artificial neural network and bayesian belief network in a computer-assisted diagnosis scheme for mammography. *International Joint Conference on Neural Networks (IJCNN)*, 6:4181–4185.

- Zheng, J. et Regentova, E. (2006). Cad system for detecting clustered microcalcifications in digital mammograms using independent component analysis and neural network. *International Journal of Computer Assisted Radiology and Surgery*, 1:517–527.
- Zhou, X. H. et Gordon, R. (1989). Detection of early breast cancer : an overview and future prospects. *Critical Reviews in Biomedical Engineering (CRBE)*, 17(3):203–255.
- Zhu, S. et Yuille, A. (1996). Region competition : unifying snakes, region growing, and Bayes/MDL for multiband image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI)*, 18(9):884–900.
- Zighed, D. et Rakotomalala, R. (2000). *Graphes d'induction : apprentissage et data mining*. Hermes Science Publications.
- Zou, F., Zheng, Y., Zhou, Z. et Agyepong, K. (2008). Gradient vector flow fields and spiculated mass detection in digital mammography images. *International Workshop on Digital Mammography (IWDM)*, pages 299–306, Berlin, Heidelberg. Springer-Verlag.

Table des figures

1.1	Géométrie d'acquisition d'une mammographie	6
1.2	Différentes formes pour les opacités.	7
1.3	Différents types de contours pour les opacités.	7
1.4	Forme schématique d'une distorsion architecturale.	7
1.5	Différents exemples de signes radiologiques suspects en mammographie standard.	8
1.6	Principe de la tomosynthèse.	9
1.7	Exemple de réduction de la superposition de structures	11
1.8	Illustration de la distortion d'objets dans un volume reconstruit.	12
1.9	Schéma général d'une chaîne de détection en mammographie 2D.	13
1.10	Différents schémas de détection en tomosynthèse.	20
1.11	Schéma global de l'approche de détection proposée.	22
2.1	Exemple d'opérateurs connexes et non connexes.	26
2.2	Exemple d'une image d'ombres floues définie sur un domaine Ω 1D.	29
2.3	Exemple d'images floues.	30
2.4	Illustration de $\mathcal{H}^1$	32
2.5	Hyperconnexion $\mathcal{H}_r^1$	32
2.6	Hyperconnexion $\mathcal{H}_r^2$	33
2.7	Composantes connexe floues (b), (c), (d) incluses dans l'ensemble (a).	34
2.8	Composantes connexes incluses dans une image floue.	34
2.9	Persistence d'une composante connexe floue.	37
2.10	Agrégation (c) du résultat du filtrage (b) d'une image d'ombres floues (a).	39
2.11	Filtrage d'un ensemble flou par décomposition en α -coupes.	41
2.12	Exemple de calcul de $\text{Vol}_p^{g,I}(N)$	43
2.13	Exemple de calcul de $\text{Vol}_p^{g,I}(I)$	43
2.14	Exemple de calcul d'une moyenne floue pour une IOF (a) et une IIF (b).	47
2.15	Filtrage d'une coupe de volume de tomosynthèse du sein par $\delta_{\Psi_{mass}}$	50
2.16	Exemple de filtrage de zones plates floues	51
3.1	Représentation d'un ensemble flou en utilisant un arbre de coupes.	56
3.2	Illustration du processus de croissance.	57
3.3	Opérateur $\overline{M}$ dans le cas 1.	60
3.4	Opérateur $\overline{M}$ dans le cas 2.	61
3.5	Opérateur $\overline{M}$ dans le cas 3.	62
3.6	Opérateur $\overline{M}$ dans le cas 4.	63
3.7	Exemple de nœuds compatibles.	63
3.8	Exemple de résultat de l'opérateur <i>subst</i>	64
3.9	Exemple d'utilisation de l'opérateur M	65
3.10	Fonctionnement de l'opérateur $\mathcal{Y}$	67
3.11	Représentation d'un nombre flou par codage de ses transitions.	69
3.12	Exemple d'augmentation de la valeur d'un ensemble flou en un point.	71
3.13	Exemple de diminution de la valeur d'un ensemble flou en un point.	73

3.14	Filtrage d'une image de nombres flous.	74
3.15	Structure de données utilisée pour représenter un arbre.	75
4.1	Construction d'une image d'intervalles flous en déployant un patron sur les niveaux de gris.	78
4.2	Exemple d'extraction de zones plates dans une image en fonction d'un critère d'aire.	79
4.3	Construction d'une image de nombres flous à partir de filtres de rang.	81
4.4	Exemple d'extraction de zones plates dans une image en fonction d'un critère d'aire.	82
4.5	Construction d'un niveau de gris flou en utilisant le principe d'extension.	83
4.6	Décomposition en ondelettes sur 3 niveaux d'une image 2D.	84
4.7	Fonctions de transfert couramment utilisées pour le débruitage par ondelettes.	85
4.8	Représentation LR d'un intervalle flou.	86
4.9	Construction d'un coefficient d'ondelettes flou.	87
4.10	Echantillonnage sur $\mathcal{G}$ d'un niveau de gris flou.	88
4.11	Détection de disques dans une image synthétique corrompue avec un bruit gaussien.	90
4.12	Performances de détection dans une image corrompue avec un bruit gaussien.	91
4.13	Performances de détection dans une image corrompue avec un bruit de Poisson.	92
4.14	Extraction des zones plates floues contenues dans une coupe de tomosynthèse.	93
5.1	Illustration de l'imprécision pour la définition d'un contour.	95
5.2	Illustration de l'incertitude pour la définition d'un contour.	96
5.3	Segmentation par multi-seuillage d'une lésion circonscrite présentant une zone mal définie.	98
5.4	Ensembles de niveaux flous.	100
5.5	Représentation en domaine polaire d'une image contenant une lésion circonscrite.	101
5.6	Exemple de calcul d'une matrice de coûts.	102
5.7	Illustration du principe de matrice de coûts cumulés.	104
5.8	Exemple de pénalisation n'empêchant pas le croisement de deux contours.	104
5.9	Exemple de pénalisation par bande empêchant le croisement de deux contours.	105
5.10	Calcul d'une carte de densité.	107
5.11	Exemple de segmentation d'une lésion circonscrite	109
5.12	Exemple de segmentation d'une lésion spiculée.	109
5.13	Segmentation floue par la méthode de pénalisation par bande.	110
5.14	Image synthétique illustrant la manipulation de l'incertitude par l'algorithme 5.2.	110
5.15	Exemple de segmentation floue d'une lésion spiculée comportant de l'incertitude.	110
5.16	Illustration du calcul de degrés d'appartenance en utilisant l'équation 5.11.	111
5.17	Illustration du calcul de degrés d'appartenance en utilisant l'équation 5.12.	111
5.18	Exemple d'extraction de plusieurs contours flous en utilisant l'algorithme 5.3.	112
6.1	Schéma global de l'algorithme.	116
6.2	Calcul d'attributs flous en utilisant le principe d'extension.	117
6.3	Calcul d'un attribut cumulé à partir d'attributs flous provenant de trois projections.	118
6.4	Description d'un arbre de décision floue.	118
6.5	Construction d'un histogramme flou.	120
6.6	Exemple de fonctions de test au nœud.	120
6.7	Résultat de segmentation.	122
6.8	Agrégation d'attributs flous extraits à partir de plusieurs régions d'intérêt.	123
6.9	Taux d'erreurs de classification.	124
6.10	Exemple de résultats pour le modèle à deux hypothèses.	124
7.1	Exemple d'une lésion spiculée et d'une distorsion architecturale en tomosynthèse du sein.	126
7.2	Dérivée seconde de gaussienne selon trois directions.	127
7.3	Mesures de convergence d'un point vers autre utilisée par Karssemeijer et te Brake (1996).	128
7.4	Comparaison des différentes définitions de convergence d'un point vers un autre.	129
7.5	Différentes approximations d'un sinus pour de petits angles.	130
7.6	Rapport (β) entre $\ \vec{c}\vec{q}\ $ et R_1 tel que $\tan(\theta) = \frac{R}{\ \vec{c}\vec{q}\ }$, $\theta = \frac{R}{\ \vec{c}\vec{q}\ }$ et $\sin(\theta) = \frac{R}{\ \vec{c}\vec{q}\ }$	130
7.7	Différences entre les distances D_1 , D_2 et D_3 normalisées par R_1 en fonction de θ	131

7.8	Comparaison visuelle des différentes mesures de convergence.	132
7.9	Critère de convergence	134
7.10	Extraction d'une carte d'orientations à partir d'une carte de gradients non corrélés.	136
7.11	Illustration des quantités $\delta_{c,r}$ et $\psi_{c,r}$ pour une configuration r et $r - 1$ donnée.	136
7.12	Exemple de B_γ et $\bar{B}_\gamma$ pour différentes orientations avec $R_{\max} = 130$ et $\alpha = 0.25$	138
7.13	Calcul de $ B_\gamma $ et $ \bar{B}_\gamma $	140
7.14	Analyse des orientations dans une zone de l'image.	141
7.15	Détection d'une lésion à forte convergence.	142
7.16	Exemple d'histogrammes d'angles	142
8.1	Schéma global de la chaîne de détection pour les opacités et les distorsions architecturales.	146
8.2	Exemple de lésion non incluse dans l'évaluation des performances.	147
8.3	Illustration de l'étape de marquage de noyaux dense.	149
8.4	Histogrammes normalisés du nombre de détections pour des seins de densités différentes.	150
8.5	Performances de l'étape de marquage pour la détection de noyaux denses suspects.	151
8.6	Illustration des SVM dans le cas de données linéairement séparables	153
8.7	Validation croisée du canal de detection de densités.	154
8.8	Courbe de performance du canal de détection de noyaux denses suspects.	155
8.9	Exemple d'une d'un traitement complet du canal de détection d'opacités.	155
8.10	Courbe de performance du modèle a contrario pour des lésions uniquement spiculées.	156
8.11	Courbe de performance du modèle a contrario.	156
8.12	Agrégation des marqueurs pour le canal de détection de convergences.	157
8.13	Performance du canal de détction de convergence suspectes.	158
8.14	Exemple d'une d'un traitement complet du canal de détection de convergences.	158
8.15	Fusion des deux canaux de détection.	159
A.1	Exemple de construction d'une image d'ombres.	167
A.2	Exemple de 4-connexité et 8-connexité.	168
A.3	Problème de connexité entre fond et objet.	168
A.4	Exemple d'une ouverture d'aire et d'un filtre d'amincissement.	169
B.1	Exemples de t-normes/t-conormes.	171
C.1	Preuve par récurrence du théorème 3.2.1.	174

Liste des Algorithmes

3.1	Algorithme de mise à jour d'un arbre par croissance.	70
3.2	Algorithme de mise à jour d'un arbre par décroissance.	72
3.3	Algorithme de mise à jour d'un ensemble flou en plusieurs points.	72
5.1	Extraction de différents contours par la méthode de pénalisation par bande.	105
5.2	Extraction de différents contours par la méthode des distances.	106
5.3	Extraction de contours flous à partir d'une image.	107
6.1	Propagation d'une particule dans un arbre flou.	119
6.2	Algorithme récursif de construction d'un arbre flou.	119
7.1	Détection rapide de convergences dans un volume de tomosynthèse du sein.	139

Liste des tableaux

2.1	Notations générales sur les images floues.	28
5.1	Performances de la segmentation nette par extraction de chemins minimaux.	108
8.1	Mesures utilisées pour caractériser les lésions à noyaux denses.	152
8.2	Performances de la chaîne complète après fusion des deux canaux.	159
8.3	Performances des chaînes de détection proposées par Chan <i>et al.</i> (2008a).	160
C.1	Énumération de toutes les configurations possibles pour l'opérateur $\overline{M}$	173

Index

- α -coupe, 28, 56, 60, 61, 66–68
- Abdel-Mottaleb *et al.* (1996), 15
- ACR (2003), 6, 149
- Analyse multi-résolution, 14, 83
- Andersen et Kak (1984), 10, 50, 121
- Angiomammographie, 8
- Anscombe (1948), 91
- Arasement volumique, 43
- Arbach *et al.* (2003), 17
- Arbre de décision, 17
- Arbre de décision flou, 117
- construction, 119
 - degré de confiance, 121
 - fonctionnement, 118
- Asymétrie, 18
- Attribut flou, 117
- Attribut flou cumulé, 117, 122
- Avinash *et al.* (2006), 11
- Ayres et Rangayyan (2005), 18
- Ayres et Rangayyan (2007), 18
- Ball et Bruce (2007a), 15
- Ball et Bruce (2007c), 15
- Ball et Bruce (2007b), 15
- Bankman *et al.* (1994), 16, 17
- Berger *et al.* (2007), 53, 74
- Bernard *et al.* (2006), 13, 20
- Bernard *et al.* (2007), 20
- Bernard *et al.* (2008), 20, 145
- Beutel *et al.* (2000), 148
- Blessing *et al.* (2006), 11
- Bloch et Maitre (1995), 46
- Boone *et al.* (2001), 8
- Bornefalk (2005), 14, 17
- Bothorel (1996), 14, 117, 119–121, 164
- S. Bothorel (1997), 17
- Bourrely et Muller (1989), 17
- Braccialarghe et Kaufmann (1996), 13
- Braga-Neto et Goutsias (2003a), 25
- Braga-Neto et Goutsias (2003b), 31, 72, 74
- te Brake et Karssemeijer (1999), 14, 125
- Breen et Jones (1996), 27, 39, 49
- Bruynooghe (2006), 15, 16
- Brzakovic *et al.* (1990), 15
- Bustince *et al.* (2007), 30
- Campanini *et al.* (2004), 17
- Cancer du sein, 5
- Canny (1986), 15
- Cao *et al.* (2004), 17
- Champ de Markov, 15
- Chan et Vese (2001), 98, 99
- Chan *et al.* (2004), 20
- Chan *et al.* (2005), 20, 159, 160
- Chan *et al.* (2007), 159
- Chan *et al.* (2008b), 21, 160
- Chan *et al.* (2008a), 1, 21, 160
- Chan *et al.* (1987), 13
- Chang *et al.* (1979), 8
- Chang *et al.* (1982), 8
- Chang et Laine (1997), 13
- Chassery et Montanvert (1991), 167
- Chemin optimal, 15
- Chen et Lee (1997), 13
- Cheng *et al.* (2003), 16
- Cheng *et al.* (2006), 15
- Cheng *et al.* (1994), 17
- Cheng *et al.* (1998), 14
- Chitre *et al.* (1994), 17
- Claus *et al.* (2002), 1
- Claus et Eberhard (2002), 1
- Coefficient d'atténuation, 6
- Cohérence intrinsèque, 103
- Compacité flou, 48
- Compatibilité radiométrique, 34
- Composante connexe, 26
- Composante connexe floue, 33
- Connexité floue, 25
- Connexité floue, 31
- Contour
- circonscrit, 6
 - mal défini, 7
 - microlobulé, 7
 - partiellement visible, 7
 - spiculé, 7
- Contour actif, 15
- Contour flou, 2, 95, 96, 121
- Contraste flou, 47
- Convergence, 126, 156

- Convergence ϵ -significative, 134
 Correspondance de motifs, 21
 Cortes et Vapnik (1995), 153
 Couto *et al.* (2008), 30
 Cover et Hart (1967), 17
 Crespo *et al.* (1997), 53
 Cristianini et Shawe-Taylor (2000), 151, 153
 Croissance de région, 15
 Croissance d'arbre, 56

 Décroissance d'arbre, 66
 Détection assistée par ordinateur, 1, 5
 en mammographie conventionnelle, 11
 en mammographie standard, 12
 en tomosynthèse, 19
 dans les projections, 19
 dans le volume, 19, 145
 en 2D et 3D, 19
 Daubechies (1988), 89
 Degré d'appartenance d'un contour flou, 97, 105
 Degré de connexité
 définition de Braga-Neto et Goutsias, 32
 définition de Nempont *et al.*, 33
 définition de Rosenfeld, 31
 Dengler *et al.* (1993), 16
 Densité de sein, 149
 Desolneux *et al.* (2000), 21, 125, 132, 134, 143
 Desolneux *et al.* (2001), 21
 Desolneux *et al.* (2003), 21
 Dhawan *et al.* (1996), 17
 Dilatation floue, 46
 Distance entrdeux eontour, 104
 Distorsion architecturale, 1, 7, 18
 Dobbins III et Godfrey (2003), 9
 Domínguez et Nandi (2007), 102, 108
 Domínguez et Nandi (2009), 16
 Donoho et Johnstone (1995), 85
 Donoho *et al.* (1995), 85
 Donoho (1995), 85
 Dose, 19
 Dubois et Prade (1980), 36, 119
 Duffy *et al.* (2003), 5

 Eltonsy *et al.* (2007), 14
 Emboîtement, 55
 van Engeland *et al.* (2006), 18
 van Engeland et Karssemeijer (2007), 18
 Ensembles de niveaux, 15, 98
 Ensembles de niveaux flous, 99
 Ensemble connexe, 26
 Ensemble flou, 28
 croissance, 68
 réduction, 69
 représentation arborescente, 53
 Évènement ϵ -significatif, 134

 Extension floue, 39
 Extraction de caractéristiques, 16

 Famille d'opérateurs connexes flous, 34
 EOCFS, 35
 EOCFE, 35
 EOCFO, 35
 EOCFF, 35
 Filtre connexe, 25, 53
 Filtre connexe flou, 2, 21, 31
 Filtre de nivellement, 27
 Filtre de rang, 80
 Floyd Jr. *et al.* (1994), 17
 Fogel *et al.* (1998), 17
 Fonction génératrice, 133
 Fonction d'appartenance, 28
 Fonction rampe, 48
 Fonction trapèze, 48
 Freer et Ulissey (2001), 12

 Galloway (1975), 16
 Gavrielides *et al.* (2002), 13
 Gennaro *et al.* (2008), 10
 Géraud *et al.* (2004), 27
 Gilbert *et al.* (2006), 12
 Gilbert (1972), 10
 Gisvold *et al.* (1979), 8
 Gordon *et al.* (1970), 10
 Gordon et Herman (1974), 10
 Gout *et al.* (2005), 98
 Grimaud (1991), 14
 Grosjean (2007), 125
 Guliato *et al.* (1998), 15

 Hara *et al.* (2006), 18
 Haralick *et al.* (1973), 16
 Hassanien et Ali (2004), 13
 Hatanaka *et al.* (2001), 15
 Heath et Bowyer (2000), 14
 Heijmans (1997), 25, 53
 Histogramme d'orientation, 140
 Holland *et al.* (1982), 1
 Holland (1992), 17
 Huo *et al.* (2001), 18
 Huo *et al.* (1998), 16
 Hyperconnexion, 31

 Imagerie par résonance magnétique (IRM), 8
 Image d'épaisseur radiologique, 9
 Image d'intervalles flous, 29
 Image d'ombres binaires, 28
 Image d'ombres floues, 28, 97, 148
 Image d'ombres floues nettes, 77
 Image de nombres flous, 30
 Image floue, 2, 28

- agrégation, 39, 45, 79, 148
- construction, 77, 82
 - avec un patron, 78
 - par débruitage, 80
- Imprécision, 2, 95, 106
- Imprécision de débruitage, 2, 80
- Incertitude, 2, 95, 106
- Indice de gradient radial (RGI), 15, 152
- Janikow (1998), 17
- Jansen (2001), 83, 84, 91
- Jansen *et al.* (1997), 85
- Jeunehomme (2005), 8
- Jiang et Wang (2003), 10
- Jiang *et al.* (1997), 17
- John et Ewen (1989), 8
- Kallergi *et al.* (1992), 15
- Kanzaki (1992), 15
- Karssemeijer *et al.* (2003), 12
- Karssemeijer (1993), 13
- Karssemeijer et te Brake (1996), 14, 18, 21, 125, 126, 128, 129, 131, 140, 143, 165
- Karssemeijer et Te Brake (1997), 129, 130
- Kass *et al.* (1988), 15
- Kegelmeyer *et al.* (1994), 16, 17
- Kerlikowske *et al.* (1995), 1
- Kilday *et al.* (1993), 16
- Kim et Kim (2005), 16
- Kim et Park (1999), 16, 17
- Klein (1985), 25
- Klette et Rosenfeld (2004), 26, 167
- Koenderink et van Doorn (1992), 126
- Kramer et Aghdasi (1999), 17
- Kupinski et Giger (1997), 17
- Kupinski et Giger (1998), 15, 152, 157
- Kupinski et Giger (1999), 16
- Kwan *et al.* (2004), 8
- K plus proches voisins, 17
- Lésion
 - irrégulière, 150
 - lobulée, 150
 - spiculée, 150
- Lai *et al.* (1989), 18
- Lau *et al.* (2008), 11
- Lau (1991), 17
- Leclerc (1989), 25
- Lee *et al.* (2000), 15
- Lee *et al.* (1999), 120
- Li *et al.* (1997), 16
- Liu *et al.* (2001), 14
- Llobet *et al.* (2005), 16, 17
- Loi de Beer-Lambert, 5
- Lopez-Diaz et Gil (1997), 80
- De Luca et Termini (1972), 87
- Machine à vecteurs de support, voir (SVM)
- Macrocalcification, 11
- Mallat (1989), 83
- Mallat et Zhong (1992), 15
- Mammographie, 1, 5
 - Pré-traitement, 13
- Marquage, 12
- Marqueur, 148
- Marqueur (Extraction), 14
- Martinez *et al.* (1999), 15
- Masse, voir Opacité
- Matheron (1974), 36
- Matrice
 - de coûts, 101
 - de coûts cumulés, 103
 - de coûts cumulés étendue, 103
- Matrice de cooccurrence, 16
- Matsubara *et al.* (2005), 18
- Meijster et Wilkinson (2002), 53
- Mesure de convergence, 128
- Meyer et Maragos (2000), 51
- Meyer (2005), 27
- Meyer (1998b), 53
- Meyer (1998a), 53
- Meyer et Maragos (1999), 51
- Microcalcification, 1, 6
 - détection en mammographie standard, 13
 - détection en tomosynthèse, 19
- Modélisation a contrario, 2, 21, 125, 132, 156
- Modèles déformables, 15
- Modèle naïf, 133
- Moisan (2003), 21
- Moisan et Stival (2004), 21
- Moyenne floue, 46
- Mudigonda *et al.* (2001), 16, 103
- Muller *et al.* (1983), 8
- Multi-seuillage, 96, 151
- Mumford *et al.* (1988), 25
- Nachtegaal *et al.* (2007), 30
- Najman et Couprie (2006), 53
- National Cancer Institute Consensus Development Panel (1998), 1
- Nempont *et al.* (2008), 31–33, 72, 74
- Niklason *et al.* (1997), 1
- Nœuds compatibles, 60
- Ondelettes, 13, 16
 - coefficients flous, 85
 - débruitage, 84
 - débruitage flou, 82
 - décomposition, 82
- Opérateur connexe, 39

- binaire, 26
- flou pour images à niveaux de gris, 36
- pour images à niveaux de gris, 27
- Opérateur d'extinction, 36
- Opacité, 1, 6, 147
 - détection en mammographie standard, 14
 - détection en tomosynthèse, 20
 - irrégulière, 7
 - lobulaire, 7
 - ovale, 7
 - ronde, 7
- Osher et Fedkiw (2002), 98
- Osher et Sethian (1988), 15, 98
- Ouverture d'aire, 43, 44
- Ouverture d'attribut, 27
- Pénalisation, 104
- Pearl (1988), 17
- Pelote de compression, 9
- Perceptron, 17
- Persistance floue, 36, 43
- Peters *et al.* (2005), 20
- Peters *et al.* (2006a), 14
- Peters *et al.* (2006b), 21
- Peters (2007), 14, 20, 50, 96, 104, 113, 117, 120, 164
- Peters *et al.* (2007), 21, 95, 98, 99, 116, 164
- Pfisterer et Aghdasi (1999), 13
- Pisano *et al.* (2005), 1
- Principe d'extension, 80, 117
- Produit de contraste, 8
- Programmation dynamique, 15, 21, 100, 103, 136, 151
- Puong (2008), 8
- Quantité de flou par pixel, 88
- Région homogène, 33
- Réseau de croyance bayésien, 17
- Réseau de neurones, 17
- Radon (1917), 10, 155
- Rangayyan *et al.* (1997), 16
- Rankin (2000), 8
- Rao et Schunck (1989), 103
- Reconstruction, 9
 - ART, 10
 - artéfacts, 11
 - itérative, 10
 - MLEM, 10
 - rétroprojection filtrée, 10
 - SART, 10
 - SIRT, 10
- Reiser *et al.* (2005), 20
- Reiser *et al.* (2006), 20
- Remontet *et al.* (2003), 5
- Ren *et al.* (2005), 8
- Représentation, 56
- Représentation polaire d'une image, 100
- Rick (1999), 14
- Rifqi (1996), 89
- Rosenfeld (1979), 31, 72, 74
- Rosenfeld (1984), 31
- Rosenfeld et Haber (1985), 47, 48, 79
- Rosenfeld (1992), 31
- Rosenfeld (1998), 31
- Rumelhart *et al.* (1986), 17
- Séparateur à Vaste Marge (SVM), 17, 21, 151
- Safavian et Landgrebe (1991), 17
- Sahiner *et al.* (2001), 16
- Sahiner *et al.* (2001), 15
- Sahiner *et al.* (1996b), 15
- Sahiner *et al.* (1996a), 17
- Sahiner *et al.* (1998), 16
- Salembier et Serra (1995), 25, 27, 53
- Salembier *et al.* (1998), 25, 53
- Sameti et Ward (1996), 15
- Sampat et Bovik (2003), 14
- Saslow *et al.* (2007), 8
- Schulte *et al.* (2006a), 30
- Schulte *et al.* (2006b), 30
- Segmentation, 15, 150
- Segmentation floue par seuillages, 96
- Segmentation par région, 15
- Serra (1982), 36
- Serra et Salembier (1993), 25, 53
- Serra (1998), 31
- Sibala *et al.* (1981), 8
- Singh *et al.* (2008a), 160
- Singh (2008), 1, 160
- Singh *et al.* (2008b), 21
- Soille (2008), 51
- Stabilité radiométrique, 34
- Strickland et Hahn (1997), 13
- Superposition de tissus, 1, 6, 10
- Sur-densité, voir Opacité
- Tahmoush et Samet (2006a), 18
- Tahmoush et Samet (2006b), 18
- Tahami *et al.* (2007), 30
- Timp et Karssemeijer (2004), 15, 21, 95, 100, 101, 103, 108, 164
- Tizhoosh (2008), 30
- Tomographie, 8
- Tomosynthèse numérique du sein, 1, 7, 145
 - algorithme de reconstruction, voir Reconstruction
 - ouverture angulaire, 9, 11
 - projection 2D, 9
- Torre et Poggio (1986), 15
- Tourassi *et al.* (2003), 18

- Transformée d'Anscombe, 91
Tube à rayons X, 6
- Vachier (2001), 27, 35, 53
Vachier (1995), 14–16, 169
Vapnik (1995), 17, 153
Veldkamp et Karssemeijer (1999), 17
Verma et Zakos (2001), 17
Vese et Chan (2002), 98
Vincent (1993), 25, 39, 53
Viton *et al.* (1996), 16
- Wei *et al.* (1997), 16
Weszka *et al.* (1976), 16
Wheeler *et al.* (2006), 19
Wilkinson (2007), 25
Wilkinson *et al.* (2008), 53, 57, 69, 75
Wilson *et al.* (1997), 13
Woods *et al.* (1993), 17
Wu *et al.* (2003a), 10, 20
Wu *et al.* (2003b), 8
Wu *et al.* (2004a), 154
Wu *et al.* (2004b), 10
Wu *et al.* (2004c), 10
- Xu et Prince (1997), 14
- Yang *et al.* (2007), 8
Yoshida *et al.* (1996), 13
- Zadeh *et al.* (2001), 17
Zadeh (1975), 81, 117
Zhang *et al.* (2006), 11
Zhen et Chan (2001), 14, 15
Zheng et Regentova (2006), 17
Zheng *et al.* (1999), 17
Zhou et Gordon (1989), 13
Zhu et Yuille (1996), 25
Zighed et Rakotomalala (2000), 17
Zone plate, 25, 27
Zone plate floue, 33, 78, 91
Zou *et al.* (2008), 14, 17