



Contributions à l'étude de la classification spectrale et applications

Sandrine Mouysset

► To cite this version:

Sandrine Mouysset. Contributions à l'étude de la classification spectrale et applications. Mathématiques [math]. Institut National Polytechnique de Toulouse - INPT, 2010. Français. NNT: . tel-00573433

HAL Id: tel-00573433

<https://theses.hal.science/tel-00573433>

Submitted on 3 Mar 2011

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.



Université
de Toulouse

THÈSE

En vue de l'obtention du DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par :

Institut National Polytechnique de Toulouse (INP Toulouse)

Discipline ou spécialité :

Mathématiques Appliquées

Présentée et soutenue par :

Sandrine MOUYSET

le : mardi 7 décembre 2010

Titre :

Contributions à l'étude de la classification spectrale
et applications

JURY

Mario ARIOLI, Directeur de Recherche Rutherford Appleton Laboratory

Jean-Baptiste CAILLAU, Professeur à l'université de Bourgogne

Vincent CHARVILLAT, Professeur à l'ENSEEIH

Pierre HANSEN, Professeur à HEC Montréal

Joseph NOAILLES, Professeur Emérite à l'ENSEEIH

Daniel RUIZ, Maître de Conférence à l'ENSEEIH

Clovis TAUBER, Maître de Conférence à l'université François Rabelais

Ecole doctorale :

Mathématiques Informatique Télécommunications (MITT)

Unité de recherche :

Institut de Recherche en Informatique de Toulouse (IRIT)

Directeur(s) de Thèse :

Joseph NOAILLES, Professeur Emérite à l'ENSEEIH

Daniel RUIZ, Maître de Conférence à l'ENSEEIH

Rapporteurs :

Mario ARIOLI, Directeur de Recherche Rutherford Appleton Laboratory

Jean-Baptiste CAILLAU, Professeur à l'université de Bourgogne

Pierre HANSEN, Professeur à HEC Montréal

Table des matières

Table des matières	i
Liste des figures	v
Liste des tableaux	vii
Introduction	3
1 Classification spectrale : algorithme et étude du paramètre	9
1.1 Présentation de la classification spectrale	9
1.1.1 Algorithme de classification spectrale	9
1.1.2 Problème du choix du paramètre	12
1.2 Heuristiques du paramètre	16
1.2.1 Cas d’une distribution isotropique avec mêmes amplitudes suivant chaque direction	16
1.2.2 Cas d’une distribution isotropique avec des amplitudes différentes suivant les directions	17
1.3 Validations numériques	19
1.3.1 Mesures de qualité	19
1.3.2 Exemples géométriques de dimension $p \geq 2$	26
1.4 Méthodes de classification spectrale modifiées	27
1.4.1 Cas de données volumiques	27
1.4.2 Traitement d’images	27
2 Classification et éléments spectraux de la matrice affinité gaussienne	33
2.1 Diverses interprétations	33
2.2 Présentation du résultat principal	35
2.3 Propriétés de classification dans le continu	41
2.3.1 Quelques résultats d’Analyse Fonctionnelle	41
2.3.2 Classification via l’opérateur de Laplace	43
2.3.3 Classification via l’opérateur de la chaleur avec conditions de Dirichlet	45
2.3.4 Classification via l’opérateur de la chaleur dans $\mathbb{R}^p$	48
2.4 Propriété de classification en dimension finie	52
2.4.1 Éléments finis de Lagrange	52
2.4.2 Interprétation des éléments de la classification spectrale	54
2.4.3 Propriété de classification du produit $(A + \mathbb{I}_N)M$	57
2.4.4 Condensation de masse	61
2.5 Discussions	66

2.5.1	Expérimentations numériques	67
2.5.2	Choix du paramètre gaussien	68
2.5.3	Passage du discret au continu : Triangulation de Delaunay	74
2.5.4	Etape de normalisation	77
2.5.5	Cas limites de validité de l'étude	79
3	Parallélisation de la classification spectrale	83
3.1	Classification spectrale parallèle : justification	83
3.1.1	Etude dans le continu	84
3.1.2	Passage au discret	86
3.1.3	Définition du nombre de classes k	88
3.2	Classification spectrale parallèle avec interface	90
3.2.1	Implémentation : composantes de l'algorithme	90
3.2.2	Justification de la cohérence de la méthode	92
3.2.3	Expérimentations parallèles	95
3.3	Classification spectrale parallèle avec recouvrement	101
3.3.1	Principe	101
3.3.2	Justification de la cohérence de la méthode	103
3.3.3	Expérimentations numériques : exemples géométriques	105
3.4	Segmentation d'images	109
3.4.1	Boîte affinité 3D rectangulaire	109
3.4.2	Boîte affinité 5D rectangulaire	114
4	Extraction de connaissances appliquée à la biologie et l'imagerie médicale	119
	Partie 1 : Classification de profils temporels d'expression de gènes	120
4.1	Présentation du problème biologique	120
4.2	Méthode Self Organizing Maps	122
4.2.1	Présentation de la méthode Self Organizing Maps	122
4.2.2	Adaptation de la méthode Self Organizing Maps pour données temporelles d'expression de gènes	125
4.2.3	Simulations	128
4.3	Spectral Clustering sur les données génomiques	130
4.4	Discussion	136
	Partie 2 : Imagerie médicale : segmentation d'images de la Tomogra-	
	phie par Emission de Positons	137
4.5	L'imagerie Tomographique par Emission de Positons	137
4.6	Classification en TEP dynamique	139
4.7	Segmentation non supervisée basée sur la dynamique des voxels	141
4.8	Validation par simulation de courbes temps-activité	142
4.8.1	Simulation TACs	142
4.8.2	Spectral clustering sur TACs	143
4.9	Validation par simulation d'images TEP réalistes	145
4.9.1	Simulation des images TEP réalistes	146
4.9.2	Critères d'évaluation quantitative	147
4.9.3	Résultats du spectral clustering	149
4.9.4	Comparaison avec la méthode k -means	152

Conclusion et perspectives	155
Bibliographie	159

Table des figures

1.1	Illustration des étapes du clustering spectral	10
1.2	Influence du paramètre dans l'espace de projection spectrale	11
1.3	Paramètre global : exemples avec le paramètre σ proposé par <i>Brand</i> [20]	14
1.4	Paramètre local : exemples avec le paramètre σ proposé par <i>Perona</i> [116]	15
1.5	Principe de l'heuristique (1.2) dans le cas 2D	16
1.6	Principe de l'heuristique (1.3) dans le cas 3D	17
1.7	Exemples avec l'heuristique 1 (σ) pour $p = 2$	20
1.8	Exemples avec l'heuristique 2 (σ) pour $p = 2$	21
1.9	Pourcentages d'erreur de clustering en fonction de σ	23
1.10	Ratio de normes de Frobenius fonction du paramètre σ	25
1.11	Exemple 1 & 2 : 6 blocs et 3 portions de N -spheres	26
1.12	Test sur images avec une boîte rectangulaire 3D d'affinité	31
1.13	Test sur images avec une boîte produit 2D par 1D d'affinité	32
2.1	Définitions du clustering	36
2.2	Exemples des donuts : Données, maillage et fonctions propres	39
2.3	Exemple des donuts : Corrélacion et différence en norme entre V_i et respectivement les vecteurs propres de $(A + \mathbb{I}_N)M$ et de A , pour $i \in \{1, 2\}$	40
2.4	Exemple des donuts	45
2.5	Approximation de l'ouvert Ω_1 par un ouvert $\mathcal{O}_1$	49
2.6	Maillage du fermé $\bar{\mathcal{O}} : x_i \in \overset{\circ}{\Omega}, \forall i \in \{1, ..N\}$	53
2.7	Exemple des donuts	55
2.8	Exemple des donuts	60
2.9	Exemple des donuts	65
2.10	Exemple 1 : Fonction propre discrétisée $V_{1,i}$ et produits $(A + \mathbb{I}_N)MV_{1,i}$, $AV_{1,i}$ pour $i \in \{1, 2, 3\}$	69
2.11	Exemple 1 : Corrélacion et différence en norme entre $V_{1,i}$ et, respectivement, les vecteurs propres X_l de $(A + \mathbb{I}_N)M$ et Y_l de A , pour $i \in \{1, 2, 3\}$	70
2.12	Exemple 1 : Influence de la discrétisation sur la corrélacion entre $V_{1,i}$ et les vecteurs propres Y_l de A , pour $i \in \{1, 2, 3\}$	70
2.13	Exemple 2 : Fonction propre discrétisée $V_{1,i}$ et produits $(A + \mathbb{I}_N)MV_{1,i}$, $AV_{1,i}$ pour $i \in \{1, 2, 3\}$	71
2.14	Exemple 2 : Corrélacion et différence en norme entre $V_{1,i}$ et respectivement les vecteurs propres X_l de $(A + \mathbb{I}_N)M$ et Y_l de A , pour $i \in \{1, 2, 3\}$	72
2.15	Exemple 2 : Influence de la discrétisation sur la corrélacion entre $V_{1,i}$ et les vecteurs propres Y_l de A , pour $i \in \{1, 2, 3\}$	72
2.16	Estimation de la distance δ	73

2.17	Exemple 2 : Mesure de qualité de clustering en fonction de t	74
2.18	Exemple 3 : Fonction propre discrétisée $V_{1,i}$ et produits $(A + \mathbb{I}_N)MV_{1,i}$, $AV_{1,i}$ pour $i \in \{1, 2\}$	75
2.19	Exemple 3 : Corrélation et différence en norme entre $V_{1,i}$ et respectivement les vecteurs propres X_l de $(A + \mathbb{I}_N)M$ et Y_l de A , pour $i \in \{1, 2\}$	76
2.20	Principe de la triangulation de Delaunay	76
2.21	Triangulation de Delaunay sur divers exemples	78
2.22	Exemple 1 : Vecteurs propres avec étape de normalisation	80
2.23	Exemple 2 : Vecteurs propres avec étape de normalisation	80
2.24	Exemples de cas limites [80]	81
3.1	Exemple de pavage d'un ouvert Ω	84
3.2	Exemple de $\omega_1^1 \mathcal{R}^{P_C} \omega_2^1$	85
3.3	Exemples pour le choix du nombre de cluster k	89
3.4	Principe du spectral clustering parallèle divisé en $q = 2$ sous-domaines	90
3.5	Exemple cible : Résultat du spectral clustering pour l'interface et les sous-domaines	93
3.6	Règles pour l'absence de points dans une zone d'intersection	94
3.7	Résultats du spectral clustering parallèle avec interface sur l'exemple des 2 sphères découpées en 2 sous-domaines	97
3.8	2 sphères tronquées : Speed up et efficacité	98
3.9	Exemple géométrique (gauche) et zooms (droite) : $N = 4361$	98
3.10	Résultats du spectral clustering parallèle sur l'exemple des 3 sphères découpées en 2 sous-domaines	99
3.11	3 sphères tronquées : Speed up et efficacité	100
3.12	Principe de l'alternative clustering spectral parallèle avec recouvrement pour $q = 2$	102
3.13	Exemple cible : Résultat du spectral clustering avec recouvrement après étape de regroupement et par sous-domaines	103
3.14	Résultats du spectral clustering sur le cas des 2 sphères découpées en 2 clusters	105
3.15	2 sphères tronquées : Speed up et efficacité	106
3.16	Résultats du spectral clustering parallèle sur l'exemple des 3 sphères découpées en 2 sous-domaines	107
3.17	3 sphères tronquées : Speed up et efficacité	108
3.18	Exemple de segmentation d'image testée sur Hyperion	110
3.19	Ensemble des données 3D de l'image bouquet	111
3.20	Résultats du spectral clustering sur le bouquet	112
3.21	Autres exemples de segmentation d'images testées sur Hyperion	113
3.22	Test sur un exemple de mahua	115
3.23	Champs colorés dans le Vaucluse, $N = 128612$ points	117
3.24	Geyser du grand Prismatic, $N = 155040$ points	118
4.1	Principe des puces à ADN	121
4.2	Root pathosystems in Medicago Truncatula	122
4.3	Initialisations de la méthode SOM (t en abscisse, niveau d'expression en ordonnée)	126
4.4	Exemple de profils temporels différents mais de même norme (t en abscisse, niveau d'expression en ordonnée)	126
4.5	Résultats pour l'espèce A17 (t en abscisse, niveau d'expression en ordonnée)	128
4.6	Résultats pour l'espèce F83 (t en abscisse, niveau d'expression en ordonnée)	129
4.7	Différence entre les espèces A17 et F83 (t en abscisse, niveau d'expression en ordonnée) 130	

4.8	Etude avec les données des profils temporels de l'espèce A17	131
4.9	Etude avec les données normalisées des profils temporels de l'espèce A17	133
4.10	Etude avec les données normalisées des profils temporels de l'espèce F83	134
4.11	Distribution et normalisation des données de la différence entre les profils temporels des espèces A17 et F83	135
4.12	Balayage des différents profils	136
4.13	Principe de la tomographie par émission de positons	138
4.14	Courbes temps activité de deux régions cérébrales	139
4.15	Différentes régions du cerveau	142
4.16	Exemple de TACs affectées à différentes régions cérébrales	143
4.17	Ensemble des profils de S	144
4.18	Application du spectral clustering sur des TACS bruitées et perturbées en amplitude	144
4.19	Mesures de qualité en fonction des valeurs de σ	145
4.20	Estimation du nombre de clusters	146
4.21	Simulation d'un imageur TEP dans GATE	147
4.22	Fantôme de cerveau Zubal utilisé pour la simulation	148
4.23	Résultat du spectral clustering pour $k = 6$	150
4.24	Résultat du spectral clustering sur les TACs normalisées	151
4.25	Comparaison avec la méthode k -means pour $k = 6$	152
4.26	Comparaison des TACs du thalamus et du lobe frontal	153

Liste des tableaux

1.1	Exemples géométriques	28
1.2	Exemples géométriques avec l’affinité (1.5)	28
3.1	2 sphères tronquées : cas avec interface	96
3.2	3 sphères tronquées : cas avec interface	100
3.3	2 sphères tronquées : cas avec recouvrement	106
3.4	3 sphères tronquées : cas avec recouvrement	108
3.5	Nombres de clusters suivant leur cardinal	114
4.1	Catégories génomiques présentes sur la puce dédiée	123
4.2	Critères d’évaluation quantitative	151

Remerciements

Je tiens tout d'abord à remercier mon directeur de thèse Joseph Noailles pour m'avoir proposé de travailler dans ce nouvel axe de recherche que constitue la fouille de données au sein de l'équipe APO et ce dès le stage du master 2. Je lui suis reconnaissante de m'avoir permis aussi d'explorer d'autres axes pour l'étude de ce sujet et de m'avoir soutenue dans ces choix. En outre, il m'a encouragée à créer des collaborations interdisciplinaires et ainsi m'a fait découvrir une autre facette particulièrement intéressante de la recherche.

Je remercie ensuite mon co-directeur de thèse Daniel Ruiz pour son aide et ses encouragements constants dans le travail. Au delà de l'encadrement, j'ai particulièrement apprécié le côté amical et humain dont il sait faire preuve dans certaines circonstances.

Je tiens à remercier M. Pierre Hansen, M. Mario Arioli et M. Jean-Baptiste Caillau pour avoir accepté de me faire l'honneur de rapporter cette thèse. Je leur sais gré de l'attention qu'ils ont portée à la lecture de ce manuscrit ainsi que des remarques qu'ils ont formulées à ce propos. Je tiens également à remercier M. Vincent Charvillat et M. Clovis Tauber pour avoir bien voulu faire partie de mon jury de thèse.

Dans le cadre de cette thèse, j'ai eu la chance de collaborer avec des personnes de différents domaines avec lesquelles les échanges furent agréables et fructueux dont notamment Clovis Tauber (Université François Rabelais), Laurent Gentzbittel (ENSAT), Jean-Marie François (INSA), Alice Senescau (DendriS) et Marc Dubourdeau (AMBIOTIS).

Je souhaite remercier tous les membres de l'équipe APO pour la convivialité et la disponibilité dont ils ont fait preuve au cours de ces années. En particulier, j'adresse mes remerciements à Michel Daydé, Patrick Amestoy et Serge Gratton pour leurs conseils avisés. Par ailleurs, je voudrais remercier les membres de l'équipe VORTEX côtoyés au fil des enseignements pour leur accueil chaleureux. Enfin, j'ai bien apprécié les pauses café en compagnie des membres de l'équipe IRT.

J'adresse aussi tout particulièrement mes remerciements à Géraldine Morin, Ronan Guivarch, Xavier Thirioux, Jérôme Ermont et Philippe Queinnec pour leur soutien, leurs conseils amicaux en enseignement et en recherche et pour tous les bons moments passés en leur compagnie.

Mes remerciements s'adressent également à Frédéric Messine et Pierre Spiteri, mes voisins de bureau, pour leur soutien et leur bonne humeur constants.

Je remercie mes amis doctorants et stagiaires rencontrés au fil des années Pascaline, Nassima, Sébastien, Benoît C., Stéphane, Benoît B., Michael, Ahmed et tous mes co-bureaux, notamment Jordan et Olivier, pour la bonne ambiance de travail !

Je remercie mes amis proches Victoria, Nicolas, Jérôme et Laurent pour m'avoir encouragée pendant toutes ces années.

Enfin, je remercie énormément mes parents et mon frère pour leur très précieux soutien et leur présence au quotidien.

Introduction

Les domaines des biologies cellulaire et moléculaire connaissent de grandes avancées avec l'arrivée notamment des nanotechnologies et des biopuces. A l'instar de la génétique, l'imagerie médicale utilise des nouvelles techniques telles que la scintigraphie pour étudier l'activité métabolique d'un organe grâce à l'injection d'un radiotraceur dont on connaît le comportement et les propriétés biologiques. Dans ces deux domaines, de grands flots de données à analyser sont générés. Ces données sont par conséquent multidimensionnelles, très nombreuses et souvent bruitées. Elles représentent, suivant le domaine d'applications, des intensités lumineuses pour transcrire l'expression de gènes, des courbes d'activités temporelles, des images 3D pour l'imagerie médicale. L'extraction de connaissances issues de ces données massives biologiques ou médicales devient alors un réel problème. Son processus se décompose en plusieurs étapes. Tout d'abord, un prétraitement des données est réalisé pour mettre à l'échelle et normaliser les données. Une analyse exploratoire des données est en général faite en vue d'une éventuelle réduction de la cardinalité de l'ensemble des variables [100, 34, 103]. Suit alors l'étape de fouille de données, étape mathématique du processus d'Extraction de Connaissance. A partir d'une formulation mathématique basée sur les objets et le choix d'une mesure de similarité, une résolution numérique du problème de classification est effectuée. Des étapes d'évaluation et de validation de la classification obtenue sont requises pour juger la pertinence de l'étape fouille de données. Enfin, la visualisation et l'interprétation des résultats constituent l'étape de post-traitement du processus.

Dans ce processus, nous nous intéressons à l'étape de fouille de données. Le manque de supervision lié au manque d'expérience, de connaissances a priori sur ces nouvelles techniques d'expérimentations mettent à mal de nombreuses méthodes d'analyse de données. Sont alors privilégiées les méthodes non supervisées (ou *clustering*) [35, 61, 53, 113, 17, 52]. Le problème de la classification non supervisée peut s'énoncer comme suit : *étant donné m items (ou objets) définis par les variables numériques de n variables (ou attributs), on regroupe ces items à l'aide d'une mesure de similarité (ou proximité), en classes de telle sorte que les objets de la même classe soient le plus semblable possible et des objets de classes différentes le moins semblable possible.*

Les méthodes non supervisées sont grossièrement regroupées en deux grandes catégories : les méthodes hiérarchiques, de nature ensembliste, et les méthodes de partitionnement basées sur des approches probabilistes, d'optimisation ou bien de théorie des graphes. Les méthodes hiérarchiques consistent à transformer une matrice de similarité en une hiérarchie de partitions emboîtées. La hiérarchie peut être représentée comme un arbre dendrogramme dans lequel chaque cluster est emboîté dans un autre. Plusieurs variantes sont distinguées [61, 65]. Le clustering ascendant [72, 73] débute en considérant chaque objet comme un cluster et itérativement réduit le nombre de clusters en fusionnant les objets les plus proches. La méthode descendante [71] débute, quant à elle, avec un seul cluster regroupant tous les objets et divise les clusters afin que l'hétérogénéité soit la plus réduite possible. Les partitions évoluent au cours des itérations via un critère de qualité comme le critère de Ward [111], critère d'agrégation selon l'inertie, alliant à la fois la dispersion à l'intérieur d'une

classe et la dispersion entre les classes. Ces processus se terminent par un critère d'arrêt à définir (en général, un nombre a priori de clusters finaux). Ces méthodes sont performantes et s'adaptent à tout type de données et pour toute mesure de similarité. Par contre, le manque d'évaluation au sein des clusters itérativement construits et le manque de critère d'arrêt précis restent les points faibles de ces algorithmes.

L'autre grande catégorie de méthodes non supervisées est basée sur le partitionnement des données initiales et la réallocation dynamique de partitions. Parmi elles, l'optimisation de fonctions objectifs reste le fondement mathématique des plus anciennes méthodes de classification. La méthode k -means est l'une des plus utilisées [77, 56, 95] et fonctionne sur deux estimateurs [90]. Le premier est basé sur l'estimation du centre de gravité de chaque cluster comme l'isobarycentre de l'ensemble des points appartenant au cluster. Le second est fondé sur l'estimation globale de la partition des données *i.e* l'inertie des nuages de points par cluster. Ce dernier représente les moindres carrés ordinaires entre les points et leurs centres correspondants pour la norme L^2 . Il existe deux principales versions à l'optimisation du k -means itératif. La première version du k -means [43] est semblable, d'un point de vue probabiliste, à un maximum de vraisemblance et se décompose en deux étapes : assigner les points à leur centre le plus proche puis recalculer les centres des nouveaux ensembles de points jusqu'à satisfaire un critère d'arrêt (par exemple, l'invariance de la composition des clusters entre deux itérations). L'autre version de cette méthode utilise la fonction objectif des moindres carrés ordinaires pour permettre le réassignement des points. Seuls les points améliorant cette fonction objectif seront déplacés et leurs centres respectifs recalculés. Ces variantes pour la norme L^2 ont la même complexité numérique mais dépendent de la mesure de similarité utilisée et de la nature des données. De nombreux autres algorithmes permettent d'optimiser la fonction objectif [9, 54] ou de définir heuristiquement un nombre optimal k de clusters [45, 15, 97] et l'initialisation des centres [4, 19, 57]. Cependant, la méthode k -means reste sensible à l'initialisation des centres de clusters, aux optima locaux, aux données aberrantes et à la nature des données. Mais le principal problème de cette méthode basée sur la séparation linéaire est que, par définition, le k -means ne partitionne pas correctement les clusters couvrant un domaine non-convexe c'est-à-dire dans le cas de séparations non linéaires.

Traiter des domaines non-convexes revient à considérer les clusters comme des ouverts et à ajouter des notions topologiques de connexité, de densité et de voisinage. En effet, dans un espace Euclidien, un ouvert peut être divisé en un ensemble de composantes connexes. Ainsi les clusters seront caractérisés de par leur densité comme des composantes connexes denses et des points appartenant au même cluster seront définis par leur voisinage. Les approches probabilistes utilisent ces propriétés car les éléments des différents clusters sont supposés être générés par des distributions de probabilités différentes ou bien de même famille avec des paramètres différents. On distingue donc deux principales approches parmi les méthodes de partitionnement basées sur la densité de points [53]. La première [37] juge la densité autour d'un voisinage d'un objet comme suffisante suivant le nombre de points situés à un certain rayon de cet objet. Elle nécessite donc la définition de deux paramètres : le rayon du voisinage et celui du cardinal minimum pour constituer un cluster. Ainsi les données aberrantes, bruitées, sont considérées et le processus est performant pour de faibles dimensions. Par contre, ces deux paramètres restent difficiles à choisir de manière automatique c'est-à-dire sans apport d'informations extérieures. L'autre approche [59, 53] est basée sur un ensemble de distribution de fonctions de densités. A partir de fonctions d'influences (en général gaussiennes ou signaux carrés) décrivant l'impact d'un point sur son voisinage, la densité globale de l'espace des données est modélisée de manière itérative comme la somme des fonctions d'influences de tous les points. En identifiant les densités d'attractions c'est-à-dire les maxima locaux de la densité globale, les clusters sont alors déterminés mathématiquement. Les paramètres de voisinage et de la fonction d'influence restent à nouveau à choisir. De plus, les lois de probabilités utilisées ne sont pas toujours

adéquates dans un cadre non supervisé et pour de grandes dimensions.

Pour renforcer l'efficacité du clustering face aux données de grandes dimensions, l'approche via les grilles utilise, par définition, une structure de grille de données et la topologie inhérente à ce type d'espace. Cette méthode [88, 110] quantifie l'espace en un nombre fini de cellules formant ainsi une structure de grille. L'information individuelle des données est remplacée par une information condensée sur les données présentes dans une cellule de la grille. Cette méthode est très performante numériquement comparée aux méthodes basées sur la densité de points et, de plus, totalement indépendante du nombre de données. Par contre, elle est dépendante du nombre et des tailles de cellules sur chaque dimension de l'espace discrétisé. Un mauvais choix de ces paramètres occasionne la perte d'information engendrée par la compression des données.

Parmi les méthodes récentes de partitionnement, les méthodes à noyau et le clustering spectral [27, 41] peuvent aussi séparer non-linéairement des clusters et sont très employées dans les domaines comme la biologie, la segmentation d'images. Elles reposent sur le même principe : elles utilisent des relations d'adjacence (de similarité) entre tous les couples de points sans a priori sur les formes des clusters. L'utilisation des propriétés inhérentes aux noyaux de Mercer [7, 58] permettent de projeter implicitement les données dans un espace de grande dimension dans lequel les données seront linéairement séparables. Ces deux approches diffèrent dans la définition de l'espace dans lequel les clusters sont constitués. Les méthodes à noyaux exploitent ce plongement dans un espace de grande dimension pour séparer les points alors que le clustering spectral utilise les éléments spectraux de ces noyaux pour plonger les données transformées dans un espace de dimension réduite. L'utilisation des espaces de Hilbert à noyau confère un caractère non-linéaire à de nombreuses méthodes originellement linéaires. Ainsi les méthodes à noyau [48] regroupent les méthodes classiques de partitionnement de classification supervisée ou non comme le k -means, les cartes auto organisatrices [68], méthodes probabilistes à noyaux et clustering à vecteurs supports [16] et introduisent une fonction noyau pour mesure de similarité [23], [74]. Les méthodes spectrales, quant à elles, ont été formulées à l'aide de la théorie spectrale des graphes et des coupes de graphes [25]. En effet, la matrice d'affinité peut être identifiée à un graphe : les noeuds du graphe sont les données et le poids des arêtes formées entre chaque paire de points représente alors la mesure de similarité locale entre deux points. Le problème de clustering peut alors être formulé à partir de la construction d'une coupe normalisée (ou conductance) pour mesurer la qualité d'une partition de graphe en k clusters. Minimiser ce critère, d'après Shi et Malik [91], est équivalent à maximiser un problème de trace NP-complet. Par une relaxe de la contrainte d'égalité, les vecteurs propres de la matrice affinité permettent de calculer la partition discrète des points. Dhillon *et al* [31, 29, 30] ont montré que la méthode k -means à noyau et les méthodes spectrales sont équivalentes : le critère de coupe normalisée définissant les méthodes spectrales peut être interprété comme un cas particulier de la fonction objectif caractérisant le k -means à noyau. Dans les deux cas, la solution revient à résoudre un problème de maximisation de trace d'une matrice.

Dans le cadre de notre étude, nous voulons considérer des méthodes performantes pouvant s'appliquer sur des données sans information a priori sur la partition finale ni la forme des classes. Les méthodes spectrales et à noyaux, de par leurs principes, permettent d'aboutir à une partition des données avec un minimum de paramètres à fixer. Nous avons aussi affaire à des données multidimensionnelles. Ainsi l'approche spectrale via une réduction de dimension de l'espace permet de traiter efficacement ce type de données. De plus, l'algorithme associé à cette méthode est simple à implémenter et les classes dans l'espace de projection spectral (ou spectral embedding) sont déterminés en appliquant une méthode k -means. Notre étude portera donc sur la classification spectrale (ou clustering spectral) et les problèmes inhérents à cette méthode.

Un premier problème apparaît avec la mesure de similarité. Le noyau gaussien est communément

choisi pour mesure de similarité et, à l’instar des méthodes précédemment présentées, il dépend d’un paramètre scalaire σ qui affecte les résultats [84]. Or, ce paramètre représente le rayon du voisinage [105] : il sert de seuil à partir duquel deux points sont déclarés proches ou non. Divers points de vues ont été adoptés pour le définir. Des approches globales [20, 84] et locales ont aussi été établies via des interprétations physiques [42, 116] pour définir ce paramètre. Mais il reste à montrer l’influence de ce paramètre sur la qualité du clustering à travers des mesures quantitatives et à en dégager une définition à même de donner des informations sur la partition finale.

D’un autre côté, le fonctionnement même de la méthode de spectral clustering requiert des justifications théoriques. En effet, des limites sur cette méthode ont été recensées [107, 80, 78]. Malgré l’explication générale à travers la conductance d’un graphe, il reste à justifier comment le regroupement dans un espace de plus petite dimension définit correctement le partitionnement des données d’origines. Plusieurs travaux ont été menés pour expliquer le fonctionnement du clustering spectral. Comme son principe est de regrouper des données à travers une notion de voisinage, le Clustering spectral peut être interprété comme la discrétisation d’un opérateur Laplace-Beltrami et d’un noyau de chaleur défini sur des variétés sous l’hypothèse d’un échantillonnage uniforme de la variété [11, 12, 13]. D’autres raisonnements probabilistes basés sur des modèles de diffusion [82, 83, 81] ont établi des résultats asymptotiques pour un grand nombre de points. La consistance de cette méthode a aussi été étudiée en considérant un grand échantillon de données [63, 106]. Des propriétés sous des hypothèses standards pour les matrices du Laplacien normalisé ont été prouvées, incluant la convergence du premier vecteur propre vers une fonction propre d’un opérateur limite. Cependant, d’un point de vue numérique, le Clustering spectral partitionne correctement un ensemble fini de points. Considérer un ensemble fini de points pose le problème du sens à donner à la notion de classe et de comment lier les classes finales à des éléments spectraux (valeur propre/vecteur propres extraits d’une matrice affinité). Il reste à expliquer comment le clustering dans un espace de projection spectral caractérise le clustering dans l’espace d’origine et étudier le rôle du paramètre du noyau Gaussien dans le partitionnement.

Enfin, dans le cadre de la biologie et de la segmentation d’images, nous avons affaire à un grand flot de données. Le calcul de la matrice affinité sur l’ensemble des données puis de son spectre devient alors très coûteux. Une parallélisation de cette méthode, notamment des travaux principalement basés sur des techniques d’algèbre linéaire pour réduire le coût numérique, ont été développés [44, 92, 38]. Cependant, les algorithmes développés ne s’affranchissent pas de la construction de la matrice affinité complète. Cette étape reste très coûteuse en temps de calcul et en stockage mémoire. De plus, déterminer le nombre de clusters reste un problème ouvert. Il reste donc à définir une stratégie parallèle pour traiter un grand nombre de données et automatiser le choix du nombre de clusters et les paramètres inhérents au clustering spectral et à la stratégie parallèle.

Plan de la Thèse

Cette thèse se découpe autour de quatre chapitres suivant les points d’étude précédemment évoqués. Nous testerons la méthode sur des exemples géométriques en 2D et 3D (ou challenges de clustering) et sur des cas de segmentation d’images sur les quatre premiers chapitres. En effet, l’aspect visuel pourra nous aider à juger de la performance de la méthode avant de l’appliquer sur des cas réels de biologie ou d’imagerie médicale dans le quatrième chapitre.

Chapitre 1 : Classification spectrale : algorithme et étude du paramètre

Cette première partie introduira l'algorithme de classification spectrale et le paramètre de la mesure d'affinité gaussienne. En effet, ce paramètre influence directement la séparabilité des données projetées dans l'espace de dimension réduite. Les diverses définitions du paramètre de l'affinité gaussienne seront présentées. Ensuite, nous proposerons une heuristique géométrique, comme paramètre global, différente sous l'hypothèse d'une distribution isotropique des points.

L'influence du paramètre sur la partition finale et notre heuristique seront alors testées à travers deux mesures de qualité sur des cas géométriques. Une première basée sur les ratios de normes de Frobenius permettra de mesurer relativement les affinités inter-clusters par rapport à celles intra-clusters. La seconde, supervisée, est une matrice de confusion évaluant le pourcentage exact de points mal-classés.

Enfin, nous appliquerons ces heuristiques sur des cas géométriques de densités, distributions et dimensions différentes et sur des cas de segmentations d'images. Nous étudierons par ailleurs deux différentes mesures d'affinité gaussienne pour considérer à la fois les coordonnées géométriques et la luminance.

Chapitre 2 : Classification et éléments spectraux de la matrice affinité gaussienne

Tout d'abord, nous rappellerons quelques propriétés sur les opérateurs Laplacien et de la chaleur avec condition de Dirichlet. Nous dégagerons ainsi une propriété de clustering sur les fonctions propres de ce dernier opérateur. Un lien sera établi entre le noyau Gaussien et le noyau de la chaleur avec conditions de Dirichlet et ainsi un premier résultat, asymptotique en t , sur les fonctions propres de l'opérateur de la chaleur en espace libre sur un support restreint sera énoncé.

Le passage à un ensemble discret de points s'effectuera par les éléments finis. Et un nouveau résultat, asymptotique en t et h (pas du maillage), sera établi dans l'espace d'approximation.

Enfin, le passage à la dimension finie par les éléments finis faisant apparaître une matrice de masse absente dans les algorithmes de clustering spectral, nous procéderons donc à une condensation de masse et définirons un nouveau résultat sur les vecteurs propres de la matrice affinité gaussienne.

Pour illustrer ces étapes théoriques, des exemples géométriques seront étudiés en partant de la version continue et les fonctions propres de l'équation de la chaleur avec conditions de Dirichlet jusqu'aux vecteurs propres de la matrice affinité Gaussienne non normalisée.

Chapitre 3 : Parallélisation de la classification spectrale

Il est proposé dans ce chapitre une nouvelle stratégie de parallélisation du clustering spectral basée sur une décomposition en sous-domaines. D'après le chapitre 2, les vecteurs propres issus de la matrice affinité gaussienne sont une représentation discrète des fonctions propres L^2 d'un opérateur de la chaleur dont le support est inclus sur une seule composante connexe. Une approche par décomposition en sous-domaines revient alors à restreindre le support de ces fonctions propres.

Tout d'abord, nous montrerons qu'appliquer le clustering spectral en sous-domaines géométriques pour identifier les composantes connexes aboutit à une partition "convenable" des données. L'étape de regroupement des partitions sur les divers sous-domaines est concrétisée par une zone de voisinage autour des frontières géométriques des sous-domaines. Ainsi, deux approches seront proposées et étudiées : le passage par une interface ou bien le recouvrement. Concrètement, le regroupement des clusters définis sur plusieurs sous-domaines sera réalisé via une relation d'équivalence.

Ensuite, nous proposerons une démarche pour déterminer automatiquement le nombre de clusters

k à partir de la mesure de qualité sur les normes de Frobenius du chapitre 1, et nous utiliserons l'heuristique définie au chapitre 1 pour fixer le paramètre d'affinité.

Enfin, nous testerons cette méthode sur des exemples géométriques et de segmentation d'images en variant le nombre de données et de découpe et nous en étudierons les limites.

Chapitre 4 : Extraction de connaissances appliquée à la biologie et l'imagerie médicale

Ce chapitre est consacré à la mise en application du matériel théorique et numérique dans un cadre biologique avec l'étude d'expressions de gènes issues de biopuces puis dans un cadre d'imagerie médicale avec la segmentation d'image issue de la tomographie par émission de positons.

Dans le cadre de la transcriptomie, des expérimentations basées sur l'inoculation d'une bactérie *Ralstonia Solanacearum* dans la légumineuse modèle *Medicago Truncatula* ont dévoilé l'existence de plantes résistantes (mutant HRP). L'observation des gènes concernés par la maladie foliaire et racinaire et l'analyse des gènes issus des plantes résistantes sont nécessaires. On a recours aux puces à ADN, et on observe, à divers instants, les niveaux d'expression d'une partie du génome. L'étude portera sur le classement des gènes suivant leur profil temporel d'expression.

Tout d'abord, nous introduirons une méthode de cartes auto organisatrices (ou Self Organizing Maps) [68] et l'adapterons à notre problématique en modifiant la mesure de similarité et les mises à jour de l'algorithme. Cette méthode, connue dans le traitement des données génomiques temporelles [98], servira par la suite de référence et d'outil de comparaison.

Ensuite, nous testerons la méthode de clustering spectral et nous comparerons les résultats par rapport à la méthode précédente.

Dans le cadre de l'imagerie médicale fonctionnelle quantitative, la Tomographie par Emission de Positons (TEP) dynamique permet de visualiser la concentration au cours du temps d'un traceur marqué par un atome radioactif en chaque point du cerveau étudié. Pour cela, les courbes temps-activité (TAC) calculées à partir de volume d'intérêt définis sur l'image sont utilisées pour quantifier la cible. Les données TEP sont une séquence d'images 3D temporelles traduisant l'évolution de la radioactivité dans le temps du volume correspondant correspondant au champ de vue de l'appareil. Une première étude est donc menée sur le spectral clustering appliqué aux TAC bruitées. Etant donné le caractère multidimensionnel des données, le choix du paramètre de l'affinité gaussienne et l'heuristique du choix du nombre de clusters sont étudiés en fonction de la mesure de qualité définie par les ratios de normes de Frobenius. Ensuite, le spectral clustering est appliqué sur des cas de segmentation d'images TEP dynamiques. Les résultats seront alors comparés à ceux issus du k -means, la méthode référence pour ce genre de données [112].

Chapitre 1

Classification spectrale : algorithme et étude du paramètre

Ce chapitre s'intéresse à la méthode de classification spectrale (ou clustering spectral) et à sa mise en oeuvre. Comme cette méthode repose sur la seule mesure d'affinité entre tous les couples de points, sans a priori sur les formes des classes (ou clusters), nous étudierons plus particulièrement, après une présentation de l'algorithme, le paramètre de l'affinité gaussienne. En effet, son rôle est crucial dans le partitionnement des données et il n'existe pas a priori de moyen pour définir un paramètre optimal, mais un ordre de grandeur peut être accessible. On propose donc deux heuristiques qui seront confrontées aux résultats théoriques dans le chapitre suivant. Dans un premier temps, les diverses définitions, globales et locales, basées sur des interprétations physiques seront présentées. Ensuite nous proposerons une heuristique basée sur un point de vue géométrique et nous introduirons une mesure de qualité pour étudier l'influence de ce paramètre sur les résultats de classification (ou clustering).

1.1 Présentation de la classification spectrale

Dans la suite, nous présentons un algorithme de spectral clustering et le choix du paramètre de l'affinité gaussienne sera étudié.

1.1.1 Algorithme de classification spectrale

La méthode de clustering spectral consiste à extraire les vecteurs propres associés aux plus grandes valeurs propres d'une matrice *affinité* normalisée, issue d'un noyau de Mercer [48]. Ces vecteurs propres constituent un espace de dimension réduite dans lequel les données transformées seront linéairement séparables. Deux principales classes d'algorithmes de clustering spectral ont été développées à partir de partitionnement de graphes [104]. La première est fondée sur un partitionnement bipartite récursif à partir du vecteur propre associé à la seconde plus grande valeur propre du graphe du Laplacien normalisé [63, 91], ou vecteur de Fiedler [25] dans le cas non-normalisé. La deuxième classe d'algorithmes n'utilise pas de manière récursive un seul vecteur propre mais propose de projeter les données originales dans un espace défini par les k plus grands vecteurs propres d'une matrice d'adjacence normalisée (ou matrice similaire à celle-ci), et d'appliquer un algorithme standard comme k -means sur ces nouvelles coordonnées [84, 79]. Nous porterons l'étude principalement sur cette dernière classe dans un souci de coût numérique et de simplicité algorithmique.

Y. Weiss et al (NJW) [84] présentent cette dernière classe d'algorithmes (c.f. Algorithme 1) pour partitionner un ensemble de points $S = \{x_1, \dots, x_N\} \subset \mathbb{R}^p$ en k clusters où k est fixé. NJW justifient

cet algorithme en considérant un cas idéal avec 3 clusters bien séparés. En partant de l'hypothèse que les points sont déjà ordonnés consécutivement par classe, la matrice affinité bien que dense, a alors une structure numérique proche de celle d'une matrice bloc-diagonale. Ainsi, la plus grande valeur propre de la matrice affinité normalisée est 1 avec un ordre de multiplicité égal à 3. Les lignes normalisées de la matrice constituée des vecteurs propres associés aux plus grandes valeurs propres sont constantes par morceaux. Dans l'espace défini par les $k = 3$ plus grands vecteurs propres de L , il est facile d'identifier 3 points bien séparés qui correspondent aux 3 constantes par morceaux des vecteurs propres, ce qui définit les clusters. Un exemple de cas idéal est illustré sur la figure 1.1, où les différentes étapes de l'algorithme de clustering spectral sont représentées.

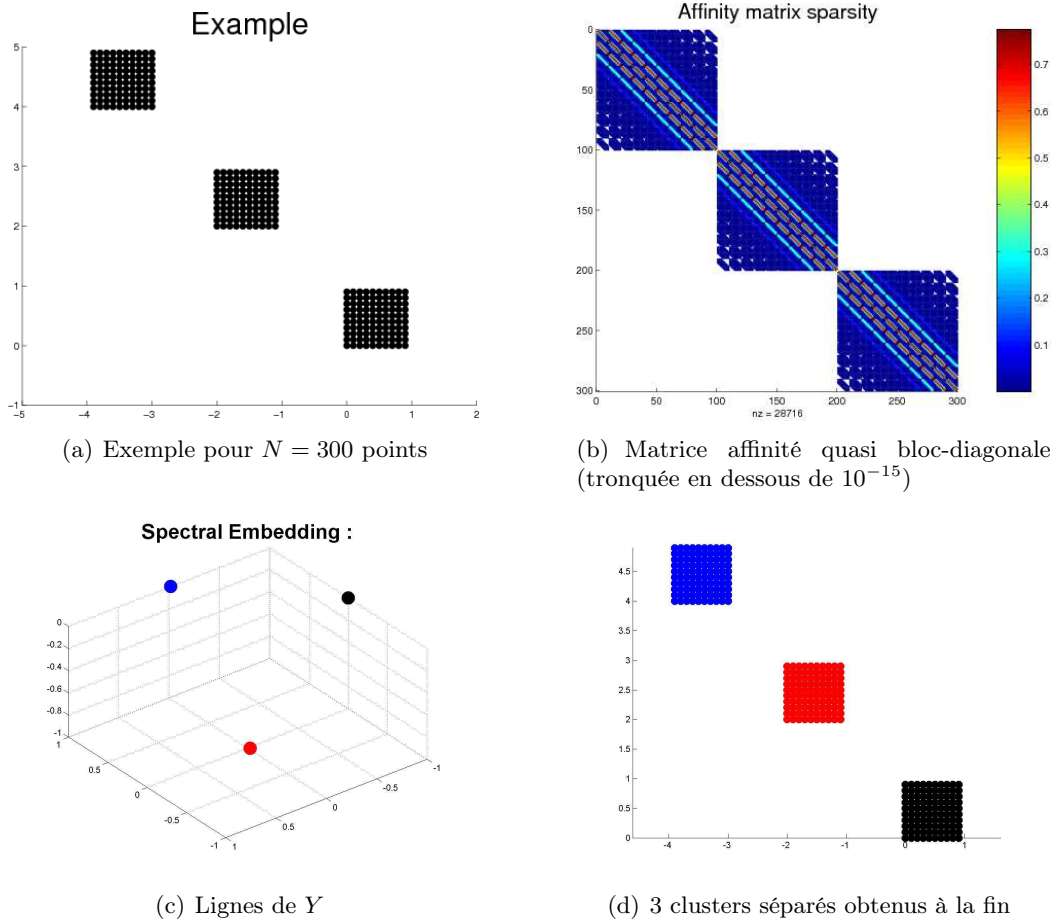


FIGURE 1.1 – Illustration des étapes du clustering spectral

Cependant, dans le cas général où les clusters ne sont pas nettement séparés par une grande distance, la structure numérique bloc-diagonale de la matrice affinité est compromise et peut être considérablement altérée par le choix de la valeur du paramètre de l'affinité (σ). Ce paramètre, en facteur

Algorithm 1 Algorithme de classification spectrale

Input : Ensemble des données S , Nombre de clusters k

1. Construction de la matrice affinité $A \in \mathbb{R}^{N \times N}$ définie par :

$$A_{ij} = \begin{cases} \exp(-\|x_i - x_j\|^2 / 2\sigma^2) & \text{si } i \neq j, \\ 0 & \text{sinon.} \end{cases}$$

2. Construction de la matrice normalisée $L = D^{-1/2}AD^{-1/2}$ où D matrice diagonale définie par :

$$D_{i,i} = \sum_{j=1}^N A_{ij}.$$

3. Construction de la matrice $X = [X_1 X_2 \dots X_k] \in \mathbb{R}^{N \times k}$ formée à partir des k plus grands vecteurs propres x_i , $i = \{1, \dots, k\}$ de L .
4. Construction de la matrice Y formée en normalisant les lignes de X :

$$Y_{ij} = \frac{X_{ij}}{\left(\sum_j X_{ij}^2\right)^{1/2}}.$$

5. Traiter chaque ligne de Y comme un point de $\mathbb{R}^k$ et les classer en k clusters via la méthode *K-means*.
 6. Assigner le point original x_i au cluster j si et seulement si la ligne i de la matrice Y est assignée au cluster j .
-

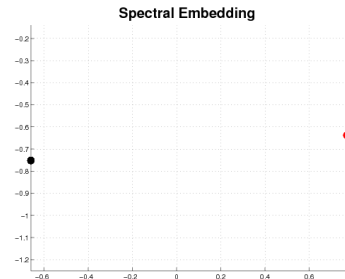
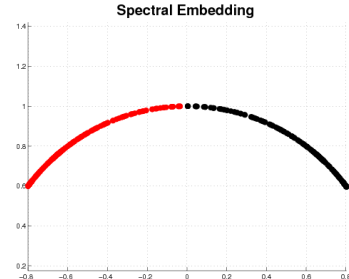
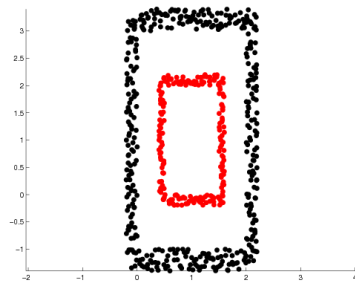
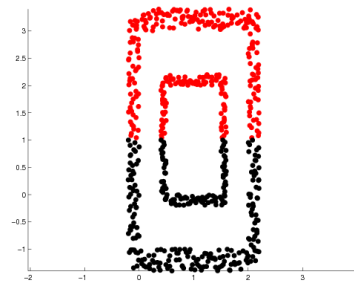
(a) $\sigma = 0.09$ (b) $\sigma = 0.8$ (c) Résultat du Clustering pour $\sigma = 0.09$ (d) Résultat du Clustering pour $\sigma = 0.8$

FIGURE 1.2 – Influence du paramètre dans l'espace de projection spectrale

de la norme entre chaque couple de points, sert de pondérateur. Il peut donc, suivant sa valeur, diminuer l'affinité intra-cluster et augmenter celle entre les clusters. Le choix de ce paramètre est donc étudié dans la suite.

1.1.2 Problème du choix du paramètre

Dans l'algorithme 1, la matrice d'affinité $A \in \mathcal{M}_{N,N}(\mathbb{R})$, A_{ij} désignant l'affinité entre les objets x_i et x_j , possède les propriétés suivantes :

$$\begin{cases} \forall (i, j), A_{ij} = A_{ji}, \\ \forall (i, j), 0 \leq A_{ij} \leq 1. \end{cases}$$

Parmi les noyaux de Mercer [7, 58], le noyau Gaussien est le plus communément choisi. L'affinité entre deux points distincts x_i et x_j de $\mathbb{R}^p$ est alors définie par :

$$A_{ij} = \begin{cases} \exp(-\|x_i - x_j\|_2^2 / 2\sigma^2) \text{ si } i \neq j, \\ 0 \text{ sinon,} \end{cases} \quad (1.1)$$

où σ est un paramètre et $\|\cdot\|_2$ est la norme euclidienne usuelle. L'expression de l'affinité gaussienne dépend donc d'un paramètre σ . Or, le principe du clustering spectral reposant sur la mesure d'affinité, ce paramètre influe directement sur la méthode. En effet, σ influe sur la séparabilité des données dans l'espace de projection spectrale (spectral embedding) comme indiqué dans l'exemple géométrique des figures 1.2 (a) et (b). L'espace de projection spectrale est tracé pour deux valeurs différentes de σ . Pour la première valeur $\sigma = 0.09$, les points projetés dans l'espace spectral forment deux paquets compacts où les points projetés associés au même cluster sont concentrés autour d'un seul et même point alors que, pour $\sigma = 0.8$, les points projetés dans cet espace décrivent un arc de cercle et aucune séparation ou agglomération de points ne sont distinguées. Dans ce cas, la méthode de k -means classique échoue et les clusters sont difficiles à déterminer. Il s'en suit les clusterings respectifs des figures 1.2 (c) et (d). Donc un mauvais choix de σ a des répercussions sur la séparabilité des clusters dans l'espace propre.

En segmentation d'images [91], ce paramètre est souvent laissé libre et doit être fixé par les utilisateurs. Cependant, plusieurs interprétations sur le rôle de ce paramètre ont été développées afin de guider le choix de manière automatique ou semi-automatique. Les principales méthodes sont les suivantes.

Seuillage

Une première approche repose sur l'idée de seuillage [85] via une analyse descriptive des données [84, 20] ou bien via la théorie des graphes [105]. En effet, *Perona et Freeman* [85] définissent σ comme une distance seuil en dessous de laquelle deux points sont considérés comme similaires et au-dessus de laquelle ils sont dissemblables.

Von Luxburg [105] interprète σ comme le rayon du voisinage. Il joue donc un rôle équivalent à celui du ϵ , distance de voisinage pour la méthode graphe à ϵ -voisinage. Ces interprétations donnent des informations sur le rôle que joue le paramètre σ . Cependant, le choix du seuil reste ouvert et non automatisable.

Approche globale

Une seconde approche repose sur des définitions globales de ce paramètre, principalement basées sur l'analyse descriptive des données [20] et des interprétations physiques. D'après *Ng, Jordan*,

Weiss [84], σ contrôle la similarité entre les données et conditionne la qualité des résultats. Ainsi, parmi les différentes valeurs de σ , on sélectionne celle qui minimise la dispersion des points d'un même cluster dans l'espace propre réduit. Cette sélection peut être implémentée automatiquement de manière non supervisée en considérant la dispersion des projections dans l'espace propre comme indicateur. Ce choix peut être défini à l'aide d'un histogramme sur les normes $\|x_i - x_j\|_2$ entre tous les points x_i, x_j , pour tout $i, j \in \{1, \dots, N\}$. En supposant l'existence de clusters, l'histogramme deviendra multi-modal : le premier mode correspondra à la moyenne intra-cluster et les suivants représenteront les distances entre les clusters. Sélectionner un σ de l'ordre du premier mode de l'histogramme revient donc à privilégier les affinités au sein des clusters et donc la structure bloc-diagonale de la matrice affinité.

Brand et Huang [20] définissent un paramètre scalaire global semblable à l'heuristique sur les histogrammes. En effet, σ doit être égal à la moyenne de la distance entre chaque point de l'ensemble S et son plus proche voisin. Cette heuristique est testée sur divers exemples géométriques 2D aux densités variées et les résultats du clustering sont représentés sur la figure 1.3. Sur certains exemples, cette estimation peut s'avérer insuffisante notamment lorsque les densités au sein même d'un cluster varient comme pour les exemples (b) et (e) de la figure 1.3. De plus, cette définition requiert de faire une boucle sur tous les points $x_i \in S$ pour résoudre $\min_j \|x_i - x_j\|_2$, pour tout $i \in \{1, \dots, N\}$, ce qui peut être coûteux numériquement dans le cas d'un nombre important de données x_j en particulier.

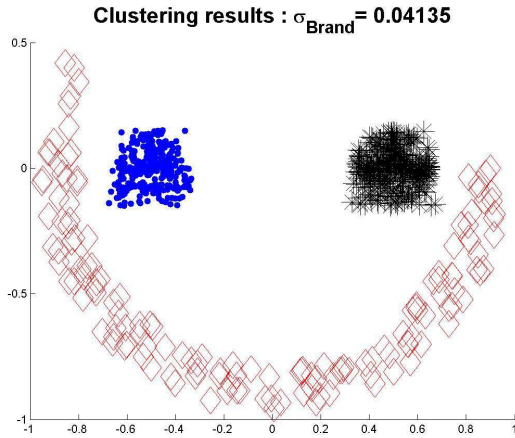
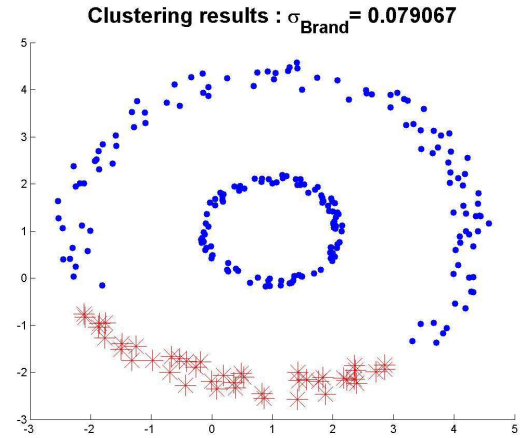
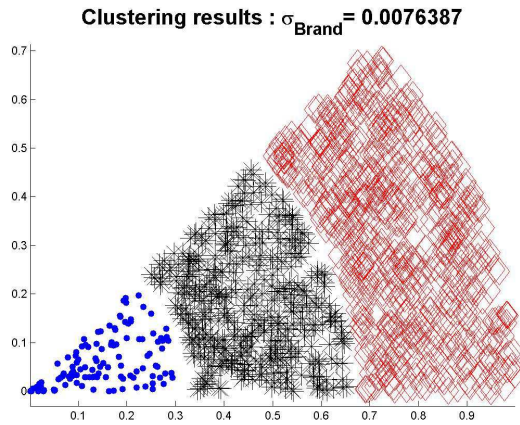
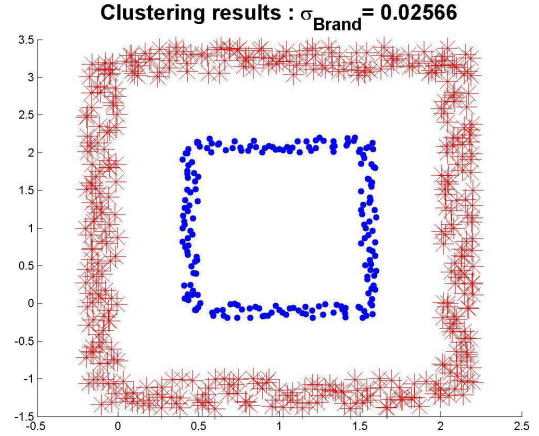
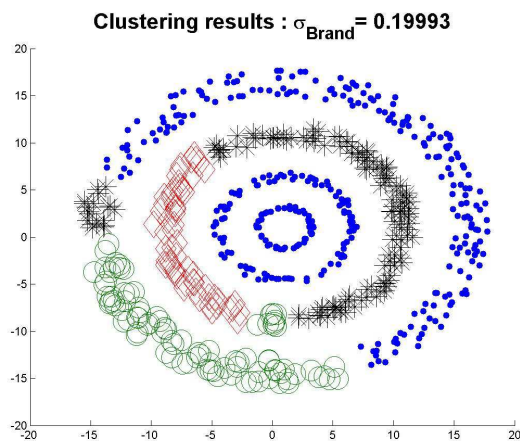
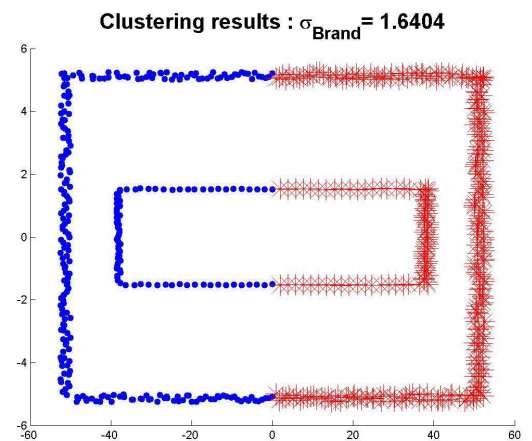
Approche locale

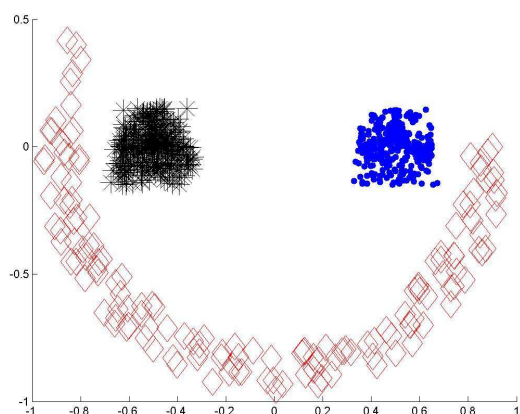
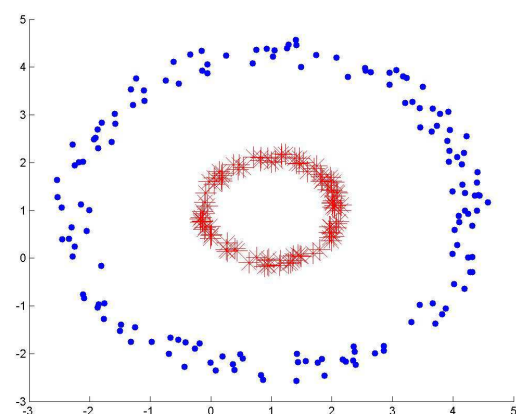
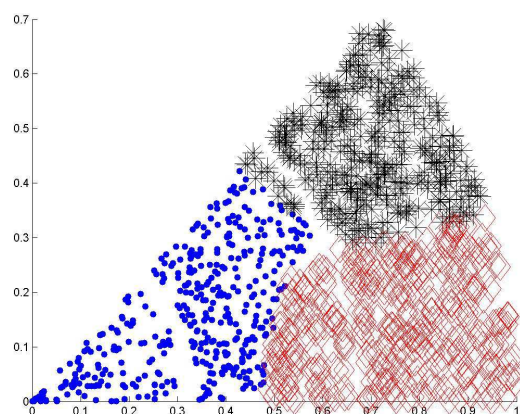
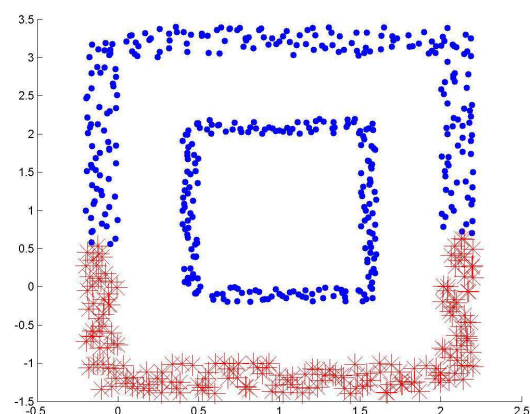
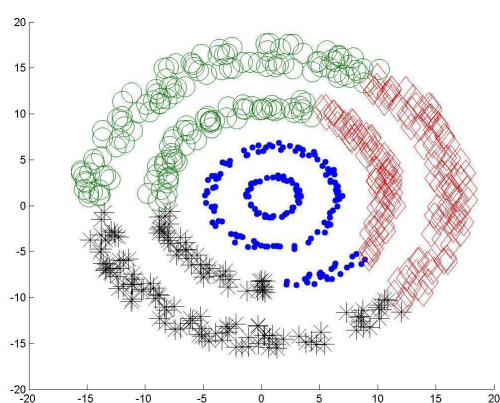
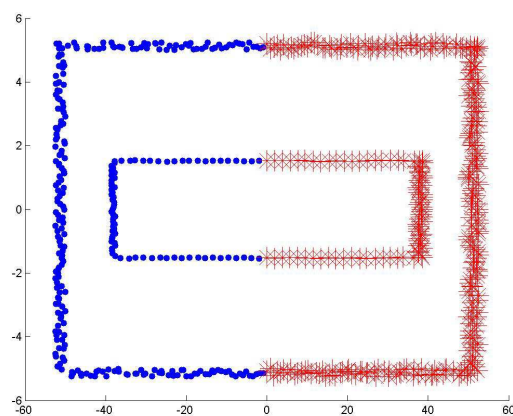
Une dernière classe de définitions basées sur des interprétations physiques [116, 42] privilégie des approches locales où un paramètre scalaire spécifique est défini pour chaque point x_i . Perona et Zelnik-Manor [116] ont adopté une approche locale, consistant à définir un scalaire pour chaque couple de point x_i, x_j . Ils assignent un paramètre scalaire σ_i différent à chaque point x_i de l'ensemble S . σ_i est égal à la distance entre le point x_i et son $P^{\text{ième}}$ voisin le plus proche. Cette méthode donne de bons résultats dans certains cas où l'effet de l'analyse locale fournit assez d'informations pour créer les clusters : par exemple, des clusters compacts plongés dans du bruit. Mais calculer une valeur de σ pour chaque point x_i peut être coûteux et la valeur P reste fixée empiriquement (à $P = 7$ pour [116]).

Une autre approche locale développée par Fischer et Poland [42] utilise la conductivité définie dans les réseaux électriques : la conductivité entre deux points dépend alors de tous les chemins entre eux. Cette définition permet de renforcer la structure numérique par bloc de la matrice affinité. Le paramètre local est fixé tel que la somme des lignes de la matrice affinité A est égale à une valeur τ . Donc σ dépend d'une autre valeur τ représentant un rayon de voisinage à fixer empiriquement.

Remarque 1.1. Ces définitions locales s'avèrent très efficaces pour des cas de données bruitées. Elles permettent de distinguer le bruit des données comme le montre la figure 1.4. Cependant, ces approches s'avèrent coûteuses et impliquent de définir de nouveaux paramètres, respectivement le nombre de voisins P pour [116] ou la valeur τ pour [42]. Ces derniers représentent le raffinement de l'étude locale.

Dans les exemples introduits sur la figure 1.4, la densité des points varie au sein des clusters. Ces résultats illustrent le fait que, sans l'information de densité globale, il peut être difficile de classer

(a) Smiley : $\sigma_{\text{Brand}} = 0.041$ (b) Cercles : $\sigma_{\text{Brand}} = 0.079$ (c) Portions de couronnes : $\sigma_{\text{Brand}} = 0.007$ (d) Carrés : $\sigma_{\text{Brand}} = 0.025$ (e) Cible : $\sigma_{\text{Brand}} = 0.199$ (f) Rectangles étirés : $\sigma_{\text{Brand}} = 1.640$ FIGURE 1.3 – Paramètre global : exemples avec le paramètre σ proposé par Brand [20]

(a) Smiley : $n = 790$ (b) Cercles : $n = 250$ points(c) Portions de couronnes : $n = 1200$ (d) Carrés : $n = 600$ (e) Cible : $n = 650$ (f) rectangles étirés : $n = 640$ FIGURE 1.4 – Paramètre local : exemples avec le paramètre σ proposé par *Perona* [116]

correctement les points dans certains cas.

Dans le cadre des applications génomiques et d'imagerie médicale, on a affaire à de grandes masses de données et nous ne disposons d'aucune information a priori sur les clusters. Donc le choix du paramètre σ ne doit pas être coûteux numériquement et doit fournir des informations sur la distribution des points. Une approche globale est donc envisagée pour définir une nouvelle heuristique.

1.2 Heuristiques du paramètre

A l'instar des méthodes probabilistes [37], on souhaite intégrer dans la définition du paramètre σ à la fois la dimension du problème et la notion de densité de points dans l'ensemble des données p -dimensionnelles. Pour garder l'efficacité de la méthode, une approche globale est privilégiée où le paramètre est fonction des distances entre les points. Dans ce but, l'ensemble de points p -dimensionnels est supposé suffisamment isotropique dans le sens où il n'existe pas de direction privilégiée avec de grandes différences de grandeurs dans les distances entre des points sur une direction. Deux distributions isotropiques avec diverses amplitudes suivant les directions sont distinguées. Soit $S = \{x_i, 1 \leq i \leq N\}$ un tel ensemble de points et

$$D_{\max} = \max_{1 \leq i, j \leq N} \|x_i - x_j\|_2,$$

la plus grande distance entre toutes les paires de points dans S .

1.2.1 Cas d'une distribution isotropique avec mêmes amplitudes suivant chaque direction

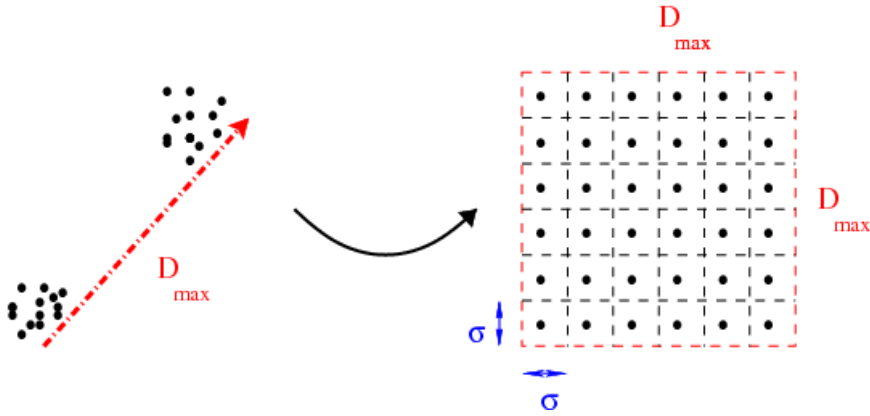


FIGURE 1.5 – Principe de l'heuristique (1.2) dans le cas 2D

Heuristique 1.2. Soit un ensemble de points $S = \{x_i, 1 \leq i \leq N\}$ dont la distribution est isotropique. Les éléments de S sont inclus dans une boîte carrée de dimension p de côté $D_{\max} = \max_{1 \leq i, j \leq N} \|x_i - x_j\|_2$. Alors, on définit une distance référence, notée σ_0 , définie par :

$$\sigma_0 = \frac{D_{\max}}{N^{\frac{1}{p}}} \quad (1.2)$$

où N est le nombre de points et p la dimension des données. Enfin, le paramètre gaussien σ est pris égal à une fraction de cette distance référence σ_0 : $\sigma = \frac{\sigma_0}{2}$.

Sous cette hypothèse d'isotropie de la distribution, on peut statuer que l'ensemble des points S est essentiellement inclus dans une boîte carrée de dimension p et de côté borné par $D_{\max}$. Si on désire être capable d'identifier des clusters au sein de S , on doit partir d'une distribution uniforme de n points inclus dans la boîte de dimension p . Cette distribution uniforme est obtenue en divisant la boîte en N briques de même taille comme représenté sur la figure 1.5 dans le cas 2D. Chaque brique est de volume égal à $D_{\max}^p/N$ et le côté des briques, noté σ_0 , est donc égal à :

$$\sigma_0 = \frac{D_{\max}}{N^{1/p}}$$

Dès lors, on peut considérer que s'il existe des clusters, il doit exister des points qui sont séparés par une distance supérieure à cette fraction σ . Sinon, les points devraient tous être à distance à peu près égale (de l'ordre de σ_0) de leurs plus proches voisins d'après l'hypothèse sur l'isotropie de la distribution des points et du fait que tous les points sont inclus dans une boîte de côté $D_{\max}$. Notre proposition consiste donc à construire une matrice affinité comme une fonction ratio de distances entre les points et une distance de référence

$$\sigma = \frac{\sigma_0}{2}$$

Remarque 1.3. Cette heuristique représente donc une distance seuil à partir de laquelle des points sont considérés comme proches. Cette interprétation rejoint aussi celles de Brand et Huang [20] qui fixent σ égal à la moyenne de la distance entre chaque point de l'ensemble S et son plus proche voisin. Cependant, cette nouvelle heuristique est moins coûteuse que l'heuristique [20] car une seule boucle est réalisée sur tous les points pour déterminer $D_{\max}$.

1.2.2 Cas d'une distribution isotropique avec des amplitudes différentes suivant les directions

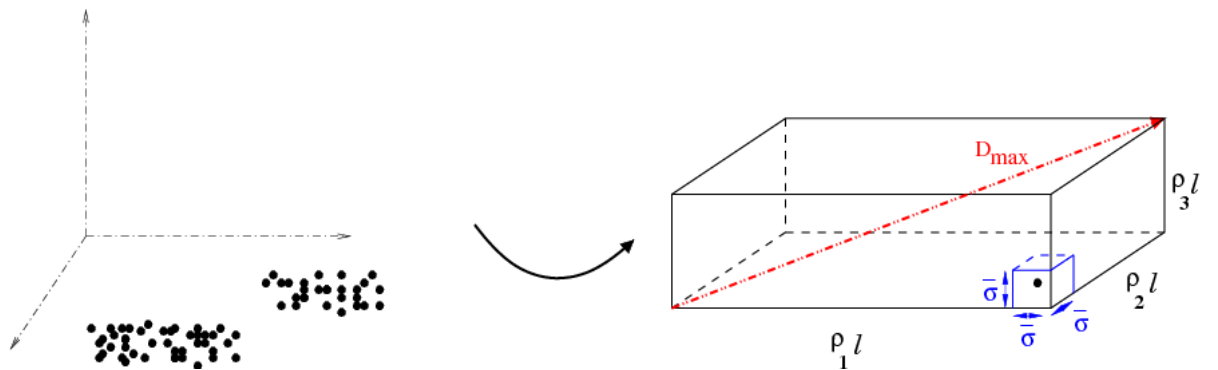


FIGURE 1.6 – Principe de l'heuristique (1.3) dans le cas 3D

Heuristique 1.4. Soit un ensemble de points $S = \{x_i, 1 \leq i \leq N\}$ dont la distribution est isotropique avec des amplitudes différentes suivant les directions. Chaque élément de S est inclus dans une boîte rectangulaire de dimension p de côté ρ_k pour la k -ième dimension défini par :

$$\rho_k = \max_{1 \leq i \leq p} x_i(k) - \min_{1 \leq j \leq p} x_j(k), k \in \{1, \dots, N\}, \text{ pour } k \in \{1, \dots, p\}.$$

Alors, le paramètre gaussien σ est une fraction de la distance référence, notée $\bar{\sigma}_0$, définie par :

$$\bar{\sigma}_0 = \frac{D_{\max} \sqrt{p}}{\|\rho\|_2} \left(\frac{\prod_{i=1}^p \rho_i}{N} \right)^{\frac{1}{p}}, \quad (1.3)$$

où n est le nombre de points et p la dimension des données. Alors, le paramètre gaussien σ_2 est égal à une fraction de la distance référence $\bar{\sigma}_0$:

$$\sigma_2 = \frac{\bar{\sigma}_0}{2}.$$

Sous l'hypothèse que l'ensemble des données de dimension p est suffisamment isotropique, il peut exister des directions dans les données avec des variations d'amplitudes. Dans ce cas, le calcul de σ est adapté en considérant que l'ensemble des points est inclus dans une boîte rectangulaire de dimension p dont les côtés sont proportionnels aux amplitudes suivant chaque direction comme le représente la figure 1.6 dans le cas 3D. Pour définir toutes les dimensions des côtés de cette nouvelle boîte, on calcule alors la plus grande distance entre toutes les paires de points appartenant à S suivant chacune des directions, notées ρ_k (pour la k^{ime} dimension), et donnée par :

$$\rho_k = \max_{1 \leq i \leq p} x_i(k) - \min_{1 \leq j \leq p} x_j(k), k \in \{1, \dots, N\}.$$

Le vecteur $\rho = (\rho_1, \dots, \rho_p)^T$ incorpore les tailles des intervalles dans lesquels chaque variable est incluse. Dans ce cas, le côté de la boîte rectangulaire reste dans le même esprit que précédemment. Il est donc fonction du vecteur ρ et du diamètre maximal $D_{\max}$. La distance $D_{\max}$ est alors égale, d'après le théorème de Pythagore à : $D_{\max}^2 = \sum_{i=1}^p \rho_i^2 l^2$,

le facteur l étant égal au ratio entre $D_{\max}$ et la norme euclidienne $\|\rho\|_2$: $l = \frac{D_{\max}}{\|\rho\|_2}$.

Le volume est alors décomposé en N volumes cubiques de côté $\bar{\sigma}_0$ défini par : $\bar{\sigma}_0 = \frac{D_{\max} \sqrt{p}}{\|\rho\|_2} \left(\frac{\prod_{i=1}^p \rho_i}{N} \right)^{\frac{1}{p}}$ où le facteur $\sqrt{p}$ permet de retrouver l'équation (1.2) quand ρ est constant et quand la boîte est carrée.

Une autre façon d'envisager ce cas de distribution aurait été d'utiliser les distances de Mahalanobis. En effet, ces distances permettent de calculer l'orientation spectrale de la dispersion des données en fixant les axes principaux et en calculant les amplitudes sur ces axes. Mais cette étape est coûteuse numériquement et elle repose sur la matrice de variance-covariance donc sur l'hypothèse que les points sont corrélés entre eux. De plus, dans le cas de données corrélées, une étape préliminaire par Analyse en Composante Principale est souvent utilisée.

Remarque 1.5. Dans les deux configurations de distributions isotropiques, ces heuristiques restent sensibles aux artefacts, au bruit et aux densités fortement variables localement. Une possibilité reste d'appliquer plusieurs fois le clustering spectral dans le cas d'artefact, en modifiant le $D_{\max}$ à chaque étape. Cependant, l'approche locale reste privilégiée dans le cas de bruitage de données entre les clusters et suppose donc une étude plus spécifique de certains clusters.

1.3 Validations numériques

Afin de valider les heuristiques (1.2) et (1.3), elles sont testées sur les 6 exemples géométriques 2D des figures 1.3 et 1.4 :

- (a) smiley constitué de $N = 790$ points et de $k = 3$ clusters,
- (b) 2 cercles concentriques avec $N = 250$ points et $k = 2$,
- (c) 3 portions de couronnes concentriques de $N = 1200$ points et $k = 3$,
- (d) 2 carrés concentriques $N = 600$ points,
- (e) une cible de $N = 650$ points constituée de $k = 4$ couronnes ;
- (f) 2 rectangles étirés de $N = 640$ points.

Ces exemples sont utilisés par [84, 20, 116] car ils représentent des domaines non-convexes, aux densités inter-clusters et intra-cluster variables et faisant échouer des méthodes classiques comme la méthode k -means. Le résultat du clustering pour chaque exemple est présenté sur les figures 1.7 et 1.8. Les heuristiques (1.2) et (1.3) partitionnent correctement les exemples géométriques (a) à (e) et leurs valeurs sont très proches voire égales car la distribution est isotropique avec approximativement les mêmes amplitudes suivant chaque direction. Par contre, l'exemple (f) représente un cas où la distribution est isotropique avec des amplitudes différentes d'un facteur 10 suivant les directions. Seule l'heuristique (1.3) partitionne correctement. En effet, l'adaptation des dimensions de la boîte à la distribution des points a divisé par deux la valeur de (1.2).

Pour l'étude plus fine du paramètre, il faut maintenant introduire des critères, des mesures de qualités, permettant de valider ces heuristiques et de montrer l'influence du paramètre sur les résultats.

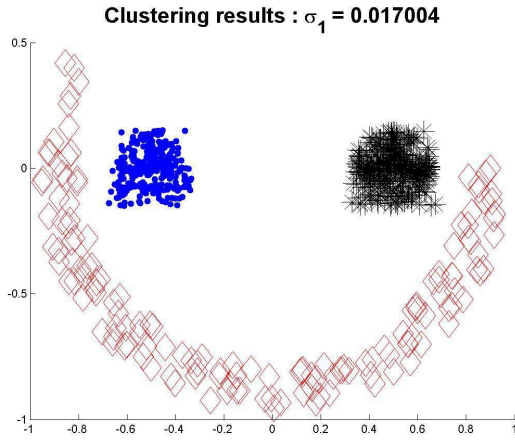
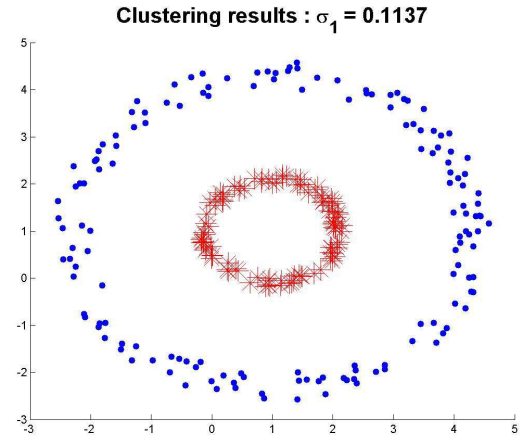
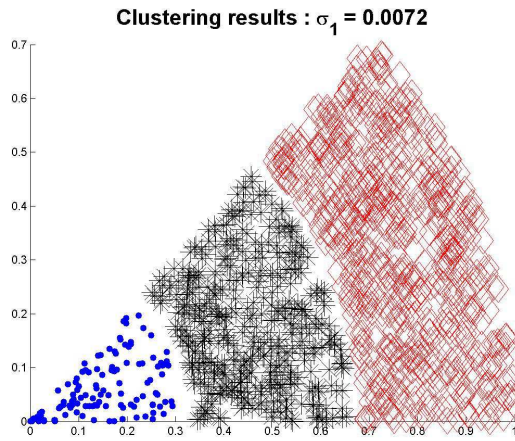
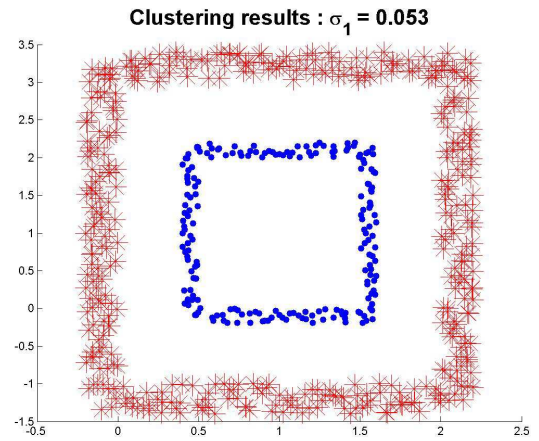
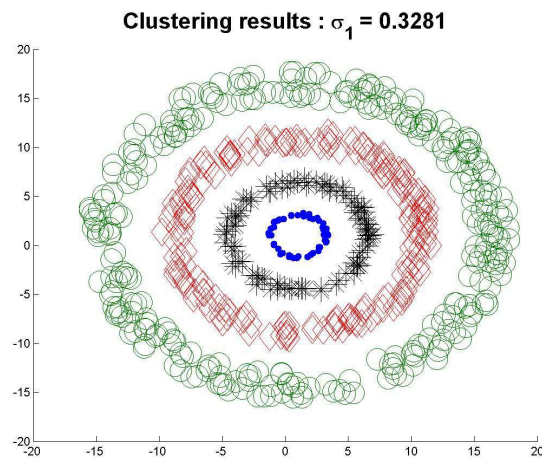
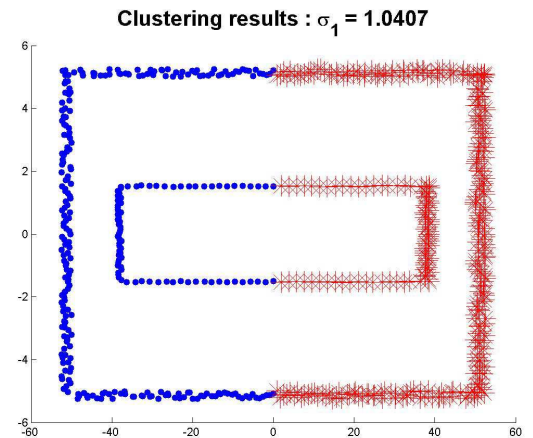
1.3.1 Mesures de qualité

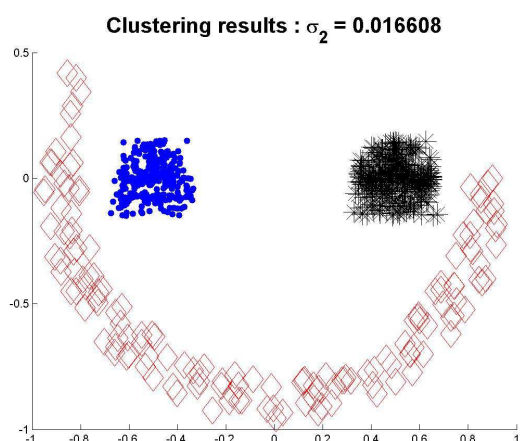
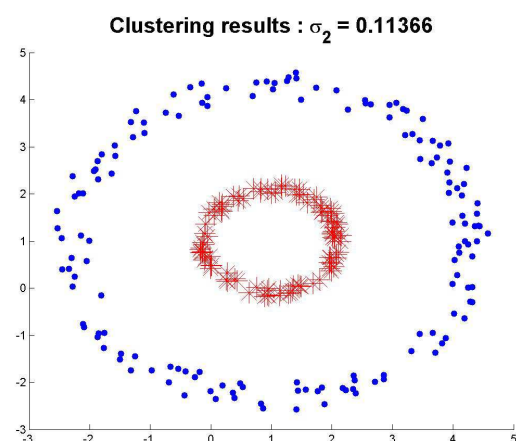
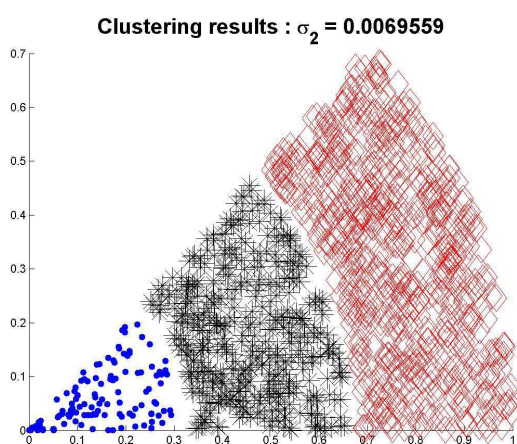
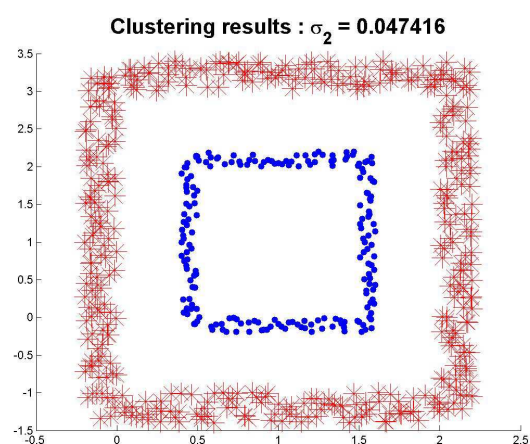
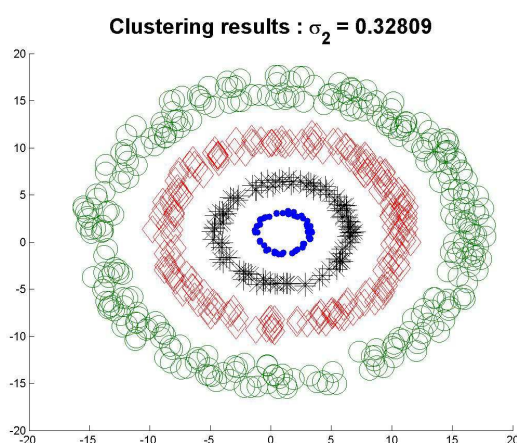
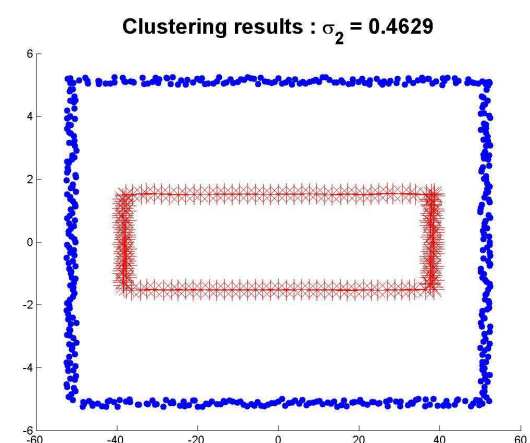
Plusieurs critères pour évaluer l'efficacité du résultat de clustering existent. Parmi eux, *Meila* [78] introduit la Variation d'Information (VI) pour comparer deux clusterings. La VI est une métrique évaluant la quantité d'information gagnée ou perdue d'un cluster à un autre utilisant l'entropie associée à un cluster et l'information mutuelle. La différence entre deux clusters peut être mesurée avec l'indice de Wallace introduit par *Wallace* [109]. Elle consiste à calculer la probabilité (donc à valeur dans $[0, 1]$) qu'un couple de points soit correctement classé. Cet indice donne la valeur 1 si le clustering n'a pas d'erreur.

Dans cette section, nous nous intéresserons à la matrice de confusion de *Verma et Meila* [104] permettant d'évaluer le pourcentage exact de points mal classés. Puis nous définissons une nouvelle mesure de qualité basée sur les normes de Frobenius de blocs d'affinités pour comparer les affinités entre les clusters et celles intra-clusters. Dans les deux cas, nous étudierons sur deux exemples géométriques l'évolution de la qualité en fonction des valeurs du paramètre σ .

Matrice de Confusion

Introduite par *Verma et Meila* [104], la matrice de confusion évalue l'erreur réelle de clustering c'est-à-dire le nombre de points mal assignés au sein des clusters. Elle suppose donc que les clusters sont connus a priori.

(a) Smiley : $\sigma = 0.017$ (b) Cercles : $\sigma = 0.113$ (c) Portions de couronnes : $\sigma = 0.007$ (d) Carrés : $\sigma = 0.053$ (e) Cible : $\sigma = 0.328$ (f) rectangles étirés : $\sigma = 1.040$ FIGURE 1.7 – Exemples avec l'heuristique 1 (σ) pour $p = 2$

(a) Smiley : $\sigma = 0.016$ (b) Cercles : $\sigma = 0.113$ (c) Portions de couronnes : $\sigma = 0.0069$ (d) Carrés : $\sigma = 0.047$ (e) Cible : $\sigma = 0.328$ (f) rectangles étirés : $\sigma = 0.4629$ FIGURE 1.8 – Exemples avec l'heuristique 2 (σ) pour $p = 2$

Mesure de qualité 1.6. Soit k le nombre de clusters. Après avoir appliqué le spectral clustering pour une valeur de k , on définit la matrice de confusion, notée $C \in \mathcal{M}_{k,k}(\mathbb{R})$, de la façon suivante : les éléments C_{ij} définissent le nombre de points qui sont assignés au cluster j au lieu du cluster i pour $i \neq j$ et C_{ii} le nombre de points correctement assignés pour chaque cluster i . On définit alors un pourcentage de points mal-classés, noté P_{erreur} , par :

$$P_{\text{erreur}} = \frac{\sum_{i \neq j}^k C_{ij}}{N}$$

où N est le nombre de points et k le nombre de clusters.

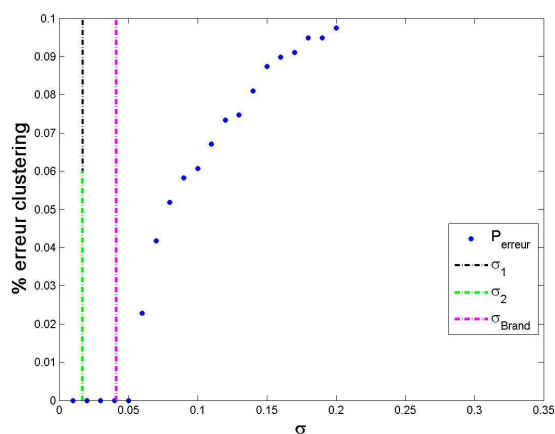
Le pourcentage d'erreur P_{erreur} issu de la matrice de confusion donne une estimation de l'erreur réelle dans la méthode de clustering. Cette mesure est donc testée sur les exemples géométriques précédemment présentés et, sur la figure 1.9, le pourcentage d'erreur P_{erreur} est tracé en fonction des valeurs de σ . Sur certains exemples comme (b), (e) et (f), les premières valeurs de σ ne sont pas testées car pour ces valeurs proches de 0, le conditionnement de la matrice affinité A est mauvais (supérieur à 10^{13}) ce qui ne permet pas de faire converger les algorithmes de recherche de valeurs propres et vecteurs propres. De plus, les valeurs de σ supérieures à l'intervalle considéré pour chaque exemple ne présentent pas d'intérêt car le pourcentage d'erreur P_{erreur} reste supérieur ou égal à celui de la dernière valeur de σ représentée sur la figure 1.9. Les lignes verticales noire, verte et magenta en pointillés indiquent respectivement la valeur du paramètre heuristique (1.2), celle du paramètre heuristique (1.3) et celle définie par Brand [20]. Suivant les exemples, l'intervalle sur lequel il n'y a pas d'erreur de clustering varie considérablement d'un cas à l'autre : par exemple, la longueur de l'intervalle peut être de l'ordre de 0.4 pour (b) ou être inférieure à 0.1 pour (a) et (c). En effet, le pourcentage d'erreur P_{erreur} varie instantanément quand σ n'appartient plus à l'intervalle adéquat. Comparées aux résultats numériques de la figure 1.3, les valeurs d'heuristiques pour lesquelles le partitionnement est incorrect appartiennent à l'intervalle où P_{erreur} est supérieure à 0%. Les valeurs des heuristiques (1.2) et (1.3) correspondent à une valeur de σ avec une erreur de clustering nulle exceptée pour l'heuristique (1.2) avec l'exemple des deux rectangles étirés figure 1.9 (f). Cette mesure valide donc l'influence du paramètre ainsi que les résultats numériques des figures 1.3, 1.7 et 1.8 pour les différentes heuristiques.

Ratio de normes de Frobenius

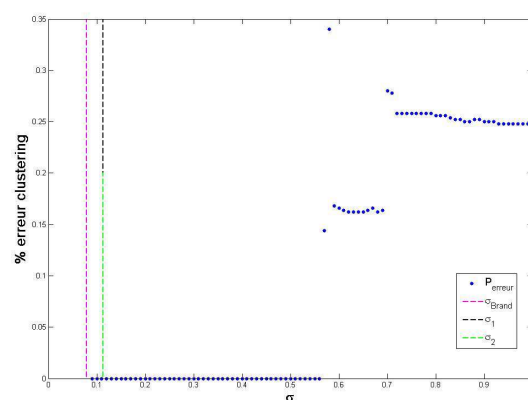
La mesure par matrice de confusion donne un très bon outil d'analyse de la qualité du cluster. Elle demande cependant de connaître l'état exact du clustering à obtenir et ne peut donc pas être utilisée pour des applications non supervisées. En particulier, on cherche à évaluer de manière automatique le bon nombre de clusters. Pour ce faire, on propose d'introduire une autre mesure de qualité calculée directement à partir des données internes au calcul. Après validation, cette mesure sera introduite par la suite comme outil de la stratégie parallèle présentée au chapitre 4.

Mesure de qualité 1.7. Après avoir appliqué le spectral clustering pour un nombre de clusters k à déterminer mais que l'on fixe a priori, la matrice affinité A définie par (1.1) est réordonnée par cluster. On obtient la matrice par bloc, notée L , telle que les blocs hors diagonaux représentent les affinités entre les clusters et les blocs diagonaux l'affinité intra-cluster. On évalue les ratios entre les normes de Frobenius des blocs diagonaux et ceux hors-diagonaux pour $i, j \in 1, \dots, k$ et $i \neq j$:

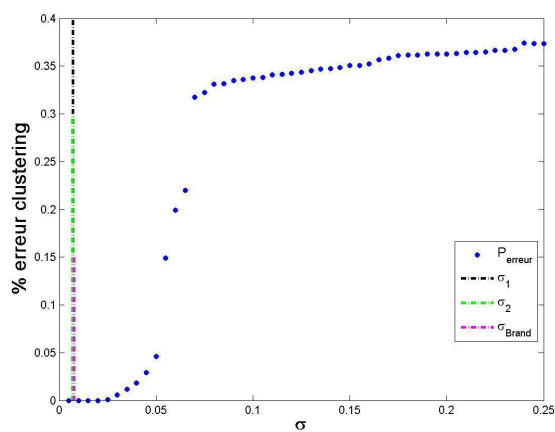
$$r_{ij} = \frac{\|L^{(ij)}\|_F}{\|L^{(ii)}\|_F}, \quad (1.4)$$



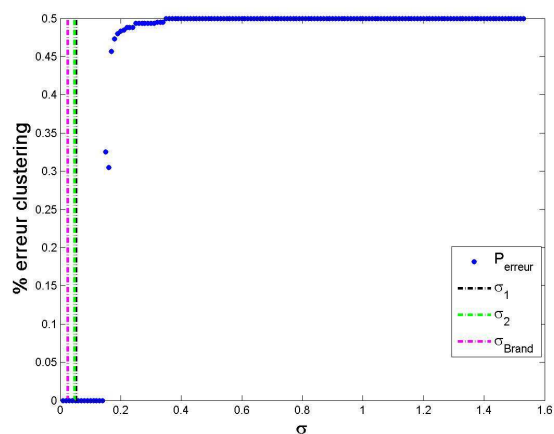
(a) Smiley



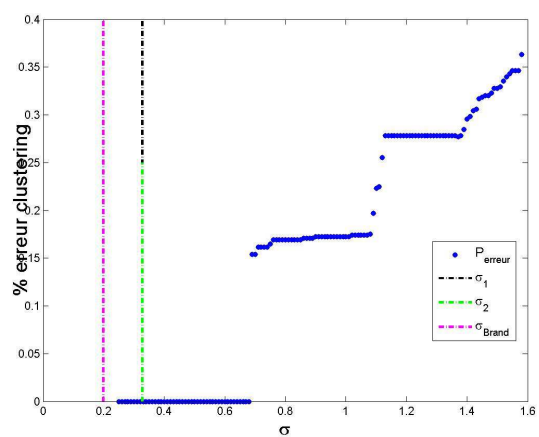
(b) Cercles



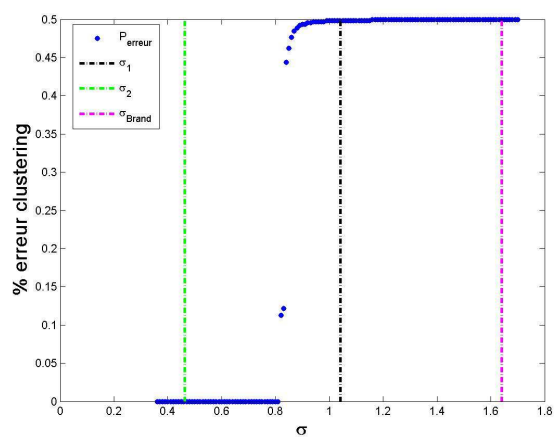
(c) Portions de couronnes



(d) Carrés



(e) Cible



(f) Rectangles étirés

FIGURE 1.9 – Pourcentages d'erreur de clustering en fonction de σ

où $\|L^{(ij)}\|_F = \left(\sum_{l=1}^{N_i} \sum_{m=1}^{N_j} |L_{lm}^{(ij)}|^2 \right)^{\frac{1}{2}}$ et où N_i et N_j sont les dimensions du bloc (ij) de L .

Si le ratio est proche de 0 alors la matrice affinité réordonnée par cluster a une structure bloc-diagonale, proche du cas idéal et donc le clustering obtenu est bon.

Prenons l'exemple du cas idéal de la figure 1.1 où l'on considère 3 blocs séparés d'une distance d . Avec $k = 3$, la matrice suivante L s'écrit donc de la façon suivante :

$$L = \begin{bmatrix} L^{(11)} & L^{(12)} & L^{(13)} \\ L^{(21)} & L^{(22)} & L^{(23)} \\ L^{(31)} & L^{(32)} & L^{(33)} \end{bmatrix}$$

En notant $d(x_l, x_m)$ la distance séparant x_l et x_m , deux points de S appartenant à deux clusters différents, on définit par ϵ_{lm} la distance telle que $\epsilon_{lm} = d(x_l, x_m) - d$ où d est la distance de séparation entre les blocs. Pour $i \neq j$, la norme de Frobenius du bloc hors-diagonal $L^{(ij)}$ est majorée par inégalité triangulaire par :

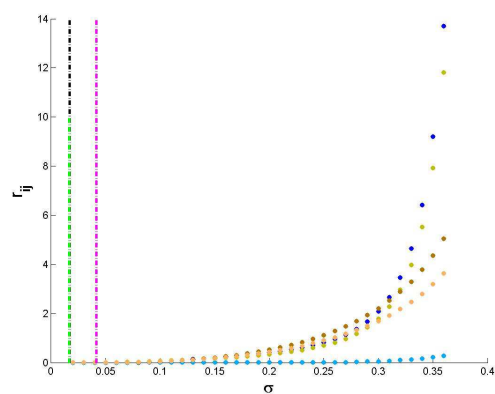
$$\|\hat{L}^{(ij)}\|_F^2 = \sum_{l=1}^{N_i} \sum_{m=1}^{N_j} e^{-\frac{\|d + \epsilon_{lm}\|_2^2}{\sigma^2}} \leq e^{-\frac{d^2}{\sigma^2}} \sum_{l=1}^{N_i} \sum_{m=1}^{N_j} e^{-\frac{\|\epsilon_{lm}\|_2^2}{\sigma^2}}.$$

Pour $i = j$, les points de S appartenant au même cluster sont séparés par une distance homogène à ϵ_{lm} . Donc, le ratio r_{ij} est fonction de $t \mapsto \exp(-\frac{d^2}{t^2})$ qui tend vers 0 lorsque t tend vers 0.

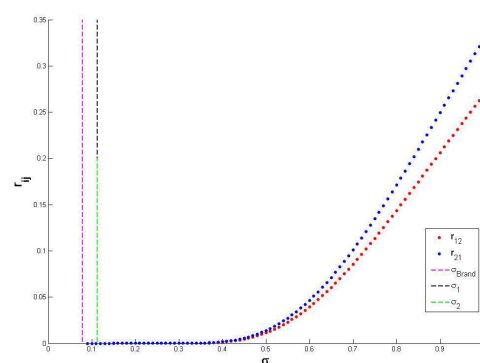
Cette mesure traduit par elle-même le principe du clustering : si le ratio est proche de 0 alors des points appartenant à des clusters différents seront le moins semblable et des points appartenant au même cluster le plus semblable possible. Si le ratio r_{ij} est proche de 0, la matrice affinité a une structure quasi bloc-diagonale. Cette situation correspond dans l'espace spectral à des clusters concentrés et séparés. Dans le cas général, les blocs hors-diagonaux de la matrice affinité normalisée L ne sont pas des blocs nuls. Dans la figure 1.10 où l'on considère tous les exemples géométriques précédemment présentés, les valeurs des ratios r_{ij} en fonction des valeurs de σ sont tracées pour les diverses valeurs de $(i, j) \in \{1, \dots, k\}^2$ avec $i \neq j$. Les lignes verticales noire, verte et magenta en pointillés indiquent respectivement la valeur du paramètre heuristique (1.2), celle du paramètre heuristique (1.3) et celle définie par Brand [20]. A l'instar de la mesure de qualité basée sur la matrice de confusion, les valeurs de σ proches de 0 ne sont pas testées pour les figures (b), (e) et (f) à cause du mauvais conditionnement de la matrice affinité A (supérieur à 10^{13}).

D'après les variations de ces mesures suivant les valeurs de σ , les intervalles sur lesquels la matrice affinité s'approche d'une structure bloc diagonale coïncident avec ceux du pourcentage d'erreur de la précédente mesure de qualité. Cet intervalle dépend de la nature du problème et diffère suivant les cas comme on l'a observé avec la précédente mesure de qualité. Les résultats sur les six exemples géométriques montrent que l'intervalle de valeurs pour un choix approprié du paramètre σ est approximativement le même que pour la mesure de qualité basée sur le ratio matrice de confusion.

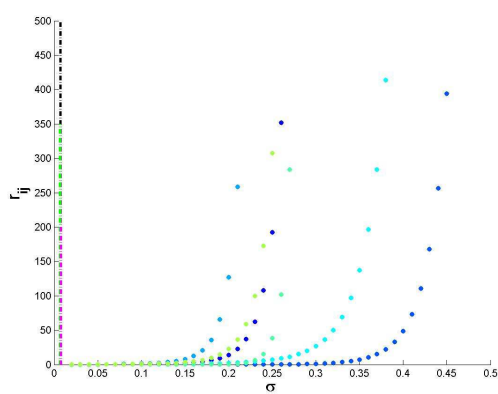
Remarque 1.8. Cette mesure permet donc d'évaluer la partition finale du clustering spectral à partir de l'affinité entre les points. Etant non supervisée par nature, cette mesure peut être utilisée pour déterminer le nombre de clusters k : le critère à minimiser serait un ratio moyen sur tous les blocs de la partition pour diverses valeurs de k (cf chapitre 4).



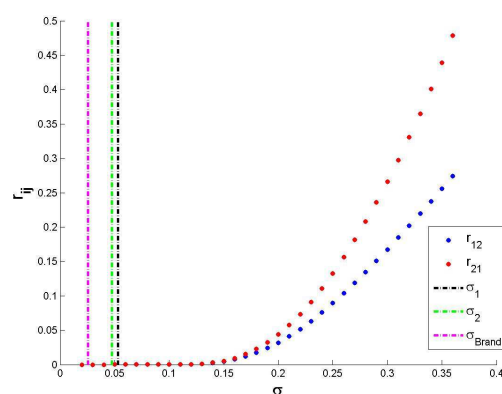
(a) Smile



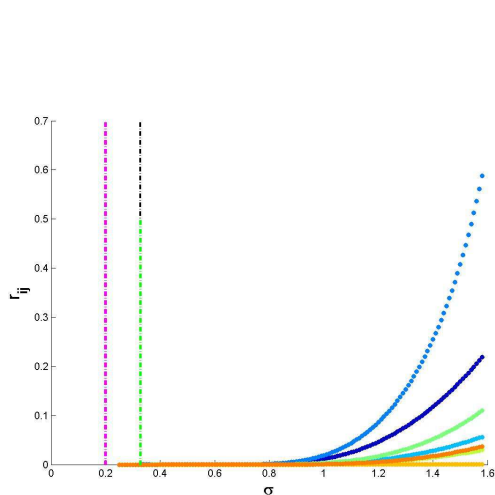
(c) Cercles



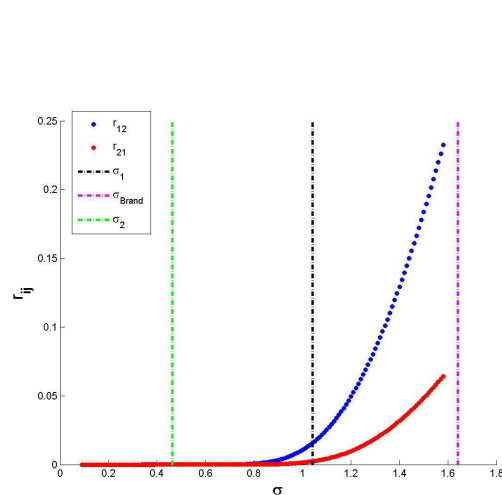
(d) Portions de couronnes



(f) Carrés



(g) Cible



(i) Rectangles étirés

FIGURE 1.10 – Ratio de normes de Frobenius fonction du paramètre σ

Remarque 1.9. D'après les figures 1.9 et 1.10, il existe un majorant dans les valeurs de σ mais il existe aussi un minorant. Le conditionnement de la matrice affinité (1.1), fonction d'une norme matricielle de A et de son inverse, impose une valeur minimale strictement positive pour éviter une matrice nulle.

1.3.2 Exemples géométriques de dimension $p \geq 2$

Afin de valider l'approche géométrique détaillée précédemment, on considère deux exemples de dimension p : le premier avec six blocs identiques de dimension p faiblement séparés entre eux et le deuxième avec des morceaux de p -sphère de $\mathbb{R}^p$ (cf figure 1.11). Les résultats présentés dans la suite sont obtenus en prenant consécutivement des valeurs de σ entre 0.01 et 1.5 et en calculant les deux mesures introduites précédemment (section 1.4). Des intervalles de faisabilité pour les valeurs de σ sont définis. Le but est de vérifier si les heuristiques (1.2) et (1.3) appartiennent aux intervalles appropriés ou non.

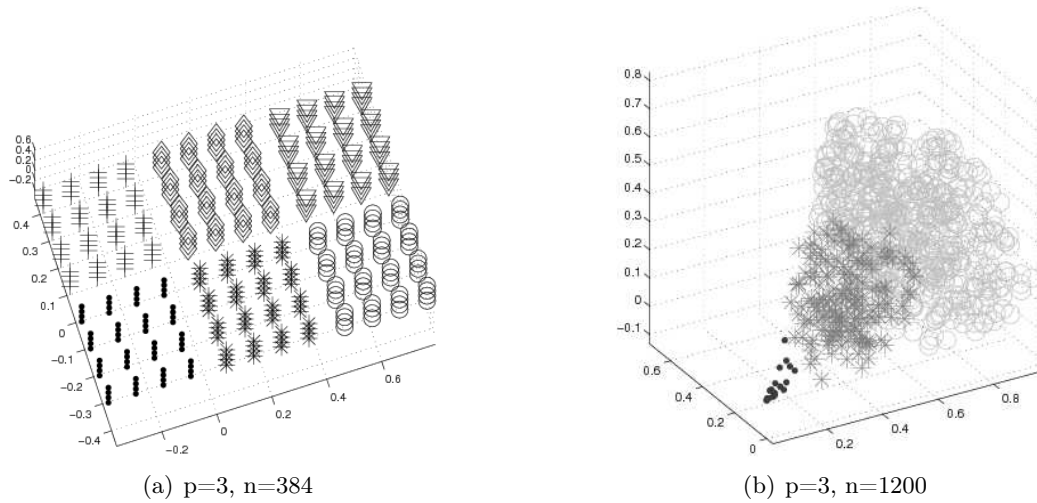


FIGURE 1.11 – Exemple 1 & 2 : 6 blocs et 3 portions de N -spheres

Premier exemple : 6 blocs

Cette géométrie, voir figure 1.11 (a), est constituée de 6 blocs de dimension p avec une distribution uniforme pour chacun. Cet exemple est en parfait accord avec l'hypothèse d'isotropie évoquée dans la section 3. Chaque bloc est composé de N^p points avec un pas de 0.1 dans chaque direction, et avec $N = 4$ dans le cas de $p = \{2, 3\}$ et $N = 3$ dans le cas $p = 4$. Les blocs sont séparés entre eux d'une distance de 0.13. Cet exemple correspond au cadre de l'heuristique (1.3) pour σ . Sur le tableau 1.1, sont indiqués les résultats pour les trois dimensions p du problème et aussi les valeurs σ_1 et σ_2 correspondant respectivement aux heuristiques (1.2) et (1.3). Pour déterminer les intervalles dans le cas des ratios de normes de Frobenius entre les blocs, la mesure de qualité est considérée comme acceptable lorsque le minimum des r_{ij} défini par (1.4), pour $i \in \{1, \dots, k\}$ et $i \neq j$, est inférieur ou égal à 0.15. En raison de la géométrie de l'exemple où l'on représente des clusters séparables linéairement, l'intervalle du pourcentage d'erreur pour lequel il n'existe pas d'erreur de clustering a un diamètre nettement plus grand que celui pour lequel le ratio de Frobenius est inférieur à 0.15.

Cet exemple montre, à nouveau, l'impact sur l'affinité intra-cluster et inter-clusters que peut avoir le paramètre σ . Les résultats du tableau 1.1 (a) montrent aussi que les deux heuristiques σ_1 et σ_2 appartiennent aux intervalles appropriés dans tous les cas.

Deuxième exemple : 3 portions d'hyperpshères

Cet exemple, représenté sur la figure 1.11 (b), est construit dans le même esprit que le précédent. De même, le tableau 1.1 (b) regroupe les résultats du clustering spectral via les deux mesures de qualité en fonction de p . En comparaison avec l'exemple des 6 blocs, chaque cluster présente un volume différent et la structure sphérique met à défaut des techniques comme le k -means car les clusters ne sont pas séparables linéairement. Ainsi, les intervalles où $P_{erreur} = 0\%$ sont approximativement les mêmes que ceux où le ratio de normes de Frobenius est inférieur à 0.15. Concernant les valeurs des heuristiques σ_1 et σ_2 , on obtient les mêmes conclusions que pour l'exemple avec 6 blocs.

1.4 Méthodes de classification spectrale modifiées

Dans cette section, de nouvelles définitions d'affinité, dérivées de l'affinité gaussienne (1.1), sont considérées. Tout d'abord, nous proposons de considérer des volumes au lieu de distances euclidiennes. Ensuite, dans un cadre de segmentation d'images où les informations géométriques et de luminance peuvent être considérées séparément, nous proposons d'englober ces deux notions dans un volume.

1.4.1 Cas de données volumiques

Dans le même esprit que le point de vue géométrique développé pour l'heuristique (1.2), nous avons considéré des ratios de volumes à la place de ratio de normes définis par la relation (1.1). Ainsi, nous incorporons une puissance dans l'exponentielle et la matrice affinité est définie par :

$$A_{ij} = \exp \left(- \left(\frac{\|x_i - x_j\|_2}{\sigma/2} \right)^d \right), \quad (1.5)$$

où d est pris respectivement entre 1 et 5 sur ces exemples afin de vérifier expérimentalement l'adéquation de la puissance d en fonction de la dimension du problème p comme suggéré précédemment (cf tableau 1.2). Pour chaque dimension, nous avons donc fait varier la puissance d dans le calcul de la matrice affinité de (1.5) et les intervalles de faisabilité sont calculés pour chaque cas pour les valeurs de σ par rapport aux deux mesures de qualité basées sur le pourcentage d'erreur de clustering et les ratios de norme de Frobenius. Le diamètre des intervalles, particulièrement pour la première mesure de qualité est plus grand pour une valeur de d proche de p . Ceci a tendance à abonder dans le fait de considérer des ratios de volumes plutôt que des distances au carré. Mais cette définition connaît une limite pour de grandes dimensions : la puissance p implique que les points doivent balayer chaque dimension de façon assez uniforme pour obtenir une matrice d'affinité nulle.

1.4.2 Traitement d'images

On considère maintenant un exemple de segmentation d'images. Dans ce cas, la matrice affinité est définie avec deux approches différentes :

TABLE 1.1 – Exemples géométriques

(a) 6 blocs p-dimensionnel				
p		2	3	4
σ_1		0.0349	0.0482	0.0558
σ_2		0.0344	0.0377	0.0391
Ratio de normes de Frobenius < 0.15		[0.02 ; 0.56]	[0.02 ; 0.58]	[0.02 ; 0.22]
P_{erreur}		[0.06 ; 1.24]	[0.06 ; 2.2]	[0.04 ; 0.78]

(b) 3 portions sphères de dimension p				
p		2	3	4
σ_1		0.0072	0.0268	0.0426
σ_2		0.0069	0.0261	0.0414
Ratio de normes de Frobenius < 0.15		[0.04 ; 0.18]	[0.02 ; 0.3]	[0.02 ; 0.04]
P_{erreur}		[0.02 ; 0.14]	[0.04 ; 0.16]	[0.04 ; 0.1]

TABLE 1.2 – Exemples géométriques avec l'affinité (1.5)

(a) 6 blocs p-dimensionnel pour $p = 2$, $p = 3$ et $p = 4$				
$p = 2$, d	1	3	4	5
Ratio de normes de Frobenius < 0.15	[0.02;0.3]	[0.02;0.6]	[0.02;0.62]	[0.02;0.6]
P_{erreur}	[0.12;1.6]	[0.08;1.06]	[0.1;1.18]	[0.12;1.1]
$p = 3$, d	1	3	4	5
Ratio de normes de Frobenius < 0.15	[0.02;0.3]	[0.02;0.64]	[0.02;0.66]	[0.02;0.66]
P_{erreur}	[0.02;3]	[0.04;1.4]	[0.04;1.2]	[0.04;1.2]
$p = 4$, d	1	3	4	5
Ratio de normes de Frobenius < 0.15	[0.02;0.3]	[0.02;0.26]	[0.02;0.28]	[0.02;0.28]
P_{erreur}	[0.02;1.2]	[0.02;0.62]	[0.12;0.52]	[0.14;0.48]
(d) 3 portions d'hyperphères pour $p = 2$, $p = 3$ et $p = 4$				
$p = 2$, d	1	3	4	5
Ratio de normes de Frobenius < 0.15	[0.04;0.14]	[0.6;0.2]	[0.06;0.2]	[0.06;0.18]
P_{erreur}	[0.04;0.08]	[0.06;0.18]	[0.6;0.18]	[0.06;0.18]
$p = 3$, d	1	3	4	5
Ratio de normes de Frobenius < 0.15	[0.02;0.12]	[0.04;0.5]	[0.06;0.5]	[0.08;0.5]
P_{erreur}	[0.02;0.12]	[0.04;0.16]	[0.08;0.18]	[0.08;0.18]
$p = 4$, d	1	3	4	5
Ratio de normes de Frobenius < 0.15	[0.04;0.04]	[0.04;0.08]	[0.06;0.1]	[0.1;0.16]
P_{erreur}	[0.02;0.02]	[0.06;0.16]	[0.08;0.17]	[0.1;0.2]

une boîte 3D rectangulaire

Si l'image donnée peut être considérée comme suffisamment isotropique c'est-à-dire si les pas entre chaque pixel et chaque niveau de luminance varie avec la même amplitude, on peut inclure les données de l'image dans une boîte 3D rectangulaire et incorporer l'heuristique (1.3) pour σ dans la matrice affinité donnée par (1.1).

un produit d'une similarité spatiale et une similarité de luminescence

La seconde possibilité est de considérer que l'image de départ est composée de deux ensembles distincts de variables, chacune avec une amplitude et une densité spécifique. Alors la distribution spatiale des pixels est isotropique mais la luminance est divisée en niveaux (256 maximum) et la densité de luminance ne peut pas provenir du nombre de points. Dans les articles traitant de segmentation d'images (par exemple [85, 91]), le produit d'une matrice affinité pour les données spatiales et d'une matrice d'affinité pour les valeurs de luminance est habituellement considéré. Chacune de ces deux matrices a un paramètre σ spécifique. La matrice affinité est construite de la manière suivante :

$$A_{ij} = \exp \left(-\frac{\|x_i - x_j\|^2}{(\sigma_G/2)^2} - \frac{|I(i) - I(j)|}{(\sigma_B/2)} \right), \quad (1.6)$$

où $I(i)$ est la valeur de luminance dans $\mathbb{R}$ et x_i les coordonnées du pixel i dans $\mathbb{R}^2$. Le paramètre σ_G est donné par l'heuristique (1.2) appliquée uniquement sur les données spatiales, et σ_B est fixé à $(I_{\max}/\ell)$ avec ℓ un nombre caractéristique des niveaux de luminance. Ce dernier paramètre correspond à l'heuristique 1.2 dans le cas de données 1D. Par exemple, dans l'exemple 1.12, ℓ est égal au nombre de seuils de l'image et σ_B définira le diamètre des intervalles dans lesquels les valeurs de luminance devraient être groupées.

Cette définition reste dans l'esprit des développements de la section 1.3 car σ défini par (1.2) reflète la distance référence de clustering dans le cas d'isotropie locale et divise suffisamment la distribution des points. Avec 256 niveaux de luminance maximum, la distribution ne peut pas être considérée comme localement dispersée (beaucoup de valeurs sont égales entre elles) et on doit donner une distance caractéristique à partir de laquelle les valeurs de luminance peuvent être clusterisées. On note aussi que la solution de prendre $\sigma_B = I_{\max}/256$ revient à re-échelonner les valeurs de luminance en clusters, équivalent à une décomposition très fine du grain de l'image, alors que la segmentation d'image doit requérir une analyse de tous les clusters constitués de pixels proches entre eux et avec le même niveau de luminance.

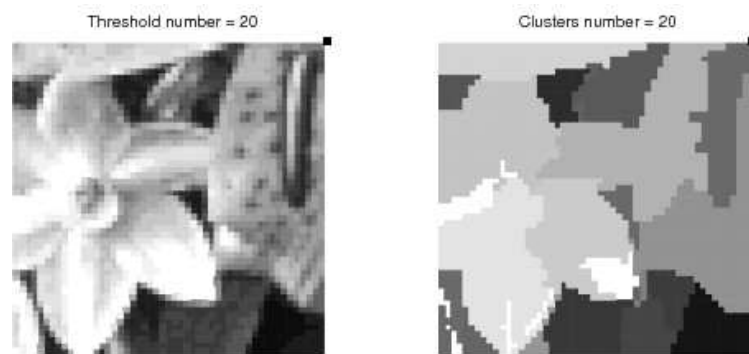
Dans les résultats suivants, les approches (1.1) et (1.6) sont testées pour le calcul de la matrice affinité sur deux images de 50×50 pixels soit $N = 2500$ points. Ces images présentent des caractéristiques intéressantes : l'image de gauche représente plusieurs fleurs aux niveaux de gris nettement différents tandis que l'image de droite représente une seule fleur, une rose, dont les niveaux de gris restent très proches et les nuances liées au relief des pétales sont difficiles à cerner.

Les figures 1.12 et 1.13 représentent sur la gauche l'image seuillée et sur la droite, les résultats du clustering obtenus avec :

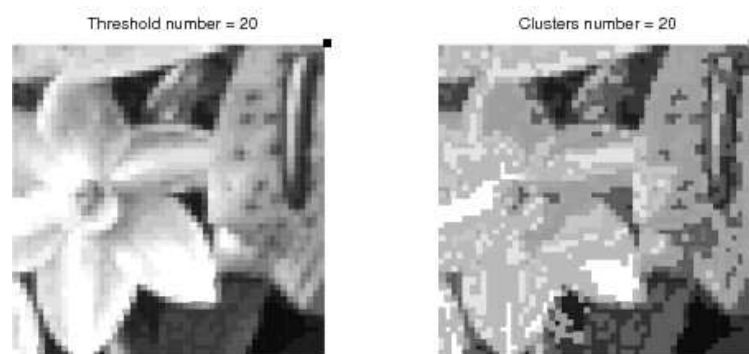
- le produit d'une similarité spatiale et une similarité de luminescence (1.6) pour les figures 1.12 et 1.13 (a) et (c),
- la boîte 3D affinité (1.1) pour les figures 1.12 et 1.13 (b) et (d).

Dans les deux cas, les résultats sont visuellement acceptables. L'approche 3D semble donner de meilleurs résultats que celle avec le produit 2D par 1D mais cela nécessite de plus amples études pour s'assurer qu'une approche est meilleure en général et affiner les résultats. Cependant, le calcul du spectre de la matrice affinité pleine limite cette étude et nous restreint à considérer de faibles dimensions pour l'image.

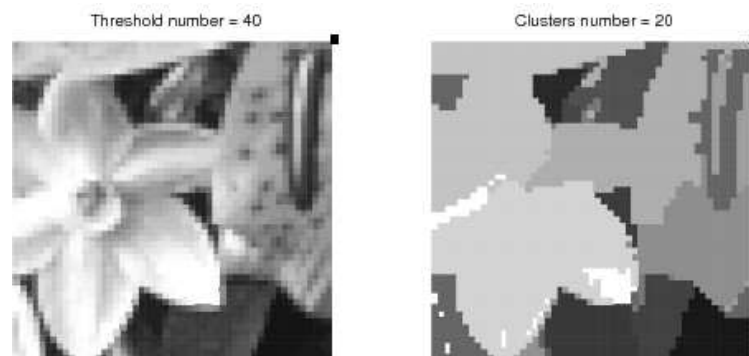
De plus, à ce stade, il manque un aspect théorique, un critère topologique, permettant de valider une heuristique pour le paramètre de l'affinité gaussienne ou, dans le cas de segmentation d'image, une définition de l'affinité gaussienne. Il reste à étudier le fonctionnement de la méthode de clustering spectral et montrer comment le paramètre influe sur la méthode.



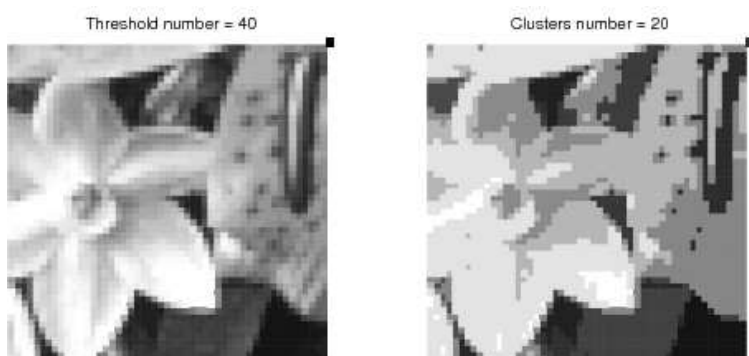
(a) Boîte produit 2D par 1D d'affinité : 20 seuils de luminances et 20 clusters



(b) Boîte rectangulaire 3D d'affinité : 20 seuils de luminances et 20 clusters

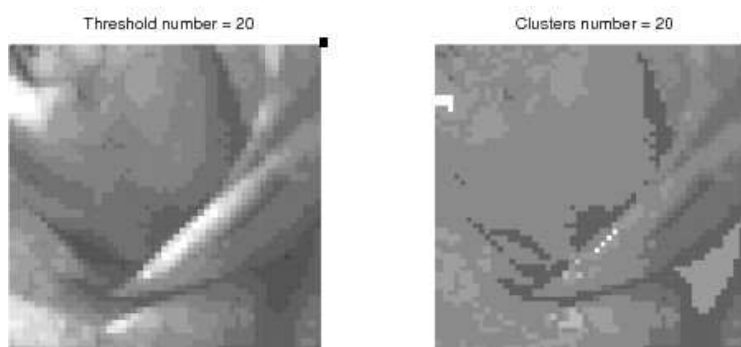


(c) Boîte produit 2D par 1D d'affinité : 40 seuils de luminances et 20 clusters

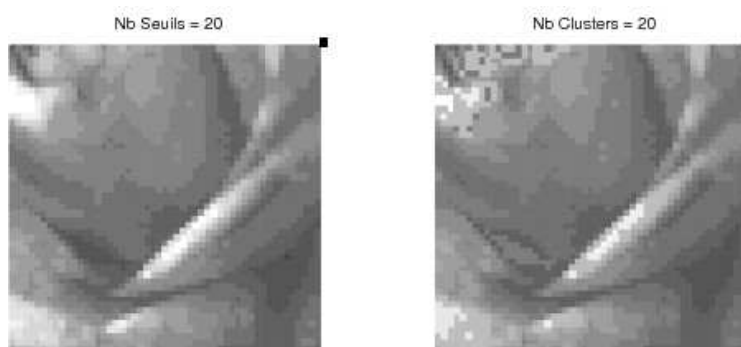


(d) Boîte rectangulaire 3D d'affinité : 40 seuils de luminances et 20 clusters

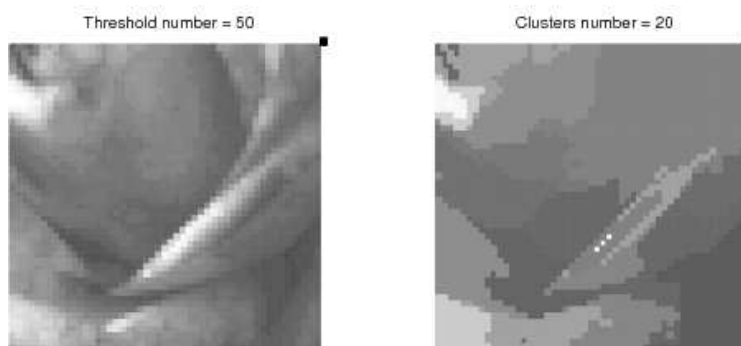
FIGURE 1.12 – Test sur images avec une boîte rectangulaire 3D d'affinité



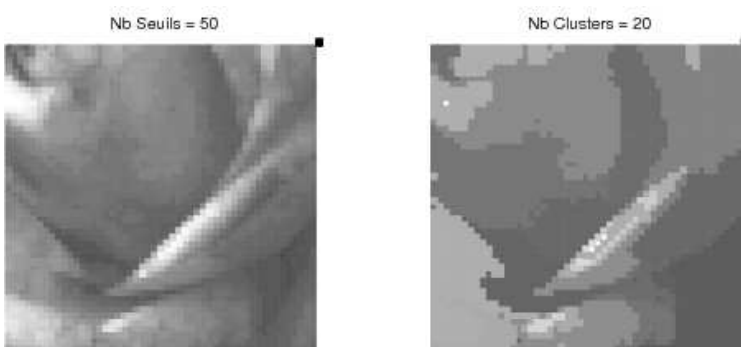
(a) Boîte produit 2D par 1D d'affinité : 20 seuils de luminances et 20 clusters



(b) Boîte rectangulaire 3D d'affinité : 20 seuils de luminances et 20 clusters



(c) Boîte produit 2D par 1D d'affinité : 50 seuils de luminances et 20 clusters



(d) Boîte rectangulaire 3D d'affinité : 50 seuils de luminances et 20 clusters

FIGURE 1.13 – Test sur images avec une boîte produit 2D par 1D d'affinité

Chapitre 2

Classification et éléments spectraux de la matrice affinité gaussienne

Dans ce chapitre, nous nous intéressons au fonctionnement de la méthode de spectral clustering pour un ensemble fini de points afin de pouvoir juger la qualité du clustering et le choix du paramètre σ . Pour cela, nous proposons d'interpréter l'algorithme de *Ng, Jordan et Weiss* [84] sans l'étape de normalisation et plus précisément l'espace de projection spectral via les équations aux dérivées partielles et la théorie des éléments finis. Cette interprétation apportera une explication sur la construction des clusters c'est-à-dire un sens à l'extraction des plus grands vecteurs propres d'une matrice affinité gaussienne.

2.1 Diverses interprétations

De nombreuses interprétations basées sur des approches probabiliste, d'optimisation ou bien de théorie des graphes et d'opérateur dans les espaces fonctionnels du spectral clustering ont été recensées [105, 32]. Rappelons que cette méthode repose seulement sur l'ensemble des données et la similarité, notée s , entre les couples de points, sans a priori sur la forme des clusters.

Partitionnement de graphe

A partir de la seule mesure de similarité entre les points, une représentation possible de ces points consiste à les identifier par un graphe de similarité, noté $G = (V, E)$. Chaque sommet $v_i \in V$ du graphe G est un point x_i . Deux sommets sont connectés si la similarité s_{ij} entre les points x_i et x_j est positive ou supérieure à un seuil prédéfini et l'arête du graphe est pondérée par la valeur de s_{ij} . A partir de cette représentation, plusieurs fonctions objectifs basées sur la minimisation de la similarité totale entre deux clusters peuvent être minimisées : le ratio de coupe [51], la coupe normalisée [91] ou bien la coupe min-max [33] incluant aussi la maximisation de la similarité totale intra-cluster. Cette interprétation est en accord avec le principe même de la classification non supervisée. Cependant, l'utilisation des premiers vecteurs propres d'un Laplacien de graphe comme espace de projection dans lequel les clusters sont constitués est, en général, justifiée avec des arguments heuristiques ou comme une relaxe du problème discret de clustering.

Marche aléatoire sur un graphe

Meila et Shi [79] utilisent le lien entre le Laplacien d'un graphe et les chaînes de Markov initié par [25] et identifient la matrice d'affinité normalisée comme une matrice stochastique représentant une marche aléatoire sur un graphe et le critère de la coupe normalisée comme la somme des probabilités de transition entre deux ensembles. Mais, seuls le cas où les vecteurs propres sont constants par morceaux pour des structures matricielles spécifiques (bloc diagonales) sont considérées. D'autres aspects des marches aléatoires sont utilisés pour proposer des variantes de la méthode de spectral clustering avec des techniques agglomératives [55] ou bien l'utilisation d'une distance euclidienne basée sur le temps moyen de commutation entre les points d'une marche aléatoire d'un graphe [114].

Perturbation matricielle

Comme évoqué dans le chapitre 1, *Ng, Jordan et Weiss* [84] expliquent le clustering spectral en considérant un cas idéal où la matrice affinité gaussienne a une structure numérique bloc diagonale. Cependant, dans le cas général, cette structure n'est pas conservée donc les auteurs utilisent des résultats sur la perturbation de matrices. La théorie de la perturbation matricielle [96] traite du comportement des valeurs propres et des vecteurs propres d'une matrice B lorsque celle-ci est sujette à de faibles perturbations additives H c'est-à-dire l'étude des éléments spectraux de $\tilde{B} = B + H$. Le théorème de Davis-Kahan [18] permet de borner la différence, via les angles principaux [49], entre les espaces propres de B et $\tilde{B}$ associés aux valeurs propres proches de 1. Cette différence dépend de l'écart entre les valeurs propres proches de 1 et le reste du spectre. Or, ces résultats sont sensibles à l'importance de la perturbation et l'écart peut être très petit.

Interprétation via des opérateurs

D'autres interprétations mathématiques de cette méthode ont été étudiées en utilisant une version continue de ce problème. Plusieurs travaux ont été menés pour expliquer le fonctionnement du clustering spectral. *Belkin et Niyogi* [11] ont montré que sur une variété de $\mathbb{R}^p$, les premiers vecteurs propres sont des approximations de l'opérateur de Laplace-Beltrami. Mais cette justification est valide lorsque les données sont uniformément échantillonnées sur une variété de $\mathbb{R}^p$.

Nadler et al [82] donnent une autre interprétation probabiliste basée sur un modèle de diffusion. Pour cela, la distance de diffusion est définie comme une distance entre deux points basée sur une marche aléatoire sur un graphe. La projection de diffusion de l'espace des données dans un espace est définie par les k premiers vecteurs propres. Il a été démontré que les distances de diffusion dans l'espace original sont égales aux distances euclidiennes dans l'espace de projection de diffusion. Ce résultat justifie l'utilisation des distances euclidiennes dans l'espace de projection pour de diffusion pour le clustering.

Tous ces résultats sont établis asymptotiquement pour un grand nombre de points. Cependant, d'un point de vue numérique, le spectral clustering partitionne correctement un ensemble fini de points avec des distributions quelconques sur les dimensions.

Nous proposons donc une nouvelle interprétation où l'ensemble fini des données représentera la discrétisation de sous-ensembles. Ainsi, les vecteurs propres de la matrice gaussienne seront, pour une bonne valeur de t , la représentation discrète de fonctions à support sur un seul de ces sous-ensembles. L'objectif est aussi d'avoir des éléments d'analyse pour juger la qualité du clustering et du choix du paramètre σ .

2.2 Présentation du résultat principal

On se propose de relier la classification (ou clustering) d'un ensemble fini de points à une partition d'ouverts de l'espace $\mathbb{R}^p$.

Définition 2.1 (Clustering).

1. On dit qu'un ouvert Ω , réunion finie de composantes connexes Ω_i , $i = \{1, \dots, k\}$ distinctes, induit un k -clustering sur un ensemble fini de points $\mathcal{P}$, si les intersections de $\mathcal{P}$ avec les Ω_i , constituent une partition $\mathcal{C} = \{C_1, \dots, C_k\}$ de $\mathcal{P}$, c'est-à-dire que $C_j = \mathcal{P} \cap \Omega_j \neq \emptyset$, $i = 1, \dots, k$, $C_i \cap C_j = \emptyset$ pour $i \neq j$ et $\mathcal{P} = \bigcup_{i=1}^k C_i$.
2. Soit $\mathcal{C} = \{C_1, \dots, C_k\}$ une partition donnée de l'ensemble fini de points de $\mathcal{P}$. On dit qu'un ouvert Ω , réunion finie de composantes connexes Ω_i , $i = \{1, \dots, k\}$ distinctes, induit un k -clustering sur $\mathcal{P}$ compatible avec la partition donnée de $\mathcal{P}$ si l'ensemble des C_j , $j = \{1, \dots, k\}$ est identique à l'ensemble des $\mathcal{P} \cap \Omega_i$, $i = 1, \dots, k$.

Si on considère l'exemple illustré dans la figure 2.1, dans lequel les quatre couronnes constituent la partition "naturelle" de l'ensemble des points de $\mathbb{R}^2$ donné, le cas (a) constitue un clustering compatible, les autres cas n'étant pas compatibles. Ce cas de clustering compatible sera aussi, dans la suite, appelé "clustering idéal".

Remarque 2.2. En d'autres termes, le clustering partitionne suivant les composantes connexes. Les exemples (b)-(d) peuvent être définis comme suit. Soit $\mathcal{C} = \{C_1, \dots, C_k\}$ un partitionnement de l'ensemble $\mathcal{P}$.

- le sous clustering (cas (b) de la figure 2.1) définit le cas où :
 $k' < k$ et $\forall i = \{1, \dots, k\}, \exists j \in \{1, \dots, k'\}, \mathcal{P}_i \subset C_j$.
- le sur-clustering (cas (c)) définit le cas où :
 $k' > k$ et $\forall j = \{1, \dots, k'\}, \exists i \in \{1, \dots, k\}, C_j \subset \mathcal{P}_i$;
- un mauvais clustering (cas (d)) représente des clusters dont les points appartiennent à plusieurs composantes connexes sans recouvrir de composante connexe entièrement.

Considérons donc maintenant un ensemble fini de points $\mathcal{P} = \bigcup_{i=1}^k \mathcal{P}_i$ et une partition en k classes disjointes 2 à 2 induisant un clustering compatible avec la partition donnée de $\mathcal{P}$.

Proposition 2.3. Supposons qu'il existe k vecteurs propres notés $X_1, \dots, X_k$ de la matrice $A \in \mathcal{M}_{N,N}(\mathbb{R})$ définie par (2.1) tels que : pour tout $l \in \{1, \dots, k\}$ et pour tout $i \in \{1, \dots, N\}$

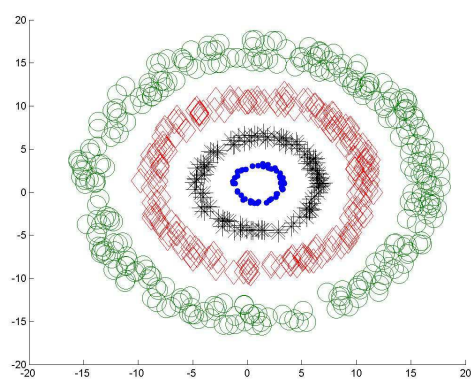
$$(X_l)_i \begin{cases} = 0 & \text{si } x_i \notin \mathcal{P}_l, \\ \neq 0 & \text{si } x_i \in \mathcal{P}_l. \end{cases}$$

Alors la partition $\mathcal{C} = \{C_1, \dots, C_k\}$ issue de l'algorithme 2 du spectral clustering définit un clustering idéal.

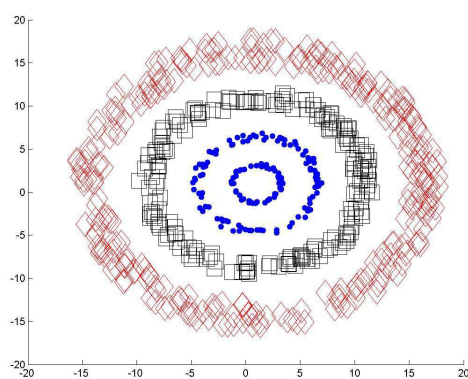
Démonstration. Supposons que les vecteurs propres $X_1, \dots, X_k$ de la matrice affinité A vérifient les hypothèses du théorème. Supposons que la partition $\mathcal{C}$ ne définisse pas un clustering idéal c'est-à-dire :

$$\exists j \in \{1, \dots, k\}, \forall i \in \{1, \dots, k\}, C_j \neq \mathcal{P}_i \iff \exists x_i \in \mathcal{P}_i, x_j \in \mathcal{P}_j \text{ avec } \mathcal{P}_i \neq \mathcal{P}_j \text{ tels que } \begin{cases} x_i \in \mathcal{C}_m, \\ x_j \in \mathcal{C}_m. \end{cases}$$

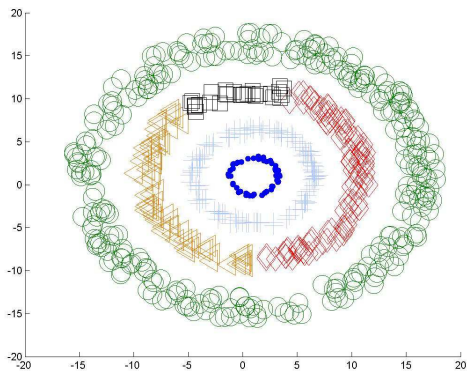
Montrons d'abord, cette équivalence. Supposons que Ω n'induit pas un k -clustering compatible. Supposons qu'il n'existe pas deux points $x_i \in \mathcal{P}_i$, $x_j \in \mathcal{P}_j$ appartenant au même cluster $\mathcal{C}_m$. Alors



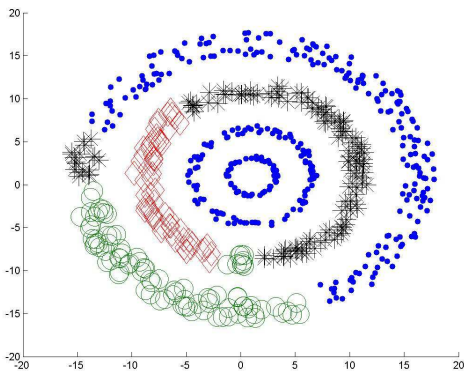
(a) Exemple de clustering idéal



(b) Exemple de sous-clustering



(c) Exemple de sur-clustering



(d) Exemple de mauvais clustering

FIGURE 2.1 – Définitions du clustering

pour tout $j \in \{1, \dots, k\}$, $\mathcal{C}_j \subset \mathcal{P}_j$. Or il existe $j \in \{1, \dots, k\}$ tel que $\mathcal{C}_j \neq \mathcal{P}_j$. Donc $\bigcup_{j=1}^k \mathcal{C}_j \neq \bigcup_{j=1}^k \mathcal{P}_j = \mathcal{P}$. Ce qui contredit l'hypothèse que $\mathcal{C}$ est une partition. La réciproque est triviale. Supposons maintenant que Ω n'induit pas un k -clustering compatible, c'est-à-dire qu'il existe deux points $x_i \in \mathcal{P}_i$ et $x_j \in \mathcal{P}_j$ avec $i \neq j$ tels que $x_i \in \mathcal{C}_1$ et $x_j \in \mathcal{C}_1$. S'ils sont assignés au même cluster $\mathcal{C}_1$ alors, d'après l'algorithme 2, $Y_{i1} \neq 0$ et $Y_{j1} \neq 0$. En d'autres termes, $(X_1)_i \neq 0$ et $(X_1)_j \neq 0$. Alors, d'après les hypothèses sur les vecteurs propres, $x_i \in \mathcal{P}_1$ et $x_j \in \mathcal{P}_1$ ce qui est faux. Donc la partition $\mathcal{C} = \{\mathcal{C}_1, \dots, \mathcal{C}_k\}$ est identique au k -clustering induit par Ω . $\square$

La proposition 2.2 énonce un résultat de clustering *immédiat* en ce sens qu'il est trivial à réaliser, sous réserve qu'on puisse trouver exactement k vecteurs dont les coordonnées sont non nulles sur une seule des k partitions de points $\mathcal{P}_j$, $j = 1, \dots, k$. Il est clair que dans la pratique ce ne sera pas le cas, mais nous allons analyser dans ce chapitre sous quelles hypothèses il est possible de se rapprocher de cette situation idéale. Avant toute chose, il est utile de rappeler que l'existence de tels vecteurs rejoint directement l'hypothèse de structure diagonale par bloc de la matrice A , telle que l'exploitent Ng, Jordan et Weiss [84]. En effet, sous l'hypothèse d'une telle structure diagonale par bloc, les vecteurs propres de A peuvent se regrouper en k sous ensembles de vecteurs ayant chacun des composantes non nulles en correspondance avec l'un des k blocs diagonaux de la matrice. La normalisation de la matrice ne sert alors qu'à éviter d'avoir à faire une décomposition spectrale complète de la matrice A et à retrouver (étape qui peut être coûteuse) dans l'ensemble des vecteurs propres la répartition par bloc des composantes non nulles de ces vecteurs (après permutation éventuelle des lignes). En effet, la normalisation garantit simplement que la valeur propre dominante égale à 1 est de multiplicité k , et que les vecteurs propres associés sont une combinaison linéaire de k vecteurs ayant des coordonnées non nulles et constantes relativement à chacun des k blocs diagonaux respectivement.

L'une des questions à laquelle nous nous intéressons dans ce chapitre est d'analyser dans quelle mesure la matrice de similarité A est proche de cette situation bloc-diagonale idéale. Ng, Jordan et Weiss [84] abordent cette question en analysant la structure des vecteurs propres de A par le biais de la théorie de la perturbation matricielle. Dans le même esprit, nous analyserons la structure de ces vecteurs propres à l'aide d'un problème continu mettant en jeu l'équation de la chaleur.

Dans le cas où l'étape de normalisation est supprimée, la méthode de spectral clustering se résume aux étapes de l'algorithme 2, dans lequel intervient la décomposition spectrale de la matrice d'affinité Gaussienne, explicitée en (2.1). Comme les éléments spectraux de la matrice d'affinité ne fournissent pas explicitement de critère géométrique relativement à un ensemble discret de données, nous nous proposons de revenir à une formulation continue où les clusters sont inclus dans un ouvert Ω fournissant un k -clustering compatible. En interprétant la matrice d'affinité gaussienne comme la discrétisation du noyau de Green de l'équation de la chaleur et en utilisant les éléments finis, on montre que, pour un ensemble fini de points, les vecteurs propres de la matrice d'affinité gaussienne sont la représentation asymptotique de fonctions dont le support est inclus dans une seule composante connexe. Ce retour à une formulation continue est effectué à l'aide des éléments finis. Ainsi, les vecteurs propres de la matrice d'affinité A sont interprétés comme la discrétisation de fonctions propres d'un opérateur. En effet, avec les éléments finis dont les noeuds correspondent aux données d'origine, une représentation d'une fonction est donnée par sa valeur nodale. Donc on peut interpréter la matrice A et ses vecteurs propres comme les représentations respectives d'un opérateur L^2 et d'une fonction L^2 .

L'opérateur dont la représentation en éléments finis concorde avec la définition de A est le noyau de l'équation de la chaleur, noté K_H , sur $\mathbb{R}^p$. Comme le spectre de l'opérateur S_H (convolution par K_H) est essentiel, les vecteurs propres de A ne peuvent pas être directement interprétés comme

Algorithm 2 Classification spectrale non normalisée

Input : Ensemble des données S , Nombre de clusters k'

1. Construction de la matrice affinité $A \in \mathbb{R}^{N \times N}$ définie par :

$$A_{ij} = \begin{cases} \exp(-\|x_i - x_j\|^2 / 2\sigma^2) & \text{si } i \neq j, \\ 0 & \text{sinon} \end{cases} \quad (2.1)$$

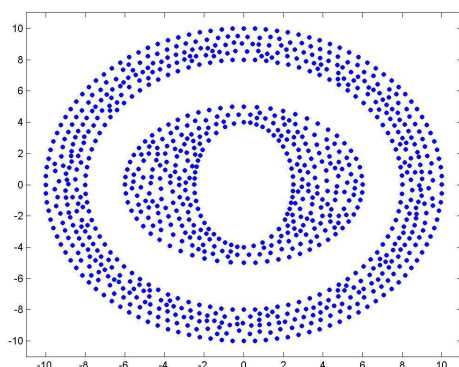
2. Construction de la matrice $X = [X_1 X_2 \dots X_{k'}] \in \mathbb{R}^{N \times k'}$ formée à partir de k' vecteurs propres x_i , $i = \{1, \dots, k'\}$ de A .
3. Construction de la matrice Y formée en normalisant les lignes de X :

$$Y_{ij} = \frac{X_{ij}}{\left(\sum_j X_{ij}^2\right)^{1/2}}$$

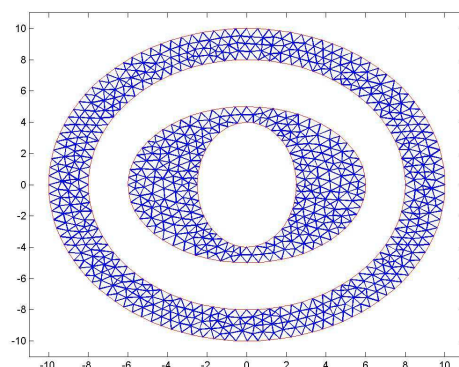
4. Traiter chaque ligne de Y comme un point de $\mathbb{R}^k$ et les classer en k clusters via la méthode k -means
 5. Assigner le point original x_i au cluster $\mathcal{C}_j$ si et seulement si la ligne i de la matrice Y est assignée au cluster $\mathcal{C}_j$.
-

une représentation des fonctions propres de S_H . Cependant, sur un domaine borné $\mathcal{O}$, K_H est proche de K_D , noyau de l'équation de la chaleur sur un domaine borné Ω avec conditions de Dirichlet, pour Ω contenant strictement $\mathcal{O}$. Maintenant, l'opérateur S_D (convolution par K_D) admet des fonctions propres $(v_{n,i})$ dans $H_0^1(\Omega)$ (proposition 2.17). Ces $v_{n,i}$ ont la propriété d'avoir leurs supports inclus sur une seule des composantes connexes. Ensuite, on montre que l'opérateur S_H admet les $v_{n,i}$ comme fonctions propres plus un résidu, noté η (proposition 2.23). Enfin, en utilisant l'approximation par les éléments finis et une condensation de masse, on montre (proposition 2.36) que les vecteurs propres de A sont une représentation de ces fonctions propres plus un résidu (noté ψ) fonction du paramètre t .

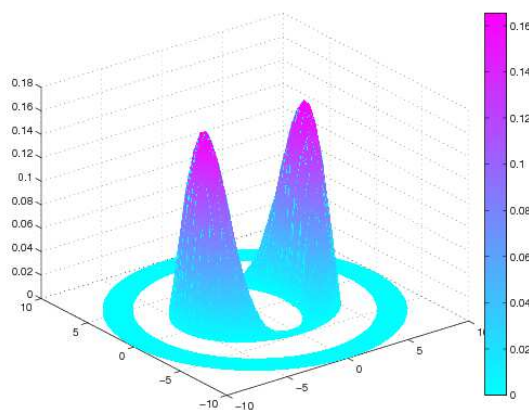
Pour illustrer les différentes étapes du résultat principal, on considère $n = 1400$ points représentant deux couronnes figure 2.2 (a). Le maillage de ces 2 composantes connexes est représenté sur la figure 2.2 (b). Les fonctions propres discrétisées de l'opérateur de la chaleur S_D à supports sur une seule couronne sont notées V_1 et V_2 et tracées sur les figures 2.2(c)-(d). Les figures (e)-(f) représentent le produit de la matrice d'affinité A et de la matrice de masse M résultant du maillage par les fonctions propres discrétisées V_1 et V_2 . Ce même produit sans la matrice de masse est représenté sur les figures 2.3 (a)-(b). Avec ou sans la matrice de masse, la propriété de support sur une seule composante connexe est préservée pour cet exemple. Enfin l'influence du paramètre de la chaleur t (variable temporelle dans l'équation de la chaleur) sur la différence entre les fonctions propres discrétisées V_1 et V_2 et les vecteurs propres de la matrice affinité AM et A est étudiée par le biais de la corrélation (figure 2.3 (d)-(f)) et de la différence en norme euclidienne (figure 2.3 (c)-(e)). Les fonctions propres discrétisées sont respectivement comparées aux vecteurs propres de AM sur les cas (c)-(d) et aux vecteurs propres de A sur les cas (e)-(f). Ces dernières fonctions ont un comportement semblable à celui du résidu $\psi(h, t)$ de la proposition 2.36. On observe donc que la propriété des vecteurs propres à support sur une seule composante connexe est préservée au mieux sur un intervalle particulier de valeurs de t .



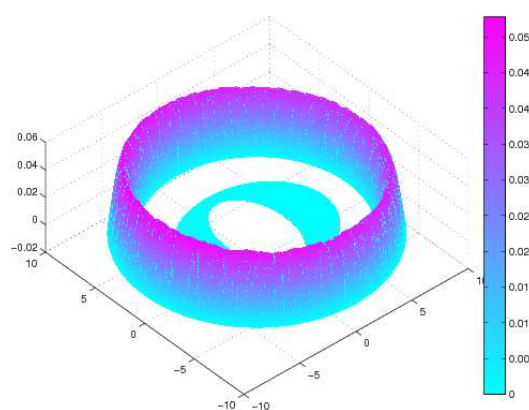
(a) Ensemble des données



(b) Maillage des donuts



(c) Fonction propre à valeur sur la première composante connexe



(d) Fonction propre à valeur sur la deuxième composante connexe

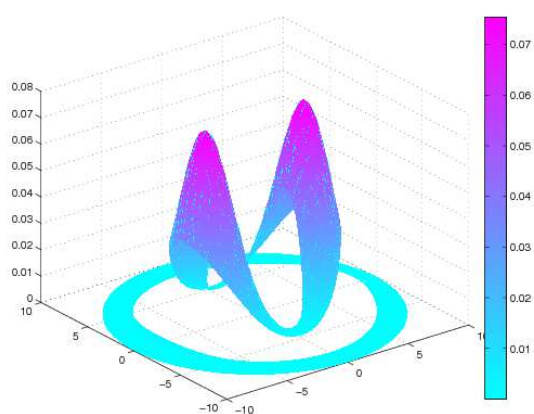
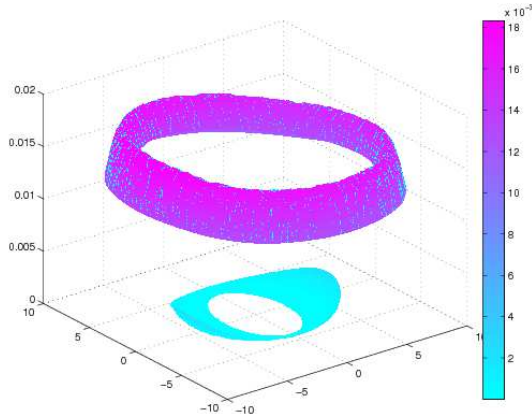
(e) Produit $AMV_{1,1}$ (f) Produit $AMV_{1,2}$

FIGURE 2.2 – Exemples des donuts : Données, maillage et fonctions propres

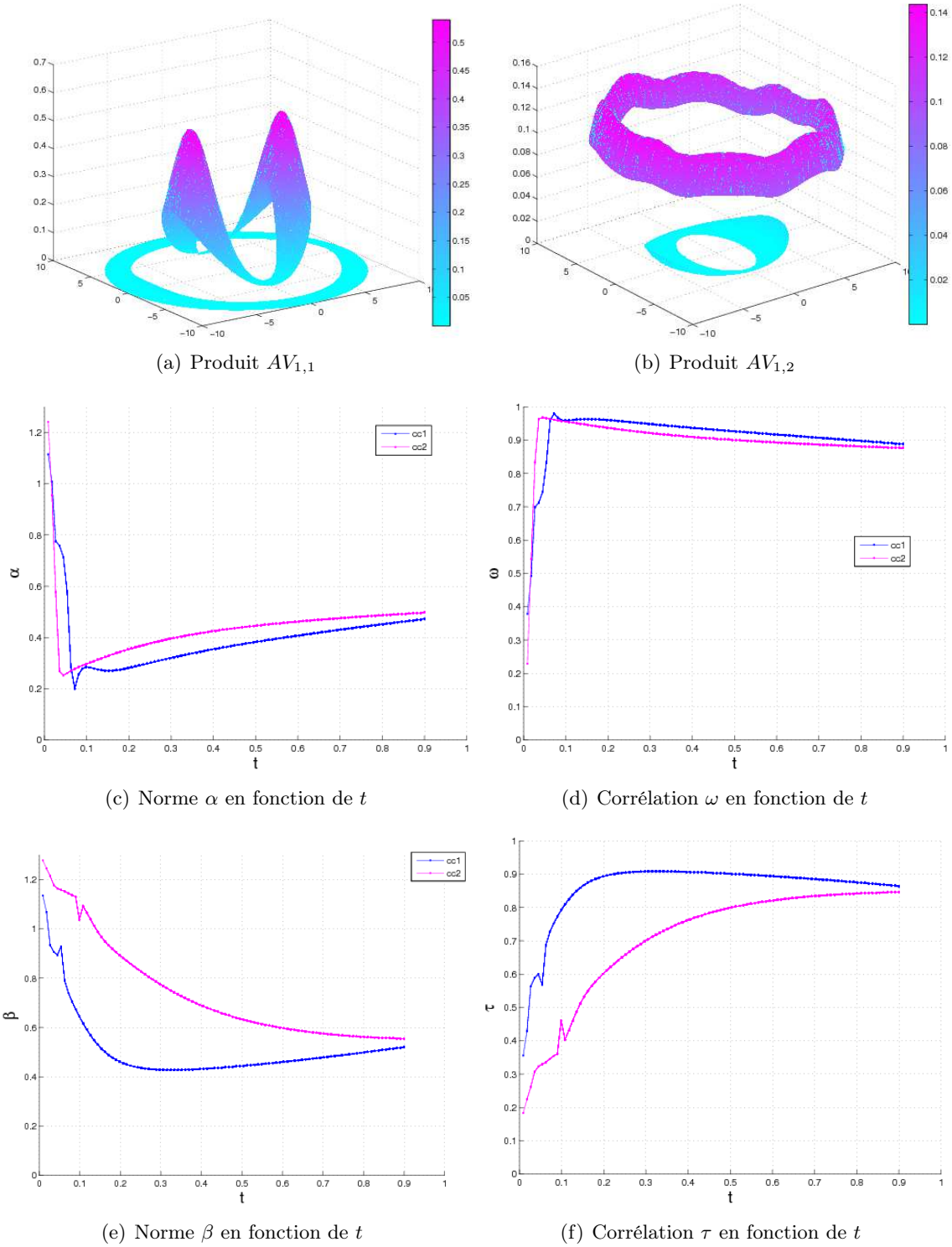


FIGURE 2.3 – Exemple des donuts : Corrélacion et différence en norme entre V_i et respectivement les vecteurs propres de $(A + \mathbb{I}_N)M$ et de A , pour $i \in \{1, 2\}$

Le résultat principal de cette analyse est de conforter la discrimination des composantes des vecteurs propres par classe de points, en fonction du temps de diffusion de la chaleur t , de la séparation moyenne h entre les points voisins au sein d'une même classe dans un ouvert Ω_i , et d'un paramètre γ minorant la distance de séparation entre $\overline{\Omega_i}$ et $\overline{\Omega_j}$, $1 \leq i \neq j \leq k$. Le temps t de propagation de la chaleur, qui peut être relié directement au paramètre σ dans la matrice d'affinité, joue en effet un rôle essentiel dans la séparation de l'information par composante connexe Ω_i , $1 \leq i \leq k$. Sachant que pour des instants t très petits, la propagation de la chaleur reste essentiellement localisée autour de chaque point possédant une unité de chaleur à l'instant $t = 0$, et que quand t devient très grand, la température a tendance à s'uniformiser globalement, il est légitime de penser qu'il existe un intervalle de temps consistant avec une situation de clustering donnée.

2.3 Propriétés de classification dans le continu

Dans la suite, quelques résultats d'Analyse Fonctionnelle seront nécessaires pour définir une propriété de clustering sur les fonctions propres de l'opérateur de Laplace. Ensuite, cette propriété permettra d'établir un résultat de clustering pour l'opérateur de la chaleur avec conditions de Dirichlet. Enfin, via un lien sur les noyaux de Green respectifs, une propriété de classification (ou propriété de clustering) pour l'opérateur de la chaleur en espace libre pourra être énoncée.

2.3.1 Quelques résultats d'Analyse Fonctionnelle

L'interprétation via l'équation de la chaleur requiert de se placer dans un contexte d'équations aux dérivées partielles. Dans la suite, nous introduirons donc quelques résultats d'analyse fonctionnelle [22] et d'équations aux dérivées partielles [60] concernant les opérateurs de Laplace et les équations de la chaleur avec conditions de Dirichlet.

Espaces de Sobolev, injections et prolongement

Soit Ω un ouvert connexe de $\mathbb{R}^p$. Dans la suite, on notera par $H^m(\Omega)$ l'espace de Sobolev défini par :

$$H^m(\Omega) = \{f \in L^2(\Omega) | \forall \alpha \in \mathbb{N}^p \text{ avec } |\alpha| \leq m, D^\alpha f \in L^2(\Omega)\}.$$

Dans $H^m(\Omega)$, le produit scalaire est défini par :

$$(u|v)_{H^m} = \sum_{|\alpha| \leq m} (D^\alpha u | D^\alpha v)_{L^2(\Omega)}.$$

On définit les espaces de Sobolev dont les éléments satisfont aussi les conditions de Dirichlet :

$$H_0^m(\Omega) = \{f \in L^2(\Omega) | f|_{\partial\Omega} = 0 \text{ et } \forall \alpha \in \mathbb{N}^p \text{ avec } |\alpha| \leq m, D^\alpha f \in L^2(\Omega)\}.$$

Suivent les injections dans les espaces de Sobolev [26] :

Théorème 2.4 (Injections de Sobolev). *L'inclusion suivante pour un ouvert Ω de $\mathbb{R}^p$ est valable :*

$$\text{si } s > p/2, H^s(\Omega) \hookrightarrow C^0(\bar{\Omega}).$$

Enfin, pour pouvoir faire le lien entre l'espace Ω et $\mathbb{R}^p$ dans la suite, on a recours au théorème suivant [60] :

Théorème 2.5. (*Prolongement de fonction*) Pour $f \in H_0^1(\Omega)$, posons :

$$\tilde{f} = \begin{cases} f & \text{sur } \Omega, \\ 0 & \text{sur } \mathbb{R}^p \setminus \Omega. \end{cases} \quad (2.2)$$

Alors $\tilde{f} \in H^1(\mathbb{R}^p)$ et l'application qui à f associe $\tilde{f}$ est un isomorphisme entre $(H_0^1(\Omega), \|\cdot\|_{H^1(\Omega)})$ et $(H^1(\mathbb{R}^p), \|\cdot\|_{H^1(\mathbb{R}^p)})$.

Opérateur et problème d'évolution dans un espace de Hilbert

Pour définir des théorèmes sur les opérateurs solutions d'un problème d'évolution dans un espace de Hilbert [22], nous avons d'abord besoin de définir des propriétés sur les opérateurs autoadjoints compacts [60] :

Définition 2.6. (*Opérateur autoadjoint*) Soit H un espace de Hilbert et que $T \in \mathcal{L}(H)$. Identifiant H' comme le dual de H , on peut considérer que $T^* \in \mathcal{L}(H)$. On dit qu'un opérateur $T \in \mathcal{L}(H)$ est autoadjoint si $T^* = T$ c'est-à-dire

$$(Tu, v) = (u, Tv) \quad \forall u, v \in H.$$

Ces propriétés sur les opérateurs engendrent des propriétés sur les fonctions propres des opérateurs compacts autoadjoints [60] :

Théorème 2.7. Soit T un opérateur autoadjoint compact sur un espace de Hilbert H séparable. Alors H admet une base hilbertienne formée de vecteurs propres de T .

Pour les problèmes d'évolution dans un espace de Hilbert, les opérateurs solutions doivent avoir les propriétés suivantes [22] :

Définition 2.8. (*Opérateur monotone maximal*) Soit un espace de Hilbert H séparable et $T : D(T) \subset H \rightarrow H$ un opérateur linéaire de domaine $D(T)$. On dit que T est monotone si

$$\forall v \in D(T), (Tv, v)_H \geq 0,$$

et que T est maximal si :

$$\forall f \in H, \exists v \in D(T), v + Tv = f.$$

Avec ces propriétés, le Théorème de Hille-Yosida [22] assure l'existence et l'unicité d'un problème d'évolution sur un espace de Hilbert :

Théorème 2.9. (*Théorème Hille-Yosida*) Soit T un opérateur monotone maximal dans un espace de Hilbert H où le domaine de T est tel que $D(T) \subset H$. Alors, pour tout $f \in C^\infty(H)$ et pour tout u_0 appartenant au domaine de T , $u_0 \in D(T)$, il existe une fonction

$$u \in H^2([0, +\infty[; H),$$

unique telle que

$$\begin{cases} \partial_t u - Tu = 0 & \text{sur } \mathbb{R}^+ \times \Omega, \\ u(t = 0) = f, & \text{sur } \Omega. \end{cases} \quad (2.3)$$

De plus, on a :

$$|u(t)| \leq |f| \quad \text{et} \quad |\partial_t u| = |Tu(t)| \leq |Tf|, \quad \forall t \geq 0.$$

2.3.2 Classification via l'opérateur de Laplace

Dans la suite, après avoir introduit l'opérateur de Laplace, ses éléments spectraux seront étudiés.

Opérateur de Laplace

D'après [60], si $f \in L^2(\Omega)$, alors une solution du Laplacien sur Ω avec pour second membre f est, par définition, un élément $u \in H_0^1(\Omega)$ tel que :

$$(\mathcal{P}_\Omega) \begin{cases} -\Delta u = f \text{ sur } \Omega, \\ u = 0 \text{ sur } \partial\Omega, \end{cases}$$

Proposition 2.10. *Si $f \in L^2(\Omega)$, les assertions suivantes sont équivalentes :*

1. $u \in H_0^1(\Omega)$ et $-\Delta u = f$,
2. $u \in H_0^1(\Omega)$ et $(u|v)_{H_0^1(\Omega)} = (f|v)_{L^2(\Omega)}$ pour tout $v \in H_0^1(\Omega)$.

L'opérateur solution du problème $(\mathcal{P})$ est introduit et possède les propriétés suivantes [60] :

Proposition 2.11. *L'opérateur défini par :*

$$\begin{aligned} T : H_0^1(\Omega) &\rightarrow H_0^1(\Omega) \\ v &\mapsto u \text{ tel que } -\Delta u = v \end{aligned}$$

est injectif, compact, autoadjoint positif sur $H_0^1(\Omega)$.

On peut donc appliquer à l'opérateur T les résultats sur le spectre des opérateurs autoadjoints compacts (Théorème 2.7). Si $\lambda \in \mathbb{C}$ et $u \in H_0^1(\Omega)$ est non nul, on a :

$$-\Delta u = \lambda u \iff T(\lambda u) = u \iff (\lambda \neq 0 \text{ et } Tu = \frac{1}{\lambda}u).$$

Donc λ est une valeur propre du Laplacien avec condition Dirichlet si et seulement si $\lambda \neq 0$ et $1/\lambda$ est une valeur propre de T dans le cas où les fonctions propres sont les mêmes.

Les fonctions propres associées à l'opérateur de Laplace ont les propriétés suivantes [60] :

Théorème 2.12. *Il existe une base Hilbertienne $(v_n)_{n \geq 1}$ de $L^2(\Omega)$ et il existe une suite $(\lambda_n)_{n \geq 0}$ de réels avec $\lambda_n > 0$ et $\lambda_n \rightarrow \infty$ tels que : $0 < \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_n \leq \dots$*

$$\begin{cases} v_n \in H_0^1(\Omega), \\ -\Delta v_n = \lambda_n v_n \text{ sur } \Omega, \\ v_n = 0 \text{ sur } \partial\Omega. \end{cases}$$

On dit que les $(\lambda_n)_n$ sont les valeurs propres de l'opérateur $-\Delta$ (avec conditions de Dirichlet) et que les $(v_n)_n$ sont les fonctions propres associées.

Enfin, on donne un dernier résultat de régularité, utile dans la suite [28].

Proposition 2.13. *Soit Ω un ouvert de $\mathbb{R}^p$ et $f \in L_{loc}^2(\Omega)$ et u une solution de l'équation :*

$$\Delta u = f \text{ sur } \Omega$$

alors ses dérivées secondes $\partial_{x_\alpha x_\beta}^2 u$ sont des fonctions appartenant à $L_{loc}^2(\Omega)$.

Propriété des fonctions propres de l'opérateur de Laplace

Dans un souci de clarté, nous fixons, pour la suite, le nombre de clusters à k . Considérons alors un ouvert Ω de $\mathbb{R}^p$ constitué de k ouverts connexes distincts, notés $\Omega_1, \dots, \Omega_k$. La frontière $\partial\Omega = \bigcup_{i=1}^k \partial\Omega_i$ avec $\bigcap_{i=1}^k \partial\Omega_i = \emptyset$ est supposée suffisamment régulière telle que l'opérateur trace soit bien posé sur $\partial\Omega$ et la décomposition spectrale de l'opérateur Laplacien soit bien définie. Dans ce cas, la frontière $\partial\Omega$ sera supposée au moins C^1 sur $\partial\Omega_i$, pour $i = \{1, \dots, k\}$. On considère le problème $(\mathcal{P}_\Omega^L)$:

$$(\mathcal{P}_\Omega^L) \begin{cases} \text{Trouver } \lambda_n \text{ et } v_n \in H_0^1(\Omega), \\ -\Delta v_n = \lambda_n v_n \text{ sur } \Omega, \\ v_n = 0 \text{ sur } \partial\Omega, \end{cases}$$

et les problèmes sur les Ω_i respectivement pour $i = \{1, \dots, k\}$:

$$(\mathcal{P}_{\Omega_i}^L) \begin{cases} \text{Trouver } \lambda_{n,i} \text{ et } v_{n,i} \in H_0^1(\Omega_i), \\ -\Delta v_{n,i} = \lambda_{n,i} v_{n,i} \text{ sur } \Omega_i, \\ v_{n,i} = 0 \text{ sur } \partial\Omega_i. \end{cases}$$

D'après [60], déterminer les valeurs propres $(\lambda_n)_{n>0}$ et les fonctions propres associées $(v_n)_{n>0}$ du problème $(\mathcal{P}_\Omega^L)$ suivant est équivalent à définir une fonction dépendant des valeurs propres $(\lambda_{n,i})_{n>0}$ des fonctions propres du Laplacien avec conditions de Dirichlet respectivement sur Ω_i , pour $i = \{1, \dots, k\}$. Ce résultat est énoncé dans la proposition suivante :

Proposition 2.14. *Si $u \in H_0^1(\Omega)$ alors $-\Delta u = \lambda u$ sur Ω si et seulement si :*

$$u_i = u|_{\Omega_i} \in H_0^1(\Omega_i) \text{ vérifie } -\Delta u_i = \lambda u_i \text{ sur } \Omega_i, \text{ pour } i = \{1, \dots, k\}.$$

Soit encore $\{\lambda_n, n \in \mathbb{N}\} = \bigcup_{i=1}^k \{\lambda_{n,i}, n \in \mathbb{N}\}$.

Démonstration. Si on considère $v_{n,i}$ solution sur Ω_i de $(\mathcal{P}_{\Omega_i}^L)$, une fonction propre de l'opérateur Laplacien Δ associée à la valeur propre $\lambda_{n,i}$ et si le support de $v_{n,i} \in \Omega_i$ est étendu du l'ensemble Ω en imposant :

$$\forall i \in \{1, \dots, k\}, \overline{v_{n,i}} = \begin{cases} v_{n,i} \text{ sur } \Omega_i, \\ 0 \text{ sur } \Omega \setminus \Omega_i \end{cases} \quad (2.4)$$

alors $\overline{v_{n,i}}$ est aussi fonction propre du Laplacien $-\Delta$ dans $H_0^1(\Omega)$ associée à la même valeur propre $\lambda_{n,i}$.

Réciproquement, tant que les Ω_i sont disjoints pour $i = \{1, \dots, k\}$, l'extension par 0 à l'ensemble entier Ω d'une fonction de $H_0^1(\Omega_i)$ est une fonction de $H_0^1(\Omega)$ d'après [26]. Par conséquent, la réunion de deux ensembles de fonctions propres $\{(v_{n,i})_{n>0}, i \in \{1, \dots, k\}\}$ est une base Hilbertienne de $H_0^1(\Omega)$. $\square$

Cette propriété est illustrée dans la figure 2.4 dans le cadre de donuts via une discrétisation par éléments finis.

Donc les fonctions propres de l'opérateur T permettent d'isoler les composantes connexes Ω_i pour $i \in \{1, \dots, k\}$ d'après [62].

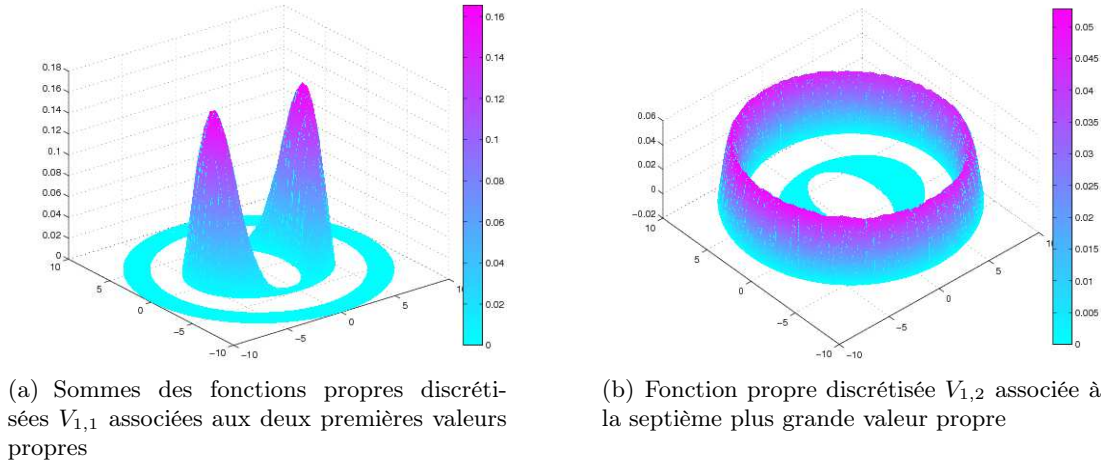


FIGURE 2.4 – Exemple des donuts

Remarque 2.15. *L'exemple précédent montre que les premières fonctions propres à support sur les différentes composantes connexes ne sont pas associées aux plus grandes valeurs propres : la deuxième composante connexe correspond (figure 2.4 (b)) à la 7ème plus grande valeur propre de $\{\lambda_n, n \in \mathbb{N}\}$. Cela permet de soulever le rôle de la normalisation dans l'algorithme 1 de spectral clustering. La normalisation permet de classer les 2 premières fonctions propres, associées aux plus grandes valeurs propres, comme étant la plus grande fonction propre sur chaque composante connexe Ω_1 et Ω_2 .*

Maintenant, il faut lier les fonctions propres de l'opérateur de Laplace avec conditions de Dirichlet avec celles de l'opérateur de la chaleur avec conditions de Dirichlet.

2.3.3 Classification via l'opérateur de la chaleur avec conditions de Dirichlet

Comme pour l'opérateur de Dirichlet, l'opérateur de la chaleur avec conditions de Dirichlet est d'abord introduit avant d'étudier une propriété de clustering sur ses fonctions propres.

Opérateur de la chaleur avec conditions de Dirichlet

Soit $f \in L^2(\Omega)$, on considère maintenant l'équation de la chaleur avec conditions de Dirichlet sur la frontière de l'ouvert borné $\bigcup_{i=1}^k \Omega_i$:

$$(\mathcal{P}_\Omega^D) \begin{cases} \partial_t u - \Delta u = 0 \text{ sur } \mathbb{R}^+ \times \Omega, \\ u = 0, \text{ sur } \mathbb{R}^+ \times \partial\Omega, \\ u(t = 0) = f, \text{ sur } \Omega. \end{cases}$$

Alors l'opérateur solution de ce problème est défini comme suit [60] :

Proposition 2.16. *Soit $f \in H_0^1(\Omega)$. Il existe une unique fonction u définie sur $[0, +\infty[$ à valeur dans $H_0^1(\Omega)$, différentiable sur $[0, +\infty[$ et satisfaisant les conditions suivantes :*

- $\partial_t u = \Delta u(t)$ pour tout $t > 0$.
- $\lim_{t \rightarrow 0^+} u(t) = f$ dans $H_0^1(\Omega)$.

En outre, cette fonction u est donnée par :

$$u(t, x) = \sum_{n=0}^{\infty} (f|v_n)_{H_0^1} e^{-\lambda_n t} v_n(x), \text{ pour tout } t > 0, x \in \Omega \quad (2.5)$$

où les v_n sont les fonctions vérifiant $(\mathcal{P}_\Omega^L)$ et les séries sont convergentes sur $H_0^1(\Omega)$. La fonction u est appelée solution du problème de la chaleur sur Ω avec une donnée initiale f et des conditions de Dirichlet.

Démonstration. L'opérateur de Laplace $-\Delta$ est injectif, compact, autoadjoint positif sur $H_0^1(\Omega)$ d'après la proposition 2.11. Donc d'après le théorème 2.9 de Hille-Yosida, il existe une unique fonction u à valeurs dans $H_0^1(\Omega)$ telle que :

$$\begin{cases} \partial_t u - \Delta u = 0 \text{ sur } \mathbb{R}^+ \times \Omega, \\ u(t = 0) = f, \text{ sur } \Omega. \end{cases}$$

Soit $u(t, x)$ définie par (2.5). Alors comme les v_n vérifient $(\mathcal{P}_\Omega^L)$ on a :

$$-\Delta u(x, t) = - \sum_{n=0}^{\infty} (f|v_n)_{H_0^1} e^{-t\lambda_n} \Delta v_n(x) = \sum_{n=0}^{\infty} (f|v_n)_{H_0^1} e^{-t\lambda_n} \lambda_n v_n(x).$$

De plus, la dérivée partielle par rapport à t est égale à :

$$\partial_t u(t, x) = - \sum_{n=0}^{\infty} (f|v_n)_{H_0^1} \lambda_n e^{-t\lambda_n} v_n(x)$$

En sommant les deux termes, on obtient : $(\partial_t - \Delta)u(x, t) = 0$ et $u(t = 0) = f$. C'est donc la solution. $\square$

Proposition 2.17. Soit l'opérateur solution à valeurs dans $H^2(\Omega) \cap H_0^1(\Omega)$ associé au problème $(\mathcal{P}_\Omega^D)$ défini, pour $f \in L^2(\Omega)$, par :

$$S_D^\Omega(t)f(x) = \int_{\Omega} K_D(t, x, y)f(y)dy, \quad x \in \mathbb{R}^p,$$

où K_D^Ω est le noyau de Green du problème $(\mathcal{P}_\Omega^D)$. Alors les fonctions propres de l'opérateur de Laplace v_n vérifiant $(\mathcal{P}_\Omega^L)$ sont fonctions propres de l'opérateur $S_D^\Omega(t)$ c'est-à-dire : pour tout $i \in \{1, \dots, k\}$ et $n > 0$,

$$S_D^\Omega(t)\overline{v_{n,i}} = e^{-\lambda_{n,i}t}\overline{v_{n,i}}. \quad (2.6)$$

Démonstration. Soit v_m vérifiant $(\mathcal{P}_\Omega^L)$ les fonctions propres de l'opérateur de Laplace, pour tout $m > 0$. D'après l'unicité de solution (proposition 2.16), on a : $S_D^\Omega v_m = \sum_{n=0}^{\infty} (v_m|v_n)_{H_0^1} e^{-t\lambda_n} v_n(x)$. Or, d'après la proposition 2.14, les $(v_n)_n$ décrivent une base Hilbertienne normée de $H_0^1(\Omega)$ donc : pour tout $m > 0$,

$$S_D^\Omega(t)v_m = e^{-t\lambda_m}v_m(x).$$

Donc les fonctions propres $(v_n)_n$ de l'opérateur $-\Delta$ sont fonctions propres de $(\mathcal{P}_\Omega^D)$. $\square$

Dans la suite, on note K_D^Ω le noyau de Green du problème $(\mathcal{P}_\Omega^D)$. Les fonctions propres de l'opérateur $S_D^\Omega(t)$ et celles de l'opérateur de Laplace Δ sont les mêmes. Ainsi, si v_n est une fonction propre de $(\mathcal{P}_\Omega^L)$ associée à la valeur propre λ_n alors v_n est aussi fonction propre de $S_D^\Omega(t)$ associée à la valeur propre $\exp^{-t\lambda_n}$ dans le sens que $S_D^\Omega(t)v_n(x) = e^{-t\lambda_n}v_n(x)$. Alors la propriété géométrique sur les fonctions propres est préservée, i.e le support des fonctions propres peut être définie sur une seule composante connexe Ω_i , pour $i \in \{1, \dots, k\}$.

Propriété de clustering

A partir de cette propriété géométrique 2.14 sur les fonctions propres de l'opérateur de la chaleur avec conditions de Dirichlet, un critère de clustering sur l'espace de projection spectrale peut maintenant être établi.

Théorème 2.18 (Critère de clustering). *Pour tout point $x \in \Omega$ et $\varepsilon > 0$, on notera ρ_x^ε une régularisée de la fonction Dirac centrée en x , c'est-à-dire telle que $\rho_x \in C^\infty(\Omega, [0, 1])$, $\rho_x^\varepsilon(x) = 1$ et $\text{supp}(\rho_x^\varepsilon) \subset \mathcal{B}(x, \varepsilon)$.*

Il existe une famille de fonctions $\overline{v_{n,i}}$ pour $i \in \{1, \dots, k\}$ et $n > 0$ qui vérifient (2.6) et telles que pour tout $x \in \Omega$ et pour tout $i \in \{1, \dots, k\}$ et pour tout $t > 0$, on obtient le résultat suivant :

$$[\exists \varepsilon_0 > 0, \forall \varepsilon \in]0, \varepsilon_0[, \exists n > 0, (S_D(t)\rho_x^\varepsilon | \overline{v_{n,i}})_{L^2(\Omega)} \neq 0] \iff x \in \Omega_i \quad (2.7)$$

où $(f|g)_{L^2(\Omega)} = \int_\Omega f(y)g(y)dy, \forall (f, g) \in L^2(\Omega)$ est le produit scalaire L^2 usuel.

Démonstration. On considère la famille des $\overline{v_{n,i}}$ fonctions propres de $S_D^{\Omega_i}$ décrits par la proposition 2.16.

On démontre par la contraposée le sens direct de l'équivalence. Soit donc $i \in \{1, \dots, k\}$ et un point $x \in \Omega_j$ avec j quelconque, $j \neq i$. Soit $d_x = d(x, \partial\Omega_j) > 0$ la distance de x au bord de Ω_j . Par hypothèse sur Ω , on a $d_0 = d(\Omega_i, \Omega_j) > 0$. Donc pour tout $\varepsilon \in]0, \inf(d_x, d_0)[$, $\mathcal{B}(x, \varepsilon) \subset \Omega_j$. Alors pour tout $t > 0$, $\text{supp}(S_D(t)\rho_x^\varepsilon) \subset \Omega_j$ et donc, d'après la proposition 2.16, pour $n > 0$, $(S_D(t)\rho_x^\varepsilon | \overline{v_{n,i}})_{L^2(\Omega)} = 0$. Donc il n'existe pas de $\varepsilon_0 > 0$ qui vérifie la première partie de l'implication (2.7).

Réciproquement, soit $x \in \Omega_i$ et $\varepsilon \in]0, \inf(d_x, d_0)[$, $\mathcal{B}(x, \varepsilon) \subset \Omega_i$. Donc le support de ρ_x^ε est dans Ω_i . Comme les $(v_{n,i})_{n>0}$ sont une base hilbertienne de $L^2(\Omega_i)$ et que $\rho_x^\varepsilon(x) = 1 \neq 0$ alors il existe un $n > 0$ tel que $(\rho_x^\varepsilon | v_{n,i}) \neq 0$ (sinon ρ_x^ε serait la fonction identiquement nulle sur Ω_i). Dans ce cas là, $(S_D(t)\rho_x^\varepsilon | \overline{v_{n,i}})_{L^2(\Omega)} = e^{-\lambda_{n,i}t}(\rho_x^\varepsilon | v_{n,i}) \neq 0$. $\square$

Remarque 2.19. *A titre d'exemple, on peut considérer ρ_x^ε une fonction régularisante positive centrée en x , C^∞ sur Ω et dont le support est inclus dans Ω_i , par exemple : pour $\varepsilon > 0$ assez petit,*

$$\rho_x^\varepsilon(z) = \begin{cases} \exp\left(-\frac{1}{-\frac{|z-x|}{\varepsilon}+1}\right), & \text{pour } |z| < \varepsilon, \\ 0, & \text{for } |z| \geq \varepsilon. \end{cases} \quad (2.8)$$

On peut encore relâcher la contrainte de l'implication (2.7) en se limitant à demander l'existence d'une suite $(\varepsilon_m)_m$ décroissant vers 0 vérifiant la propriété. Cette relaxation élargit l'éventail des fonctions régularisantes susceptibles de vérifier le critère de clustering.

Corollaire 2.20 (Critère de clustering). *Pour tout point $x \in \Omega$ et $\varepsilon > 0$, on notera ρ_x^ε une régularisée de la fonction Dirac centrée en x , c'est-à-dire telle que $\rho_x \in C^\infty(\Omega, [0, 1])$, $\rho_x^\varepsilon(x) = 1$ et $\text{supp}\rho_x^\varepsilon \subset \mathcal{B}(x, \varepsilon)$.*

Il existe une famille de fonctions $\overline{v_{n,i}}$ pour $i \in \{1, \dots, k\}$ et $n > 0$ qui vérifient (2.6) et telles que pour tout $x \in \Omega$ et pour tout $i \in \{1, \dots, k\}$ et pour tout $t > 0$, on obtient le résultat suivant :

$$[\exists \varepsilon_0 > 0, \exists (\varepsilon_m)_m \in]0, \varepsilon_0[, \lim_{m \rightarrow \infty} \varepsilon_m = 0, \exists n > 0, (S_D(t)\rho_x^{\varepsilon_m} | \overline{v_{n,i}})_{L^2(\Omega)} \neq 0] \iff x \in \Omega_i \quad (2.9)$$

où $(f|g)_{L^2(\Omega)} = \int_\Omega f(y)g(y)dy, \forall (f, g) \in L^2(\Omega)$ est le produit scalaire L^2 usuel.

Démonstration. En reprenant la preuve de l'implication directe dans la démonstration du théorème précédent, on monte que pour une suite de valeurs positives $(\varepsilon_m)_m$ telle que $\varepsilon_m \rightarrow 0$ quand $m \rightarrow \infty$, si m est tel que $\varepsilon_m < \inf(d_x, d_0)$ alors $\text{supp}(S_D(t)\rho_x^{\varepsilon_m}) \subset \Omega_j$ et donc, pour $n > 0$, $(S_D(t)\rho_x^{\varepsilon}|_{\overline{v_{n,i}}})_{L^2(\Omega)} = 0$.

Réciproquement, si $x \in \Omega_i$, par le théorème 2.18, il existe un $\varepsilon_0 > 0$ tel que $\forall \varepsilon \in]0, \varepsilon_0[, \exists n > 0$, $(S_D(t)\rho_x^{\varepsilon}|_{\overline{v_{n,i}}})_{L^2(\Omega)} \neq 0$. On construit alors la suite de terme général $\varepsilon_m = \varepsilon_0/m$ vérifie les hypothèses (2.7). $\square$

2.3.4 Classification via l'opérateur de la chaleur dans $\mathbb{R}^p$

Dans la suite, on considère l'équation de la chaleur définie sur $\mathbb{R}^p$. Un lien entre l'opérateur de la chaleur dans $\mathbb{R}^p$ et celui dans Ω est étudié en comparant les noyaux de Green respectifs et en définissant une propriété de clustering pour ce nouvel opérateur.

Approximation de noyaux de la chaleur

Le coefficient de l'affinité Gaussienne A_{ij} défini dans (1.1) peut être interprété comme une représentation du noyau de la chaleur évaluée aux points x_i et x_j . En effet, soit $K_H(t, x) = (4\pi t)^{-\frac{p}{2}} \exp\left(-\frac{|x|^2}{4t}\right)$ le noyau de la chaleur défini sur $\mathbb{R}^+ \setminus \{0\} \times \mathbb{R}^p$. Notons que l'affinité Gaussienne entre deux points distincts x_i et x_j est définie par $A_{ij} = (2\pi\sigma^2)^{-\frac{p}{2}} K_H(\sigma^2/2, x_i, x_j)$. On considère maintenant le problème $(\mathcal{P}_{\mathbb{R}^p})$ d'équation de chaleur dans $L^2(\mathbb{R}^p)$, pour $f \in L^2(\mathbb{R}^p)$:

$$(\mathcal{P}_{\mathbb{R}^p}) \begin{cases} \partial_t u - \Delta u = 0 & \text{pour } (x, t) \in \mathbb{R}^p \times \mathbb{R}^+ \setminus \{0\}, \\ u(x, 0) = f & \text{pour } x \in \mathbb{R}^p, \end{cases}$$

Par le théorème 2.9, on sait qu'il existe, pour tout $f \in C^\infty([0, +\infty[; H)$, une unique fonction $u \in H^2([0, +\infty[; H)$. On introduit alors l'opérateur solution dans $L^2(\mathbb{R}^p)$ de $(\mathcal{P}_{\mathbb{R}^p})$:

$$(S_H(t)f)(x) = (K_H(t, x, \cdot)|f)_{L^2(\mathbb{R}^p)} = \int_{\mathbb{R}^p} K_H(t, x, y)f(y)dy, \quad x \in \mathbb{R}^p. \quad (2.10)$$

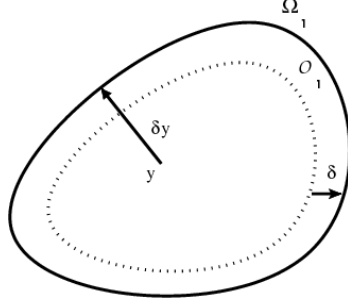
Formellement, pour $f(y) = \delta_{x_j}(y)$ où δ_{x_j} représente la fonction Dirac en x_j , on observe que :

$$\left(S_H\left(\frac{\sigma^2}{2}\right)f\right)(x_i) = K_H\left(\frac{\sigma^2}{2}, x_i, x_j\right) = (2\pi\sigma^2)^{\frac{p}{2}}(A + I_n)_{ij}. \quad (2.11)$$

Ainsi, les propriétés spectrales de la matrice A utilisées dans l'algorithme de clustering spectral semblent être reliées à celles de l'opérateur $S_H(t)$.

Contrairement à la section précédente, on ne peut pas déduire des propriétés similaires comme pour l'opérateur solution de $(\mathcal{P}_\Omega^D)$ car le spectre de l'opérateur $S_H(t)$ est essentiel et ses fonctions propres ne sont pas localisées dans $\mathbb{R}^p$ sans conditions aux frontières. Cependant, les opérateurs $S_H(t)$ et $S_D^\Omega(t)$, pour des valeurs fixes de t , peuvent être comparés et des propriétés asymptotiques sur $S_H(t)$ qui approchent les propriétés précédemment analysées de $S_D^\Omega(t)$ peuvent être établies.

Pour $\varepsilon > 0$, on définit pour tout i un ouvert $\mathcal{O}_i$ qui approche de l'intérieur l'ouvert Ω_i c'est-à-dire telle que $\mathcal{O}_i \subset \bar{\mathcal{O}}_i \subset \Omega_i$ et $\text{Volume}(\Omega_i \setminus \mathcal{O}_i) \leq \varepsilon$. Notons $\delta_i(y)$, la distance d'un point $y \in \mathcal{O}_i$ à la frontière $\partial\Omega_i$ de Ω_i et $\delta_i = \inf_{y \in \mathcal{O}_i} \delta_i(y)$ la distance de $\mathcal{O}_i$ à Ω_i , représentée par la figure 2.5. Par la suite, on notera $\delta = \inf_i \delta_i > 0$ et $\mathcal{O} = \bigcup_i \mathcal{O}_i \subset \bar{\mathcal{O}} \subset \Omega$. Enfin, l'ouvert $\mathcal{O}_i$ est supposé être choisi tel que $\delta > 0$.

FIGURE 2.5 – Approximation de l'ouvert Ω_1 par un ouvert $\mathcal{O}_1$

D'après [22], pour $t > 0$ et pour x et y deux points de Ω , $K_H(t, x, y)$ est une distribution correspondant aux solutions élémentaires de l'équation de la chaleur $(\mathcal{P}_{\mathbb{R}^p}|_{\Omega})$. D'après [62, 10], la différence entre les noyaux K_D et K_H peut être estimée sur l'ouvert $\mathcal{O} \subset \Omega$ et est présentée dans la proposition suivante.

Proposition 2.21. *Soit K_D le noyau de Green du problème $(\mathcal{P}_{\Omega}^D)$ et K_H le noyau de la chaleur défini sur $\mathbb{R}^+ \setminus \{0\} \times \mathbb{R}^p$. Alors K_D est majoré par K_H pour tout couple $(t, x, y) \in \mathbb{R}_+ \times \Omega \times \Omega$:*

$$0 < K_D(t, x, y) < K_H(t, x, y). \quad (2.12)$$

De plus, le noyau de Green K_H est proche du noyau de Green K_D^{Ω} sur chaque ouvert $\mathcal{O}$ et leur différence est majorée, pour $\forall (x, y) \in \mathcal{O} \times \mathcal{O}$, par :

$$0 < K_H(t, x, y) - K_D^{\Omega}(t, x, y) \leq \frac{1}{(4\pi t)^{\frac{p}{2}}} e^{-\frac{\delta^2}{4t}}. \quad (2.13)$$

Démonstration. Par définition, K_H est strictement positif sur la frontière $\partial\Omega$. Or d'après les conditions de Dirichlet, K_D est nul sur $\partial\Omega$. Donc en appliquant le principe du maximum [22], K_D est majoré par K_H sur $\partial\Omega$ c'est-à-dire $0 < K_D(t, x, y) < K_H(t, x, y), \forall (t, x, y) \in \mathbb{R}_+ \times \Omega \times \Omega$. De plus, soit $u_0 \in L^2(\mathcal{O})$. On a d'après le principe du maximum, pour tout $x \in \Omega$:

$$\int_{\Omega} (K_H(t, x, y) - K_D^{\Omega}(t, x, y)) u_0(y) dy \leq \sup_{x \in \partial\Omega} \int_{\Omega} (K_H(t, x, y) - K_D^{\Omega}(t, x, y)) u_0(y) dy$$

Or $\text{supp}(u_0) \subset \mathcal{O}$ donc pour tout $x \in \partial\Omega$, en utilisant l'inégalité de Cauchy-Schwarz :

$$\int_{\mathcal{O}} (K_H(t, x, y) - K_D^{\Omega}(t, x, y)) u_0(y) dy \leq \frac{1}{(4\pi t)^{\frac{p}{2}}} e^{-\frac{\delta(y)^2}{4t}} \|u_0\|_{L^2(\mathcal{O})} \leq \frac{1}{(4\pi t)^{\frac{p}{2}}} e^{-\frac{\delta^2}{4t}} \|u_0\|_{L^2(\mathcal{O})}.$$

Prenons pour u_0 une fonction régularisée de la fonction Dirac centrée en y_0 , par exemple la définition (2.8) alors $\|u_0\|_{L^2(\mathcal{O})} = 1$ et donc d'après le principe du maximum, pour tout $\forall (x, y) \in \mathcal{O} \times \mathcal{O}$:

$$0 < K_H(t, x, y) - K_D^{\Omega}(t, x, y) \leq \frac{1}{(4\pi t)^{\frac{p}{2}}} e^{-\frac{\delta^2}{4t}}. \quad (2.14)$$

□

Remarque 2.22. En étudiant le comportement, $s \mapsto (4\pi s)^{-\frac{p}{2}} e^{-\frac{\delta(y)^2}{4s}}$, t doit décroître plus vite que $\delta(y)^2$ pour tout $y \in \mathcal{O}$ pour minimiser la différence (2.13).

Propriété de classification

Dans la suite, on considère l'opérateur solution de la chaleur $S_H(t)f$ et son support est restreint à l'ouvert $\mathcal{O}_i$, pour $i \in \{1, \dots, k\}$. Pour cela on introduit les notations suivantes. Pour toute fonction $f \in L^2(\Theta)$ où Θ est un ouvert quelconque de $\mathbb{R}^p$, l'opérateur S_H restreint à Θ est défini par :

$$S_H^\Theta(t)f(x) = \int_{\Theta} K_H(t, x, y)f(y)dy, x \in \Theta. \quad (2.15)$$

Enfin, d'après (2.2), on définit $\mathcal{T}_i$, pour $i \in \{1, \dots, k\}$, un opérateur L^2 de prolongement de $L^2(\mathcal{O}_i)$ dans $L^2(\Omega)$ qui étend le support de toute fonction $u \in L^2(\mathcal{O}_i)$ à l'ouvert entier Ω avec :

$$\forall u \in L^2(\mathcal{O}_i), \mathcal{T}_i(u) = \begin{cases} u \text{ sur } \mathcal{O}_i, \\ 0 \text{ sur } \Omega \setminus \mathcal{O}_i. \end{cases} \quad (2.16)$$

Avec ces notations, on peut démontrer la proposition suivante :

Proposition 2.23. *Soient $i \in \{1, \dots, k\}$ et $\widetilde{v_{n,i}} = \mathcal{T}_i(v_{n,i}|_{\mathcal{O}_i})$ où $v_{n,i}$ est la fonction propre normée dans $H^2(\Omega_i) \cap H_0^1(\Omega_i)$ associée à la valeur propre $e^{-\lambda_{n,i}t}$ pour l'opérateur $S_D^{\Omega_i}$. Alors, pour $0 < t < \delta^2$, on a le résultat suivant :*

$$S_H^\mathcal{O}(t)\widetilde{v_{n,i}} = e^{-\lambda_{n,i}t}\widetilde{v_{n,i}} + \eta(t, \widetilde{v_{n,i}}), \quad (2.17)$$

avec $\|\eta(t)\|_{L^2(\mathcal{O})} \rightarrow 0$ quand $t \rightarrow 0, \delta \rightarrow 0$.

Démonstration. Par inégalité triangulaire, $\|\eta(t, \widetilde{v_{n,i}})\|_{L^2(\mathcal{O})} = \|S_H^\mathcal{O}(t)\widetilde{v_{n,i}} - e^{-\lambda_{n,i}t}\widetilde{v_{n,i}}\|_{L^2(\mathcal{O})}$ est majorée par les trois termes suivant :

$$\|\eta(t, \widetilde{v_{n,i}})\|_{L^2(\mathcal{O})} \leq \|S_H^\mathcal{O}(t)\widetilde{v_{n,i}} - S_D^{\Omega_i}(t)\widetilde{v_{n,i}}\|_{L^2(\mathcal{O})} + \|S_D^{\Omega_i}(t)(\widetilde{v_{n,i}} - v_{n,i})\|_{L^2(\mathcal{O})} + \|S_D^{\Omega_i}(t)v_{n,i} - e^{-\lambda_{n,i}t}\widetilde{v_{n,i}}\|_{L^2(\mathcal{O})}. \quad (2.18)$$

Le terme du membre du second membre, $\widetilde{v_{n,i}} - v_{n,i}$ a son support inclus dans $\Omega_i \setminus \mathcal{O}_i$ et tend, par continuité de $S_D^{\Omega_i}$, vers 0 quand $\delta \rightarrow 0$. De plus, l'opérateur $S_D^{\Omega_i}(t)$ tend vers l'opérateur identité quand $t \rightarrow 0$. Donc les deux derniers termes du membre de droite tendent vers 0 quand $(t, \delta) \rightarrow 0$. Pour $v_{n,i}|_{\mathcal{O}_i} \in L^2(\mathcal{O}_i)$, $\|S_H^\mathcal{O}(t)\widetilde{v_{n,i}} - S_D^{\Omega_i}(t)\widetilde{v_{n,i}}\|_{L^2(\mathcal{O})}$ peut être majoré par :

$$\leq \int_{\mathcal{O}_i} \left([S_H^\mathcal{O}(t) - S_D^{\Omega_i}(t)]v_{n,i}|_{\mathcal{O}_i}(y) \right)^2 dy + \int_{\bigcup_{j \neq i} \mathcal{O}_j} \left([S_H^\mathcal{O}(t) - S_D^{\Omega_i}(t)]v_{n,i}|_{\mathcal{O}_i}(y) \right)^2 dy.$$

Avec l'équation (2.12), K_D est majoré par K_H sur $\mathcal{O}$ et en appliquant l'inégalité de Cauchy-Schwarz dans $L^2(\mathcal{O}_i)$,

$$\leq \int_{\mathcal{O}_i} \left\| K_H(t, x, \cdot) - K_D^{\Omega_i}(t, x, \cdot) \right\|_{L^2(\mathcal{O}_i)}^2 \|v_{n,i}|_{\mathcal{O}_i}\|_{L^2(\mathcal{O}_i)}^2 dx + \int_{\bigcup_{j \neq i} \mathcal{O}_j} \int_{\mathcal{O}_i} e^{-\frac{\|x-y\|^2}{4t}} dy dx \|v_{n,i}|_{\mathcal{O}_i}\|_{L^2(\mathcal{O}_i)}^2$$

en appliquant l'équation (2.13) et en minorant $\delta_i(y)$ la distance d'un point $i \in \mathcal{O}$ à $\partial\Omega$ par $\delta_i = \inf_{y \in \mathcal{O}_i} \delta_i(y)$, on obtient :

$$\|\eta(t, \widetilde{v_{n,i}})\|_{L^2(\mathcal{O})} \leq \frac{1}{(4\pi t)^{\frac{p}{2}}} \|v_{n,i}|_{\mathcal{O}_i}\|_{L^2(\mathcal{O}_i)}^2 \left(e^{-\frac{\delta_i^2}{2t} \text{Vol}(\mathcal{O}_i)} + e^{-\frac{\gamma^2}{4t} \text{Vol}(\mathcal{O})} \right). \quad (2.19)$$

avec $\gamma = \inf_{x \in \mathcal{O}_j, y \in \mathcal{O}_i} \|x - y\| \geq 2\delta$. □

La proposition 1 montre que l'opérateur $S_H^\mathcal{O}(t)$ admet "quasiment" des fonctions propres qui sont égales à celles du problème de la chaleur avec conditions de Dirichlet restreintes à $\mathcal{O}_i \subset \Omega_i$. De plus, la proximité avec les fonctions propres du problème avec conditions de Dirichlet dépend du paramètre t . Enfin, comme ces fonctions propres $\widetilde{v_{n,i}}$ ont leur support inclus sur une seule composante connexe, cette propriété peut être utilisée pour caractériser si un point appartient à un ouvert Ω_i ou non.

Remarque 2.24. Une hypothèse sur le paramètre t doit être définie pour préserver le résultat asymptotique de la proposition 1 : t doit tendre vers 0 plus vite que δ^2 c'est-à-dire :

$$t \ll \delta^2. \quad (2.20)$$

Dans la suite, on considérera toujours des t "assez petits" c'est-à-dire qu'ils satisferont entre autres cette hypothèse (2.20).

Théorème 2.25. (Propriété de clustering) Pour tout point $x \in \mathcal{O}$ et $\varepsilon > 0$, on note ρ_x^ε une régularisée de la fonction Dirac centrée en x , c'est-à-dire telle que $\rho_x \in C^\infty(\mathcal{O}, [0, 1])$, $\rho_x^\varepsilon(x) = 1$ et $\text{supp}(\rho_x^\varepsilon) \subset \mathcal{B}(x, \varepsilon)$.

Il existe une famille de fonctions normées, notées $\widetilde{v_{n,i}}$ pour $i \in \{1, \dots, k\}$ et $n > 0$ qui vérifient (2.17) et telles que pour tout $x \in \mathcal{O}$ et pour tout $i \in \{1, \dots, k\}$ on obtient le résultat suivant :

$$\left[\begin{array}{l} \exists \varepsilon_0 > 0, \exists \alpha > 0, \forall \varepsilon \in]0, \varepsilon_0[, \exists n > 0, \forall t > 0 \text{ assez petit,} \\ \widetilde{v_{n,i}} = \operatorname{argmax}_{\{\widetilde{v_{m,j}}, m \in \mathbb{N}, j \in [1, k]\}} |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{m,j}}})_{L^2(\mathcal{O})}| \text{ et } |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})_{L^2(\mathcal{O})}| > \alpha \end{array} \right] \iff x \in \mathcal{O}_i \quad (2.21)$$

Démonstration. On considère la famille des $\widetilde{v_{n,i}}$ décrits dans la proposition 2.23.

On démontre par la contraposée le sens direct de l'équivalence. Soit donc $i \in \{1, \dots, k\}$ et un point $x \in \mathcal{O}_j$ avec j quelconque, $j \neq i$. Soit $d_x = d(x, \partial\mathcal{O}_j) > 0$ la distance de x au bord de $\mathcal{O}_j$. Par hypothèse sur Ω , on a $d_0 = d(\mathcal{O}_i, \mathcal{O}_j) > 0$. Donc pour tout $\varepsilon \in]0, \inf(d_x, d_0)[$, $\mathcal{B}(x, \varepsilon) \subset \mathcal{O}_j$. Alors pour tout $t > 0$, $\text{supp}(S_D^\mathcal{O}(t)\rho_x^\varepsilon) \subset \mathcal{O}_j$ et donc, d'après la proposition 2.16, pour $n > 0$, pour tout $m > 0$ et pour tout $j \in \{1, \dots, k\}$, $j \neq i$, on a :

$$(S_H^\mathcal{O}(t)\widetilde{v_{m,j}}|_{\widetilde{v_{n,i}}})_{L^2(\mathcal{O})} = (\eta(t, \widetilde{v_{m,j}})|_{\widetilde{v_{n,i}}})_{L^2(\mathcal{O})}.$$

Quand t tend vers 0, d'après la proposition 2.23, $(\eta(t, \widetilde{v_{m,j}})|_{\widetilde{v_{n,i}}})_{L^2(\mathcal{O})}$ tend vers 0, pour tout $m > 0$ et tout $j \in \{1, \dots, k\}$, $j \neq i$. Il n'existe donc pas de $\widetilde{w_{m',j'}}$, $j' \neq i$ et de $\alpha > 0$ tel que $(S_H^\mathcal{O}(t)\widetilde{w_{m',j'}}|_{\widetilde{v_{n,i}}})_{L^2(\mathcal{O})} > \alpha$ pour tout t .

Réciproquement, soit $x \in \mathcal{O}_i$ et soit $\varepsilon \in]0, \inf(d_x, d_0)[$, $\mathcal{B}(x, \varepsilon) \subset \mathcal{O}_i$. Le support de ρ_x^ε est dans Ω_i . Comme les $(v_{n,i})_{n>0}$ sont une base hilbertienne de $L^2(\Omega_i)$ et que $\rho_x^\varepsilon(x) = 1 \neq 0$ alors il existe un $n > 0$ tel que $(\rho_x^\varepsilon|_{v_{n,i}}) \neq 0$ (sinon ρ_x^ε serait la fonction identiquement nulle sur $\mathcal{O}_i$). Dans ce cas, $(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})_{L^2(\Omega)} = e^{-\lambda_{n,i}t}(\rho_x^\varepsilon|_{\widetilde{v_{n,i}}}) + (\eta(t, \widetilde{v_{n,i}})|_{\widetilde{v_{n,i}}}) \neq 0$. Quand t tend vers 0, $(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})_{L^2(\Omega)} \rightarrow (\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})$ donc en particulier, pour t assez petit, $|(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})_{L^2(\Omega)}| > |(\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})|/2 > 0$. A contrario, pour tout $j \neq i$, pour tout $m > 0$, $(S_H(t)\widetilde{v_{m,j}}|_{\widetilde{v_{n,i}}})_{L^2(\mathcal{O})} \rightarrow 0$ quand $t \rightarrow 0$ et donc, pour t assez petit, $|(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{m,j}}})_{L^2(\Omega)}| < |(\rho_x^\varepsilon|_{\widetilde{v_{n,i}}})|/2 = \alpha$.

Soit $T_\alpha = \{\widetilde{v_{m,j}}, |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{m,j}}})_{L^2(\mathcal{O})}| > \alpha\}$. Par ce qui précède, $T_\alpha \neq \emptyset$ et n'est composé que de $(\widetilde{v_{n,i}})_n$. Par ailleurs, T_α est fini car

$$\infty > \|S_H^\mathcal{O}(t)\rho_x^\varepsilon\| \geq \sum_{v_{n,i} \in T_\alpha} |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{m,j}}})_{L^2(\mathcal{O})}| \geq \#T_\alpha \alpha.$$

Donc $\operatorname{argmax}_{\{\widetilde{v_{m,j}}, m \in \mathbb{N}, j \in [1, k]\}} |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|_{\widetilde{v_{m,j}}})_{L^2(\mathcal{O})}|$ existe. $\square$

Remarque 2.26. Le théorème 2.18 pourrait être énoncé sous la même forme que le théorème précédent. En effet, l'argument d'imposer que $\widetilde{v_{n,i}}$ réalise le maximum des produits scalaires $(S_H^\mathcal{O}(t)\rho|\widetilde{v_{n,i}})_{L^2(\mathcal{O})}$ n'est pas requis dans l'équivalence de (2.21) car cela se simplifie par un test d'égalité à 0.

De la même manière que pour théorème 2.18, on peut relâcher la contrainte de l'implication (2.21) en se limitant à demander l'existence d'une suite $(\varepsilon_m)_m$ décroissant vers 0 vérifiant la propriété.

Corollaire 2.27 (Critère de clustering). *Pour tout point $x \in \mathcal{O}$ et $\varepsilon > 0$, on notera ρ_x^ε une régularisée de la fonction Dirac centrée en x , c'est-à-dire telle que $\rho_x \in C^\infty(\mathcal{O}, [0, 1])$, $\rho_x^\varepsilon(x) = 1$ et $\text{supp}(\rho_x^\varepsilon) \subset \mathcal{B}(x, \varepsilon)$.*

Il existe une famille de fonctions normées, notées $\widetilde{v_{n,i}}$ pour $i \in \{1, \dots, k\}$ et $n > 0$ qui vérifient (2.17) et telles que pour tout $x \in \mathcal{O}$ et pour tout $i \in \{1, \dots, k\}$ on obtient le résultat suivant :

$$\left[\begin{array}{l} \exists \varepsilon_0 > 0, \exists (\varepsilon_m)_m \in]0, \varepsilon_0[, \lim_{m \rightarrow \infty} \varepsilon_m = 0, \\ \exists \alpha > 0, \forall \varepsilon \in]0, \varepsilon_0[, \exists n > 0, \forall t > 0 \text{ assez petit,} \\ \widetilde{v_{n,i}} = \operatorname{argmax}_{\{\widetilde{v_{m,j}}, m \in \mathbb{N}, j \in \llbracket 1, k \rrbracket\}} |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|\widetilde{v_{m,j}})_{L^2(\mathcal{O})}| \text{ et } |(S_H^\mathcal{O}(t)\rho_x^\varepsilon|\widetilde{v_{n,i}})_{L^2(\mathcal{O})}| > \alpha \end{array} \right] \iff x \in \mathcal{O}_i \quad (2.22)$$

Démonstration. En reprenant la preuve de l'implication directe dans la démonstration du théorème précédent, pour une suite de valeurs positives $(\varepsilon_m)_m$ telle que $\varepsilon_m \rightarrow 0$ quand $m \rightarrow \infty$, si m est tel que $\varepsilon_m < \inf(d_x, d_0)$ alors $\text{supp}(S_D(t)\rho_x^{\varepsilon_m}) \subset \Omega_j$ et donc, pour $n > 0$, $(S_H^\mathcal{O}(t)\rho_x^\varepsilon|\widetilde{v_{n,i}})_{L^2(\mathcal{O})} = (\eta(t, \widetilde{v_{m,j}})|\widetilde{v_{n,i}})_{L^2(\mathcal{O})}$ tend vers 0 quand t tend vers 0, pour tout $m > 0$ et tout $j \in \{1, \dots, k\}$, $j \neq i$.

Réciproquement, si $x \in \mathcal{O}_i$, par le théorème 2.18, il existe un $\varepsilon_0 > 0$ tel que $\forall \varepsilon \in]0, \varepsilon_0[, \exists n > 0$, $(S_D(t)\rho_x^\varepsilon|\widetilde{v_{n,i}})_{L^2(\mathcal{O})} \neq 0$. On construit alors la suite de terme général $\varepsilon_m = \varepsilon_0/m$ qui vérifie les hypothèses (2.21). $\square$

Dans la suite, nous relierons ces propriétés spectrales aux vecteurs propres de la matrice affinité A définie par (1.1) en introduisant une représentation en dimension finie via la théorie des éléments finis.

2.4 Propriété de classification en dimension finie

Le passage à la dimension finie est réalisé par une méthode de Galerkin [87]. Comme l'objectif est d'obtenir des valeurs ponctuelles aux points x_i donnés, les méthodes par éléments finis sont alors privilégiées. Dans la suite, après une présentation des éléments finis de Lagrange, les éléments de la classification spectrale ainsi que la propriété de classification seront interprétés en dimension finie.

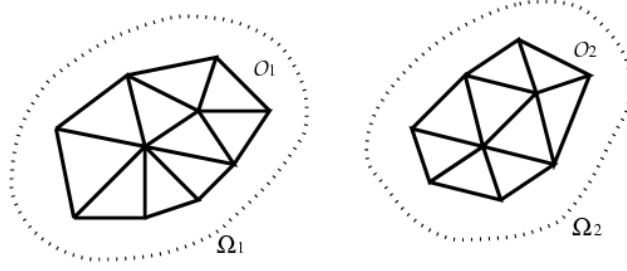
2.4.1 Éléments finis de Lagrange

Soit Σ un ensemble de noeuds dans $\bar{\mathcal{O}}$. Dans la suite, on prendra les éléments finis tels que $\Sigma = S = \{x_i\}_{i=1..N}$ l'ensemble des points utilisés dans l'algorithme 2. Soit $\mathcal{T}_h$ une triangulation ou p -simplexe sur $\bar{\mathcal{O}}$ telle que : $h = \max_{K \in \mathcal{T}_h} h_K$ où h_K est une grandeur caractéristique du triangle K .

Un maillage de chaque partie de l'ensemble $\bar{\mathcal{O}}$ est représentée par la figure 2.6.

Ainsi, on considère une décomposition finie d'un domaine : $\bar{\mathcal{O}} = \cup_{K \in \mathcal{T}_h} K$ dans lequel (K, P_K, Σ_K) satisfait les hypothèses des éléments finis de Lagrange pour tout $K \in \mathcal{T}_h$, d'après [87] :

1. K est un compact, ensemble non-vidé de $\mathcal{T}_h$ tel que $\text{supp}(K) \subset \mathcal{O}$,

FIGURE 2.6 – Maillage du fermé $\bar{\mathcal{O}} : x_i \in \bar{\Omega}, \forall i \in \{1, \dots, N\}$

2. P_K est un espace vectoriel de dimension finie engendré par les fonctions chapeaux : $(\phi_i)_i$ tel que $P_K \subset H^1(K)$,
3. $\Sigma = \cup_{K \in \mathcal{T}_h} \Sigma_K$.

Définition 2.28. (*Espace d'approximation*) Soit l'espace d'approximation de dimension finie :

$$V_h = \{w \in C^0(\bar{\mathcal{O}}); \forall K \in \mathcal{T}_h, w|_K \in P_K\}$$

Pour tout i , $1 \leq i \leq N$, soit la suite de fonctions chapeaux V_h , notée $(\phi_i)_{i \in [1, N]}$, définie sur $\mathcal{O}$ à valeur dans $\mathbb{R}$ telle que :

$$\phi_i(x_j) = \delta_{ij}, \forall x_j \in \Sigma,$$

où δ_{ij} est le symbole de Kronecker en (i, j) . De plus, $\phi_i|_K \in P_K$.

L'espace de dimension finie V_h est généré par la suite $(\phi_i)_i$ avec $V_h = \text{Vect}\{\phi_i\}_i \subset L^2(\mathcal{O})$.

Le passage de la dimension infinie $C^0(\bar{\mathcal{O}})$ à la dimension finie V_h est réalisé par un opérateur d'interpolation.

Définition 2.29. (*Interpolation de Lagrange*) Etant donné un élément fini de Lagrange (K, P, Σ) , on appelle opérateur d'interpolation de Lagrange, l'opérateur défini sur $C^0(\bar{\mathcal{O}})$ à valeurs dans V_h qui à toute fonction v définie sur K associe la fonction $\Pi_h v$ définie par :

$$\Pi_h v = \sum_{i=1}^N v(x_i) \phi_i.$$

On notera par $[\Pi_h v] \in \mathbb{R}^N$ le vecteur des coordonnées de $\Pi_h v$ dans la base des $(\phi_j)_j$, et le produit scalaire associé dans $\mathbb{R}^N$ sera donné par

$$\forall v_h, w_h \in V_h, ([v_h][w_h])_M = [w_h]^T M [v_h] = (w_h|v_h)_{L^2(V_h)},$$

où M représente la matrice de masse des éléments finis : $M_{ij} = (\phi_i|\phi_j)_{L^2}$. La norme associée à $(\cdot|\cdot)_M$ sera notée $|\cdot|_M$.

On rappelle des résultats classiques sur l'approximation par éléments finis [87] :

Proposition 2.30. Si $u \in C^0(\mathcal{O})$ alors, pour $t > 0$, il existe une constante strictement positive $C > 0$, indépendante du pas de discrétisation h telle que :

$$\|u - \Pi_h u\|_{L^2(\mathcal{O})} \leq Ch^2 |u|_{H^2(\mathcal{O})}. \quad (2.23)$$

Remarque 2.31. Pour la dimension de l'espace $p \in \{1, 2\}$, comme les fonctions $v_{n,i}$ sont solutions de $(\mathcal{P}_\Omega^L)$, elles appartiennent à $H_0^1(\Omega)$. Donc par le théorème 2.4, elles sont donc dans $C^0(\Omega)$ et vérifient l'hypothèse du théorème 2.30. Mais comme $v_{n,i}$ vérifient $\Delta v_{n,i} = \lambda_{n,i} v_{n,i} \in L^2(\Omega)$ alors par la proposition 2.13, on a $\partial_{x_\alpha, x_\beta}^2 v_{n,i} \in L^2(\Omega)$, d'où $v_{n,i} \in H^2(\Omega)$, et ainsi $v_{n,i} \in C^0(\Omega)$ si $p < 4$. Par une récurrence directe, en prenant $s > p/4$ entier, pour tout $\alpha \in \mathbb{N}^s$ avec $|\alpha| \leq 2s - 2$, on a $\Delta(\partial_{x_{\alpha_1}, \dots, x_{\alpha_s}}^{| \alpha |} v_{n,i}) = \partial_{x_{\alpha_1}, \dots, x_{\alpha_s}}^{| \alpha |} (\lambda_{n,i} v_{n,i}) \in L^2(\Omega)$ alors par la proposition 2.13, on a $\partial_{x_{\alpha_1}, \dots, x_{\alpha_s}}^{| \alpha |} v_{n,i} \in H^{2s}(\Omega)$, et ainsi par le théorème 2.4, on a $v_{n,i} \in C^0(\Omega)$. Donc, pour tout p , les $v_{n,i}$ satisfont les conditions de la proposition 2.30.

Corollaire 2.32. Soit $u, v \in L^2(\mathcal{O})$ alors, pour $t > 0$:

$$(u|v)_{L^2(\mathcal{O})} = (\Pi_h u | \Pi_h v)_{L^2(V_h)} + O(h^2). \quad (2.24)$$

Démonstration. Par linéarité du produit scalaire, $(\Pi_h u | \Pi_h v)_{L^2(V_h)} - (u|v)_{L^2(\mathcal{O})} = (\Pi_h u - u, \Pi_h v) - (u, v - \Pi_h v)_{L^2(\mathcal{O})}$. D'après l'inégalité de Cauchy-Schwarz et l'équation (2.23), il vient $|(\Pi_h u - u, \Pi_h v)_{L^2(V_h)}| \leq \|\Pi_h u - u\|_{L^2(V_h)} \|\Pi_h v\|_{L^2(V_h)} \leq C_{\Pi_h} C h^2 \|u\|_{H^2(\mathcal{O})} \|v\|_{L^2(\mathcal{O})}$. De la même façon, $|(u, v - \Pi_h v)_{L^2(\mathcal{O})}| \leq C_{\Pi_h} C h^2 \|v\|_{H^2(\mathcal{O})} \|u\|_{L^2(\mathcal{O})}$. D'où le résultat en sommant les deux termes. $\square$

Un maillage régulier est défini comme suit.

Définition 2.33 (Régularité et quasi-uniformité d'un maillage). Une famille de maillages $(\mathcal{T}_h)_h$ est régulière s'il existe une constante $\sigma_0 > 0$ telle que :

$$\forall h, \forall K \in \mathcal{T}_h, \sigma_K = \frac{h_K}{\rho_K} \leq \sigma_0, \quad (2.25)$$

où le diamètre du triangle K , noté $h_K = \text{diam}(K)$, est la longueur du plus grand côté de K , la rondeur $\rho_K = \sup\{\text{diam}(\mathcal{B}); \mathcal{B} \text{ boule de } \mathbb{R}^p, \mathcal{B} \subset K\}$ représente le diamètre de la boule inscrite de K .

Dans une triangulation isotrope, tous les triangles satisfont l'inégalité (2.25) c'est-à-dire une estimation uniforme.

2.4.2 Interprétation des éléments de la classification spectrale

On se propose maintenant d'étudier en dimension finie la discrétisation des résultats précédents. La théorie des éléments finis permet d'introduire l'ensemble de points $S = \{x_i\}_{i=1..N}$. En effet, on considère que les noeuds du maillage Σ correspondent aux points de l'ensemble $S \subset \mathbb{R}^p$ comme le montre la figure 2.7. De plus, l'hypothèse de régularité du maillage correspond à l'hypothèse d'isotropie de la distribution des points du chapitre 1 (1.2).

D'après ces notations, le projection par Π_h du noyau $y \mapsto K_H(t, x, y)$ est définie pour $t > 0$ par :

$$\Pi_h(K_H(t, x, \cdot))(y) = \frac{1}{(4\pi t)^{\frac{p}{2}}} \sum_{i=1}^N \exp\left(-\frac{\|x - x_i\|^2}{4t}\right) \phi_i(y), \quad \forall y \in \mathcal{O}.$$

Alors le projection de l'opérateur solution $\Pi_h S_H^\mathcal{O}(t)$ est définie dans la proposition suivante :

Proposition 2.34. On effectue les hypothèses de régularité suivantes du maillage :

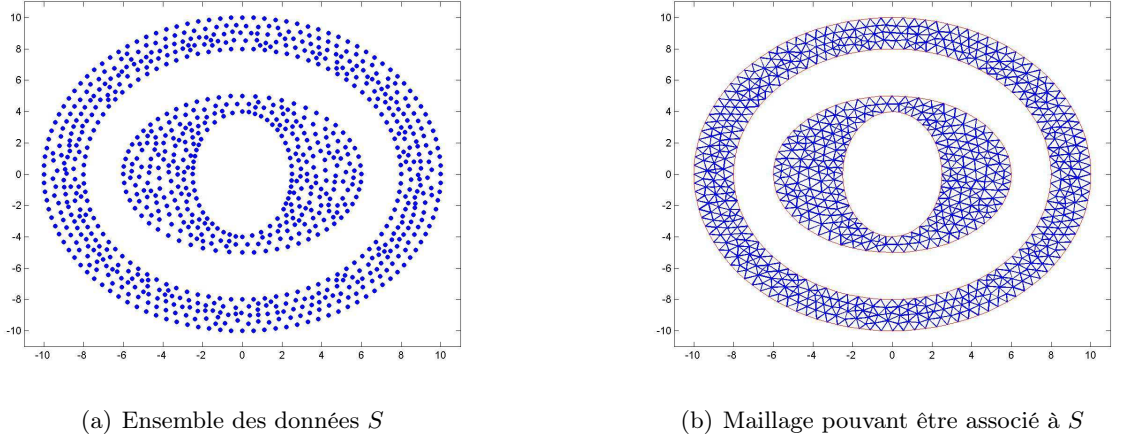


FIGURE 2.7 – Exemple des donuts

1. il existe deux constantes positives α et β ne dépendant pas de h telles que $\forall K \in \mathcal{T}_h, \alpha h^p \leq |K| \leq \beta h^p$,
2. il existe un nombre $\mu > 0$ ne dépendant pas de h tel que pour tout h un sommet quelconque du maillage Σ_h appartient au plus à μ éléments $K \in \mathcal{T}_h$.

Pour $t > 0$ assez petit, la projection de l'opérateur $S_H^\mathcal{O}(t)$ appliqué à ϕ_j , dans la base de $(\phi_k)_k$, est définie par :

$$(4\pi t)^{\frac{p}{2}} \Pi_h(S_H^\mathcal{O}(t)\phi_j)(x) = \sum_{k=1}^N ((A + \mathbb{I}_N)M)_{kj} \phi_k(x) + O\left(\frac{h^{3p+2}}{t^2}\right), \quad \forall 1 \leq j \leq N, \quad (2.26)$$

où A est la matrice affinité définie par (1.1).

Démonstration. D'après la définition (2.15) de l'opérateur $S_H(t)$ restreint à l'ouvert $\mathcal{O}$, pour tout $x \in \mathcal{O}$, on a pour tout $1 \leq j \leq N$:

$$S_H^\mathcal{O}(t)\phi_j(x) = (K_H(t, x, \cdot)|\phi_j)_{L^2(\mathbb{R}^p)} = (K_H(t, x, \cdot)|\phi_j)_{L^2(\mathcal{O})},$$

car $\text{supp}(\phi_j) \subset \mathcal{O}$ pour tout $1 \leq j \leq n$. On note F l'application de V_h dans $L^2(\mathcal{O})$ définie sur la base des $(\phi_j)_j$, pour toute fonction chapeau ϕ_j , $1 \leq j \leq N$, $\forall x \in \mathcal{O}$:

$$F(\phi_j)(x) = (\Pi_h K_H(t, x, \cdot)|\phi_j)_{L^2(\mathcal{O})} = \int_{\mathcal{O}} \Pi_h K_H(t, x, y) \phi_j(y) dy, \quad (2.27)$$

soit encore, avec les notations des éléments de la matrice de masse M ,

$$= (4\pi t)^{-\frac{p}{2}} \sum_{i=1}^N \exp\left(-\frac{\|x - x_i\|^2}{4t}\right) \int_{\mathcal{O}} \phi_i(y) \phi_j(y) dy = (4\pi t)^{-\frac{p}{2}} \sum_{i=1}^N \exp\left(-\frac{\|x - x_i\|^2}{4t}\right) M_{ij}.$$

Or, pour $t > 0$, $F(\phi_j)$ est une fonction C^∞ sur $\mathcal{O}$ et donc H^2 sur $\mathcal{O}$. Donc la projection Π_h dans l'espace d'approximation V_h de $F(\phi_j)$, $1 \leq j \leq N$ est bien définie, pour tout $x \in \mathcal{O}$ et on a :

$$\Pi_h F(\phi_j)(x) = \Pi_h \left((\Pi_h K_H(t, x, \cdot)|\phi_j)_{L^2(\mathcal{O})} \right), \quad \forall j.$$

D'après les notations de l'affinité (1.1), on obtient pour tout $x \in \mathcal{O}$ et tout j :

$$\Pi_h F(\phi_j)(x) = (4\pi t)^{-\frac{p}{2}} \sum_{k=1}^N \sum_{i=1}^N \left(\exp\left(-\frac{\|x_k - x_i\|^2}{4t}\right) M_{ij} \right) \phi_k(x) = (4\pi t)^{-\frac{p}{2}} \sum_{k=1}^N ((A + \mathbb{I}_N)M)_{kj} \phi_k(x).$$

En utilisant la proposition 2.30, l'erreur en norme L^2 de $\|\Pi_h F(\phi_j) - F(\phi_j)\|_{L^2(\mathcal{O})}$ peut être évaluée. Il existe alors une constante $C > 0$ telle que

$$\|\Pi_h F(\phi_j) - F(\phi_j)\|_{L^2(\mathcal{O})} \leq Ch^2 |F(\phi_j)|_{H^2(\mathcal{O})}.$$

Rappelons que l'application $x \mapsto K_H(t, x, y)$ est C^∞ sur $\mathcal{O}$ donc, pour tout $y \in \mathcal{O}$, $\partial_{x_\alpha} K_H(t, x, y)$ et $\partial_{x_\alpha} \partial_{x_\beta} K_H(t, x, y)$ les dérivées partielles de K_H peuvent être calculées, pour tout $x \in \mathcal{O}$, $\alpha, \beta \in [1, p]$ et $t > 0$:

$$\begin{aligned} \partial_{x_\alpha} K_H(t, x, y) &= -\frac{2}{t} (x - y)_\alpha K_H(t, x, y), \\ \partial_{x_\alpha} \partial_{x_\beta} K_H(t, x, y) &= \begin{cases} \left(\frac{4}{t^2} (x - y)_\alpha (x - y)_\beta - \frac{2}{t} \right) K_H(t, x, y) & \text{si } \alpha \neq \beta, \\ \left(\frac{4}{t^2} (x - y)_\alpha^2 - \frac{2}{t} \right) K_H(t, x, y) & \text{sinon.} \end{cases} \end{aligned} \quad (2.28)$$

Remarquons que l'application $y \mapsto \partial_{x_\alpha} \partial_{x_\beta} K_H(t, x, y)$ peut être majorée sur $\mathcal{O}$ par une fonction $g \in C^\infty(\mathcal{O})$ définie pour t assez petit par $g : y \mapsto \frac{\text{diam}(\mathcal{O})}{t^2} (4\pi t)^{-\frac{p}{2}}$ qui est intégrable sur $\mathcal{O}$. Donc le théorème de dérivation sous le signe intégrale [108] peut être appliqué à la fonction $F(\phi_j)$ avec le produit scalaire $(\cdot, \cdot)_{H^2}$:

$$\partial_{x_\alpha} \partial_{x_\beta} F(\phi_j)(x) = (\Pi_h \partial_{x_\alpha} \partial_{x_\beta} K_H(t, x, \cdot) | \phi_j)_{L^2(\mathcal{O})}. \quad (2.29)$$

En utilisant l'inégalité de Cauchy-Schwarz avec le produit scalaire $(\cdot, \cdot)_{L^2(\mathcal{O})}$:

$$\|\partial_{x_\alpha} \partial_{x_\beta} F(\phi_j)\|_{L^2(\mathcal{O})} \leq \|\Pi_h \partial_{x_\alpha} \partial_{x_\beta} K_H(t, x, \cdot)\|_{L^2(\text{supp}(\phi_j))} \|\phi_j\|_{L^2(\mathcal{O})}$$

Notons que la fonction $y \mapsto \Pi_h \partial_{x_\alpha} \partial_{x_\beta} K_H(t, x, \cdot)$ est majorée en module par g . De plus, en majorant la fonction ϕ_j par 1 sur son support, on obtient :

$$\|\partial_{x_\alpha} \partial_{x_\beta} F(\phi_j)\|_{L^2(\mathcal{O})} \leq |\text{supp}(\phi_j)|^2 \frac{\text{diam}(\mathcal{O})}{t^2} (4\pi t)^{-\frac{p}{2}}.$$

En utilisant les hypothèses de régularité du maillage, on en déduit que $|\text{supp}(\phi_j)| \leq \beta \mu h^p$ et donc

$$\|\partial_{x_\alpha} \partial_{x_\beta} F(\phi_j)\|_{L^2(\mathcal{O})} \leq \beta^2 \mu^2 h^{2p} \frac{\text{diam}(\mathcal{O})}{t^2} (4\pi t)^{-\frac{p}{2}}.$$

Calculons maintenant que $\Pi_h(S_H^\mathcal{O}(t)\phi_j)(x)$. Par définition, $\phi_j \in L^2(\mathcal{O})$ et en appliquant le théorème de Hille-Yosida, l'opérateur solution $S_H^\mathcal{O}(t)(\phi_j)$ de $(\mathcal{P}_{\mathbb{R}^p} | \mathcal{O})$ est à valeurs dans $H^2(\mathcal{O})$. Par le théorème 2.4, $S_H^\mathcal{O}(t)(\phi_j) \in C^0(\bar{\mathcal{O}})$. Donc, pour tout $1 \leq j \leq N$, d'après la continuité de l'application Π_h , il existe une constante C_{Π_h} telle que :

$$\|\Pi_h [S_H^\mathcal{O}(t)(\phi_j) - F(\phi_j)]\|_{L^2(V_h)} \leq C_{\Pi_h} \|S_H^\mathcal{O}(t)(\phi_j) - F(\phi_j)\|_{C^0(\mathcal{O})}.$$

A partir de la définition de l'opérateur $S_H^\mathcal{O}(\phi_j)$ et $F(\phi_j)$, on a

$$\|S_H^\mathcal{O}(t)(\phi_j) - F(\phi_j)\|_{L^2(\mathcal{O})} = \|(K_H(t, x, \cdot) - \Pi_h K_H(t, x, \cdot) | \phi_j)_{L^2(\mathcal{O})}\|_{L^2(\mathcal{O})}.$$

Par l'inégalité de Cauchy-Schwarz, il vient alors :

$$\|S_H^\mathcal{O}(t)(\phi_j) - F(\phi_j)\|_{L^2(\mathcal{O})} \leq \|K_H(t, x, \cdot) - \Pi_h K_H(t, x, \cdot)\|_{L^2(\mathcal{O})} \|\phi_j\|_{L^2(\mathcal{O})}.$$

En appliquant la proposition 2.30, la norme L^2 de l'erreur $\|K_H(t, x, \cdot) - \Pi_h K_H(t, x, \cdot)\|_{L^2(\mathcal{O})}$ peut être évaluée. En effet, il existe une constante $C_2 > 0$ telle que

$$\|K_H(t, x, \cdot) - \Pi_h K_H(t, x, \cdot)\|_{L^2(\mathcal{O})} \leq C_2 h^2 |K_H|_{H^2(\mathcal{O})}.$$

En utilisant respectivement les majorations de $\|\phi_j\|_{L^2(\mathcal{O})}$ et $\|\partial_{x_\alpha} \partial_{x_\beta} K_H\|$, on obtient enfin

$$\|\Pi_h(S_H^\mathcal{O}(t)(\phi_j) - F(\phi_j))\|_{L^2(V_h)} \leq C \frac{h^{3p+2}}{t^2} (4\pi t)^{\frac{-p}{2}}, \quad (2.30)$$

ce qui termine la preuve. $\square$

Remarque 2.35. Le résultat (2.26) de la proposition 2.34 permet de définir une borne minimale pour le paramètre de la chaleur t . En effet, cette approximation est valable si t satisfait la condition suivante :

$$h^{3p+2} < t^2. \quad (2.31)$$

2.4.3 Propriété de classification du produit $(A + \mathbb{I}_N)M$

La propriété 2.34 montre que le produit de la matrice affinité (1.1) de l'Algorithme 1 par la matrice de masse M issue de la discrétisation par éléments finis est interprété comme la projection par Π_h de l'opérateur solution de $(\mathcal{P}_{\mathbb{R}^p})$. On a donc un lien entre les vecteurs propres de la matrice affinité A et des éléments caractérisant les Ω_i , pour $i \in \{1, \dots, k\}$.

Donc, en supposant les mêmes hypothèses que pour la proposition 2.23 de la section précédente, l'équation (2.23) peut être formulée via l'application d'interpolation Π_h .

Proposition 2.36. Soit $v_{n,i}$ une fonction propre normée $S_D^{\Omega_i}(t)$ dans $H^2(\Omega) \cap H_0^1(\Omega)$ associée à la valeur propre $e^{-\lambda_{n,i}t}$ telle que $S_D^{\Omega_i}(t)v_{n,i} = e^{-\lambda_{n,i}t}v_{n,i}$. Soit $i \in \{1, \dots, k\}$ et $\widetilde{v_{n,i}} = \mathcal{T}_i(v_{n,i}|\mathcal{O}_i)$. Alors, pour $W_{n,i} = [\Pi_h \widetilde{v_{n,i}}] \in \mathbb{R}^N$ et $t > 0$, on a :

$$(4\pi t)^{\frac{-p}{2}} (A + \mathbb{I}_N) M W_{n,i} = e^{-\lambda_{n,i}t} W_{n,i} + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right), \quad (2.32)$$

quand $(h, t) \rightarrow 0$ avec $h^{3p+2}/t^2 \rightarrow 0$ et $\delta \rightarrow 0$.

Le lemme suivant est nécessaire pour démontrer le résultat de la proposition 2.36.

Lemme 2.37. Il existe une constante $C > 0$ indépendante de h et de t telle que, pour tout fonction $f \in H^2(\Omega)$, on a :

$$\|\Pi_h S_H^\mathcal{O} \Pi_h f - \Pi_h S_H^\mathcal{O} f\|_{L^2(V_h)} \leq C h^2 e^{-\frac{\delta^2}{4t}} \|f\|_{H^2(\Omega)} \quad (2.33)$$

Démonstration. En utilisant la linéarité et la continuité de Π_h , il existe une constante $C_{\Pi_h} > 0$ telle que :

$$\|\Pi_h S_H^\mathcal{O} \Pi_h f - \Pi_h S_H^\mathcal{O} f\|_{L^2(V_h)} \leq C_{\Pi_h} \|S_H^\mathcal{O}(\Pi_h f - f)\|_{C^0(\bar{\mathcal{O}})} \quad (2.34)$$

Avec la définition de la norme $C^0(\bar{\mathcal{O}})$, l'inégalité de Cauchy-Schwarz et la majoration de l'opérateur $S_H^\mathcal{O}$ (2.35), l'équation (2.23) de la proposition 2.30 est appliquée à $\|\Pi_h f - f\|_{L^2(\mathcal{O})}$:

$$\leq Vol(\mathcal{O}) e^{-\frac{\delta^2}{4t}} \|\Pi_h f - f\|_{L^2(\mathcal{O})} \leq Vol(\mathcal{O}) C h^2 e^{-\frac{\delta^2}{4t}} \|f\|_{H^2(\mathcal{O})}. \quad (2.35)$$

$\square$

Démonstration de la proposition 2.36. On applique l'application Π_h à l'équation de la proposition 2.23. Par inégalité triangulaire, pour $w_{n,i} = \Pi_h \widetilde{v_{n,i}} \in V_h$ et $t > 0$, la différence $\|\psi(t, h)\|_{L^2(V_h)} = \|\Pi_h S_H^\mathcal{O} w_{n,i} - e^{-\lambda_{n,i}t} w_{n,i}\|_{L^2(V_h)}$ est majorée par :

$$\|\psi(t, h)\|_{L^2(V_h)} \leq \|\Pi_h S_H^\mathcal{O} w_{n,i} - \Pi_h S_H^\mathcal{O} \widetilde{v_{n,i}}\|_{L^2(V_h)} + \|\Pi_h S_H^\mathcal{O} \widetilde{v_{n,i}} - e^{-\lambda_{n,i}t} w_{n,i}\|_{L^2(V_h)} \quad (2.36)$$

D'après le lemme 2.37 appliqué à $\widetilde{v_{n,i}}$, on a :

$$\|\Pi_h S_H^\mathcal{O} w_{n,i} - \Pi_h S_H^\mathcal{O} \widetilde{v_{n,i}}\|_{L^2(V_h)} \leq Ch^2 e^{-\frac{\delta^2}{4t}} |\widetilde{v_{n,i}}|_{H^2(\mathcal{O})}. \quad (2.37)$$

De la même façon, la linéarité et la continuité de Π_h sont utilisées pour le second terme du membre de droite (2.36) :

$$\|\Pi_h (S_H^\mathcal{O} \widetilde{v_{n,i}} - e^{-\lambda_{n,i}t} \widetilde{v_{n,i}})\|_{L^2(V_h)} \leq C_{\Pi_h} \|S_H^\mathcal{O} \widetilde{v_{n,i}} - e^{-\lambda_{n,i}t} \widetilde{v_{n,i}}\|_{C^0(\bar{\mathcal{O}})}.$$

En utilisant directement le résultat (2.19) avec l'inégalité de Cauchy-Schwarz, il vient

$$\leq C_{\Pi_h} \text{Vol}(\mathcal{O}) \|\eta(t, \widetilde{v_{n,i}})\|_{L^2(\mathcal{O})} \leq (4\pi t)^{-\frac{p}{4}} \|v_{n,i}|_{\mathcal{O}_i}\|_{L^2(\mathcal{O}_i)} \sqrt{\text{Vol}(\mathcal{O})} \left(e^{-\frac{\delta^2}{2t}} + e^{-\frac{\gamma^2}{4t}} \right)^{\frac{1}{2}} \quad (2.38)$$

avec $\gamma = \inf_{x \in \mathcal{O}_j, y \in \mathcal{O}_i} \|x - y\| \geq 2\delta$. Enfin, en remarquant que par l'inégalité triangulaire,

$$\begin{aligned} \|(4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) M W_{n,i} - e^{-\lambda_{n,i}t} W_{n,i}\|_{L^2(V_h)} &\leq \|(4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) M W_{n,i} - \Pi_h (S_H^\mathcal{O}(t) w_{n,i})\|_{L^2(V_h)} \\ &\quad + \|\Pi_h S_H^\mathcal{O} w_{n,i} - e^{-\lambda_{n,i}t} w_{n,i}\|_{L^2(V_h)} \end{aligned}$$

puis par (2.26), (2.35) et (2.38) et en normant le vecteur $W_{n,i} = [w_{n,i}]$, on en déduit (2.32). $\square$

Cette propriété montre que les fonctions propres de l'opérateur $S_H^\mathcal{O}$ sont asymptotiquement vecteurs propres du produit $(A + \mathbb{I}_N)M$. Pour δ et t suffisamment petits, chaque point dans l'espace de projection représente le coefficient de projection dans l'espace spectral de l'opérateur S_H projeté par Π_h . Suivant le cluster d'appartenance d'un point x , son coefficient de projection dans cet espace spectral sera maximal suivant le vecteur propre associé à la fonction propre caractérisant le support de son cluster. Ainsi on aboutit à la même assignation au cluster pour un ensemble de points suivant si on considère les fonctions propres dans $L^2(\Omega)$ ou bien les projections par Π_h de ces fonctions propres dans l'espace d'approximation V_h . On obtient alors le résultat résumé dans le théorème suivant.

Théorème 2.38. (*Propriété de clustering en dimension finie*) En reprenant les hypothèses de la proposition 2.36, il existe une famille de vecteurs normés, notés $W_{n,i}$ pour $i \in \{1, \dots, k\}$ et $n > 0$ qui vérifient (2.32) et telle que pour tout $x_r \in S$ et pour tout $i \in \{1, \dots, k\}$ on a le résultat suivant :

$$\left[\begin{array}{l} \exists \alpha > 0, \exists n > 0, \forall t > 0 \text{ avec } t \text{ et } h^{3p+2}/t^2 \text{ assez petits,} \\ W_{n,i} = \operatorname{argmax}_{\{W_{m,j}, m \in \mathbb{N}, j \in [1, k]\}} \left| \left(\left((4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) M \right)_{\cdot r} | W_{m,j} \right)_M \right| \\ \text{et } \left| \left(\left((4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) M \right)_{\cdot r} | W_{n,i} \right)_M \right| > \alpha \end{array} \right] \iff x_r \in \mathcal{O}_i, \quad (2.39)$$

où $((A + \mathbb{I}_N)M)_{\cdot r}$ correspond à la $r^{\text{ième}}$ colonne de la matrice $(A + \mathbb{I}_N)M$.

Démonstration. On considère la famille de vecteurs $W_{n,i}$ décrits dans la proposition 2.36 c'est-à-dire les vecteurs normés des coordonnées des fonctions $w_{n,i} : W_{n,i} = \frac{[w_{n,i}]}{||[w_{n,i}]||_M}$ avec $w_{n,i} = \Pi_h \widetilde{v_{n,i}}$. Soit un point $x_r \in S$. Pour tout $m > 0$ et $j \in \{1, \dots, k\}$, on a par le corollaire 2.32 :

$$|(S_H^\mathcal{O}(t)\rho_x^\epsilon|\widetilde{v_{m,j}})_{L^2(\mathcal{O})}| = |(\Pi_h S_H^\mathcal{O}(t)\rho_x^\epsilon|\Pi_h \widetilde{v_{m,j}})_{L^2(V_h)}| + O(h).$$

Donc par le lemme 2.37 et une inégalité de Cauchy-Schwarz, on a :

$$|(\Pi_h S_H^\mathcal{O}(t)\rho_x^\epsilon|\Pi_h \widetilde{v_{m,j}})_{L^2(V_h)}| = |(\Pi_h S_H^\mathcal{O}(t)\Pi_h \rho_x^\epsilon|\Pi_h \widetilde{v_{m,j}})_{L^2(V_h)}| + O(h).$$

Pour tout $\epsilon < \min_{s \neq r} \|x_r - x_s\|_2 = \epsilon_0$, on a $\text{supp}(\rho_{x_r}^\epsilon) \subset \mathcal{B}(x_r, \epsilon) \subset \text{supp}(\phi_r)$. Il existe alors un coefficient non nul β_ϵ tel que $\Pi_h \rho_{x_r}^\epsilon = \beta_\epsilon \phi_r$. Pour tout $\epsilon < \epsilon_0$, on a d'après la proposition 2.34 :

$$\begin{aligned} |((A + \mathbb{I}_N)M)_{\cdot,r} |[w_{m,j}]_M| &= |(\Pi_h S_H^\mathcal{O}(t)\Pi_h \rho_x^\epsilon|\Pi_h \widetilde{v_{m,j}})_{L^2(V_h)}| + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right) \\ &= |(S_H^\mathcal{O}(t)\rho_{x_r}^\epsilon|\widetilde{v_{m,j}})_{L^2(\mathcal{O})}| + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right). \end{aligned} \quad (2.40)$$

Si $x_r \in \mathcal{O}_i$ alors d'après le théorème 2.21, pour tout ϵ assez petit et pour tout $t > 0$, il existe un indice $n > 0$ tel que :

$$\widetilde{v_{n,i}} = \argmax_{\{\widetilde{v_{m,j}}, m \in \mathbb{N}, j \in [1, k]\}} |(S_H^\mathcal{O}(t)\rho_x^\epsilon|\widetilde{v_{m,j}})_{L^2(\mathcal{O})}| \text{ et } |(S_H^\mathcal{O}(t)\rho_{x_r}^\epsilon|\widetilde{v_{n,i}})_{L^2(\mathcal{O})}| > \alpha.$$

Donc d'après l'équation (2.40), on obtient : $|((A + \mathbb{I}_N)M)_{\cdot,r} |[w_{n,i}]_M| > \alpha + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right)$. Si h est assez petit alors l'inégalité précédente peut être formulée comme : il existe un $\beta > 0$ tel que : $|((A + \mathbb{I}_N)M)_{\cdot,r} |[w_{n,i}]_M| > \beta$. A contrario, pour tout $j \neq i$, on a d'après la démonstration du théorème 2.25 : $|((A + \mathbb{I}_N)M)_{\cdot,r} |[w_{m,j}]_M| = |(S_H^\mathcal{O}(t)\rho_x^\epsilon|\widetilde{v_{m,j}})_{L^2(\mathcal{O})}| + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right) = (\eta(t)|\widetilde{v_{m,j}}) + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right)$, avec $\|\eta(t)\|_{L^2(\mathcal{O})} \rightarrow 0$ quand $t \rightarrow 0, \delta \rightarrow 0$. Donc pour t, h et h^{3p+2}/t^2 assez petits, $|((A + \mathbb{I}_N)M)_{\cdot,r} |[w_{m,j}]_M| < \beta$ et ainsi l'argument maximal de (2.39) est obtenu parmi les $[w_{n,i}]$.

Si $x_r \in \mathcal{O}_j$ avec $j \neq i$ alors en reprenant ce qui précède, pour t et h assez petits, $|((A + \mathbb{I}_N)M)_{\cdot,r} |[w_{n,i}]_M| < \beta$ pour tout $n > 0$.

On finit la démonstration en normalisant les $[w_{n,i}]$, donc en obtenant les vecteurs $W_{n,i}$, en notant que :

$$\frac{w_{n,i}}{||[w_{n,i}]||_M} = \frac{w_{n,i}}{||w_{n,i}||_{L^2(V_h)}} = \frac{\Pi_h \widetilde{v_{n,i}}}{||\Pi_h \widetilde{v_{n,i}}||_{L^2(V_h)}} = \Pi_h \widetilde{v_{n,i}} + O(h) = w_{n,i} + O(h),$$

car $||\Pi_h \widetilde{v_{n,i}}||_{L^2(V_h)} - 1 = ||\Pi_h \widetilde{v_{n,i}}||_{L^2(V_h)} - ||\widetilde{v_{n,i}}||_{L^2(V_h)} \leq ||\Pi_h \widetilde{v_{n,i}} - \widetilde{v_{n,i}}||_{L^2(V_h)} \leq Ch^2$ d'après (2.23). $\square$

La figure 2.8 représente les vecteurs propres de la matrice $(A + \mathbb{I}_N)M$ qui sont le plus semblables respectivement aux fonctions propres discrétisées V_1 et V_7 pour la valeur $t = t_h$, où t_h représente l'heuristique (1.2) développée dans le chapitre 1 satisfaisant l'équation (2.11) : $t_h = \frac{\sigma_h}{4} = \frac{D_{\max}}{8N^{\frac{1}{p}}}$.

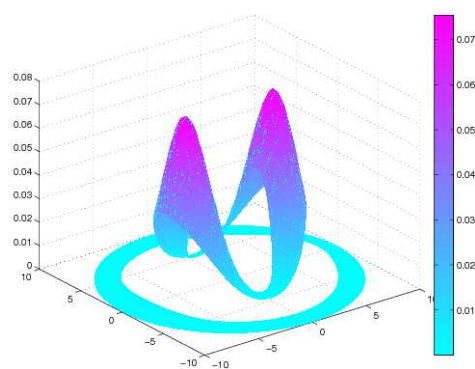
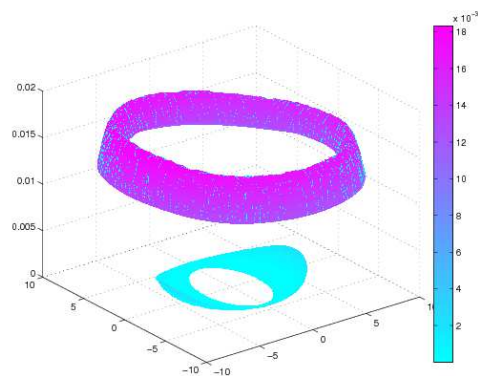
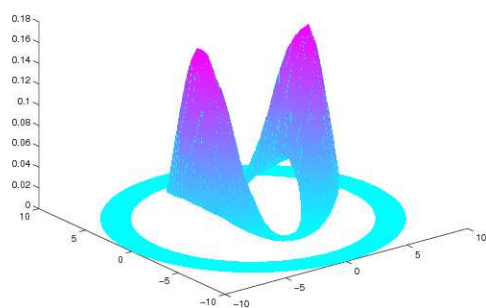
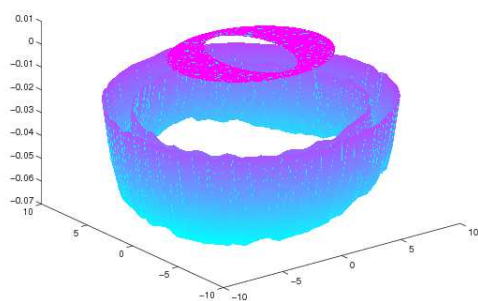
(a) Produit $(A + \mathbb{I}_N)MV_{1,1}$ (b) Produit $(A + \mathbb{I}_N)MV_{1,2}$ (c) Vecteur propre de $(A + \mathbb{I}_N)M$ associé à $V_{1,1}$ pour $t = t_h$ (d) Vecteur propre de $(A + \mathbb{I}_N)M$ associé à $V_{1,2}$ pour $t = t_h$

FIGURE 2.8 – Exemple des donuts

Le théorème 2.38 permet d'expliquer le lien entre le clustering dans l'espace de projection spectrale et le clustering des données initiales. En effet, on retrouve qu'un point x_r appartient à l'ensemble $\mathcal{O}_i$ à partir du produit scalaire de la r^{ieme} colonne de la matrice $(A + \mathbb{I}_N)M$ et des 'quasi' vecteurs propres $V_{n,i}$. Les valeurs de ces produits scalaires représentent les coefficients de la matrice Y de l'algorithme 2. L'étape de k -means revient donc à chercher l'élément qui réalise l'argument maximal de l'équation (2.39).

Les résultats du théorème 2.38 incluent la matrice de masse M dont la définition est totalement dépendante des éléments finis construits et donc de la définition des Ω_i , $i \in \{1, \dots, k\}$, ce qui ne peut pas être vraisemblablement supposé avant d'avoir effectué le clustering. Afin de supprimer cette dépendance, une condensation de masse est envisagée dans la section suivante et les résultats spectraux précédents seront formulés avec une matrice de masse diagonale à la place d'une matrice de masse pleine. L'objectif est de faire disparaître cette connaissance a priori afin de focaliser le travail sur l'étude spectrale de la matrice $A + \mathbb{I}_N$, ou de manière équivalente sur A .

2.4.4 Condensation de masse

Pour supprimer l'influence des éléments finis et de la matrice de masse dans les résultats, on s'intéresse à une condensation de masse qui approche la matrice de masse par une matrice diagonale. Dans la suite, après avoir défini la condensation de masse, les résultats de la propriété 2.36 et du théorème 2.38 sont reformulés et une propriété de clustering de la matrice affinité A sera donc énoncée.

Définition et approximation

Commençons par introduire quelques éléments d'approximation numérique.

Définition 2.39. (*Quadrature numérique*) Soit $\mathcal{I}_K$ la liste des indices de points appartenant à un élément $K \in \tau_h$ donné. Considérons la quadrature numérique pour les polynômes $\mathbb{P}_1$:

$$\int_K \phi(x) dx \approx \sum_{k \in \mathcal{I}_K} \frac{|K|}{p+1} \phi(x_{i_k}), \quad (2.41)$$

où $|K|$ représente le volume de l'élément fini K .

Proposition 2.40. Cette quadrature numérique est exacte pour les polynômes de degré ≤ 1 .

Démonstration. Cette propriété est démontrée dans [26]. □

Donc pour toute fonction $f \in C^0(\mathcal{O})$, on a :

$$\int_{\mathcal{O}} f(x) \phi_j(x) dx \approx \sum_{K \in \mathcal{T}_h} \sum_{i=1}^n f_i \int_K \phi_i(x) \phi_j(x) dx. \quad (2.42)$$

Avec la définition et le support de ϕ_i et ϕ_j , le produit $\phi_i(x_{i_k}) \phi_j(x_{i_k})$ est défini par :

$$\phi_i(x_{i_k}) \phi_j(x_{i_k}) = \delta_{i_k i} \delta_{i_k j}.$$

On obtient donc l'approximation suivante de l'équation (2.42) avec la quadrature numérique (2.41) :

$$\sum_{K \in \mathcal{T}_h} \sum_{i=1}^n f_i \sum_{k=1}^3 \frac{|K|}{3} \phi_i(x_{i_k}) \phi_j(x_{i_k}) \approx \sum_{\substack{K \in \mathcal{T}_h \\ \text{supp}(\phi_j) \cap K \neq \emptyset}} \sum_{i_k \in \mathcal{I}_K} \frac{|K|}{p+1} f_{i_k} \delta_{i_k j}.$$

Notons $\widetilde{M}$ la matrice diagonale issue de la condensation de masse d'élément $\widetilde{M}(i, j)$ défini par :

$$\widetilde{M}(i, j) = \begin{cases} \sum_{\text{supp}(\phi_j) \in K} \frac{|K|}{p+1} & \text{si } i = j, \\ 0 & \text{sinon.} \end{cases} \quad (2.43)$$

D'après [87], la quadrature numérique est exacte pour les polynômes de degré 1. Donc l'erreur d'approximation peut être estimée à l'aide de la proposition suivante.

Proposition 2.41. *Pour tout $f \in C^0(\mathcal{O})$, il existe une constante strictement positive $C_3 > 0$ telle que l'erreur d'approximation satisfait :*

$$\left| \int_K f(x) \phi_j(x) dx - \sum_{i=1}^n f_i \sum_{k \in \mathcal{I}_K} \frac{|K|}{p+1} \phi_i(x_{i_k}) \phi_j(x_{i_k}) \right| \leq C_3 h |f|_{1,K}. \quad (2.44)$$

où $|f|_{1,K} = \left(\int_K |\partial_x f|^2 dx \right)^{\frac{1}{2}}$.

Démonstration. En effectuant le développement de Taylor avec reste intégral à l'ordre 1 de f puis en appliquant l'équation (2.41) à la partie polynomiale, on obtient le résultat. $\square$

Propriété de classification de la matrice affinité A

Le résultat suivant estime l'erreur qui est commise en remplaçant la matrice de masse par la matrice condensée (2.43) dans la proposition 2.36.

Proposition 2.42. *On effectue les hypothèses de régularité du maillage suivantes :*

1. *il existe deux constantes positives α et β ne dépendant pas de h telles que $\forall K \in \mathcal{T}_h, \alpha h^p \leq |K| \leq \beta h^p$,*
2. *il existe un nombre $\mu > 0$ ne dépendant pas de h tel que, pour tout h , un sommet quelconque du maillage Σ_h appartient au plus à μ éléments $K \in \mathcal{T}_h$.*

Soient $\widetilde{M}$ la matrice condensée définie par (2.43) et les $W_{n,i}$ définis par la proposition 2.36. Alors pour $t > 0$ assez petit, on a :

$$(4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) \widetilde{M} W_{n,i} = e^{-\lambda_{n,i} t} W_{n,i} + O\left(t, h^2, \frac{h^{3p+2}}{t^2}\right) + O\left(\frac{h^{p+1}}{t}\right), \quad (2.45)$$

où quand $(h, t) \rightarrow 0$ avec $h^{3p+2}/t^2 \rightarrow 0$, $h^3/t \rightarrow 0$ et $\delta \rightarrow 0$.

Démonstration. Introduisons $\widetilde{F}$ comme l'opérateur linéaire de V_h dans $L^2(\mathcal{O})$ défini, pour tout $\phi_j \in L^2(\mathcal{O})$ et pour $x \in \mathcal{O}$, par

$$\widetilde{F}(\phi_j)(x) = (4\pi t)^{-\frac{p}{2}} \sum_{i=1}^N \exp\left(-\frac{\|x - x_i\|^2}{4t}\right) \widetilde{M}_{ij},$$

avec $\widetilde{M}_{ij}$ définie par (2.43).

Par ailleurs, $\widetilde{F}(\phi_j)$ est une fonction C^∞ et donc $\widetilde{F}(\phi_j)$ est H^2 sur $\mathcal{O}$. Ainsi la projection Π_h de $C^0(\mathcal{O})$ sur V_h appliquée à $\widetilde{F}(\phi_j)$ donne, pour tout $\phi_j \in L^2(\mathcal{O})$:

$$\Pi_h \widetilde{F}(\phi_j)(x) = (4\pi t)^{-\frac{p}{2}} \sum_{k=1}^N \left((A + \mathbb{I}_N) \widetilde{M} \right)_{kj} \phi_k(x), \quad \forall x \in \mathcal{O}.$$

Evaluons $\|\Pi_h S_H^\mathcal{O}(t)(\phi_j) - \Pi_h \tilde{F}(\phi_j)\|_{L^2(V_h)}$. Introduisons l'inégalité triangulaire suivante :

$$\begin{aligned} \|\Pi_h S_H^\mathcal{O}(t)(\phi_j) - \Pi_h \tilde{F}(\phi_j)\|_{L^2(\mathcal{O})} &\leq \|\Pi_h S_H^\mathcal{O}(t)(\phi_j) - \Pi_h F(\phi_j)\|_{L^2(\mathcal{O})} \\ &+ \|\Pi_h F(\phi_j) - \Pi_h \tilde{F}(\phi_j)\|_{L^2(\mathcal{O})}. \end{aligned} \quad (2.46)$$

La majoration de $\|\Pi_h S_H^\mathcal{O}(t)(\phi_j) - \Pi_h F(\phi_j)\|_{L^2(\mathcal{O})}$ est donnée par la proposition 2.34. Déterminer l'erreur revient alors à estimer la norme L^2 sur $\mathcal{O}$ de la différence $\|\Pi_h F(\phi_j) - \Pi_h \tilde{F}(\phi_j)\|_{L^2(\mathcal{O})}$:

$$\left(\int_{V_h} \left| \Pi_h F(\phi_j)(x) - \Pi_h \tilde{F}(\phi_j)(x) \right|^2 dx \right)^{\frac{1}{2}} = \left(\sum_{\substack{K \in \mathcal{T}_h \\ \text{supp} \phi_j \cap K \neq \emptyset}} \int_K \left| \Pi_h F(\phi_j)(x) - \Pi_h \tilde{F}(\phi_j)(x) \right|^2 dx \right)^{\frac{1}{2}}$$

D'après (2.27), pour tout $K \in \mathcal{T}_h$,

$$F(\phi_j)(x) - \tilde{F}(\phi_j)(x) = \int_K \Pi_h K_H(t, x, y) \phi_j(y) dy - (4\pi t)^{-\frac{p}{2}} \sum_{i=1}^N \exp\left(-\frac{\|x - x_i\|^2}{4t}\right) \tilde{M}_{ij}.$$

Par continuité de Π_h et en appliquant l'équation (2.44), il existe alors une constante positive $C_3 > 0$ telle que :

$$\left| F(\phi_j) - \tilde{F}(\phi_j) \right| \leq C_3 h |\Pi_h K_H|_{1,K} \leq C_3 C_{\Pi_h} h |K_H|_{1,K}.$$

Or d'après (2.28), l'application $y \mapsto \partial_{x_\alpha}$ est majorée sur $\mathcal{O}$ par la fonction $g_1 \in C^\infty$ pour t assez petit définie par $g_1 : y \mapsto \frac{\text{diam}(\mathcal{O})}{t} (4\pi t)^{-\frac{p}{2}}$ intégrable sur $\mathcal{O}$. Donc

$$\left| F(\phi_j) - \tilde{F}(\phi_j) \right| \leq C_{\Pi_h} C_3 h \sup_{x_\alpha} \partial_{x_\alpha} K_H|_{1,K} \leq C_{\Pi_h} C_3 \frac{h}{t} \text{diam}(\mathcal{O}) (4\pi t)^{-\frac{p}{2}}.$$

D'après les hypothèses de régularité du maillage, il existe $\beta > 0$ tel que $|K| \leq \beta h^p$ donc pour tout $K \in \mathcal{T}_h$, il suit :

$$\int_K \left| \Pi_h F(\phi_j)(x) - \Pi_h \tilde{F}(\phi_j)(x) \right|^2 dx \leq C_{\Pi_h}^2 \int_K |F(\phi_j)(x) - \tilde{F}(\phi_j)(x)|^2 dx \leq C_{\Pi_h}^2 \left(C_{\Pi_h} C_3 \frac{h}{t} \text{diam}(\mathcal{O}) (4\pi t)^{-\frac{p}{2}} \right)^2$$

De plus, $|\text{supp}(\phi_j)| \leq \beta \mu h^p$ donc on obtient finalement :

$$\begin{aligned} \left(\int_{V_h} \left| \Pi_h F(\phi_j)(x) - \Pi_h \tilde{F}(\phi_j)(x) \right|^2 dx \right)^{\frac{1}{2}} &\leq \left(\sum_{\substack{K \in \mathcal{T}_h \\ \text{supp} \phi_j \cap K \neq \emptyset}} C_{\Pi_h}^4 C_3^2 \beta \frac{h^{p+2}}{t^2} \text{diam}(\mathcal{O})^2 (4\pi t)^{-p} \right)^{\frac{1}{2}} \\ &\leq C \frac{h^{p+1}}{t} (4\pi t)^{-\frac{p}{2}}, \end{aligned}$$

avec $C = C_3 C_{\Pi_h}^2 \beta \sqrt{\mu} \text{diam}(\mathcal{O})$. □

Remarque 2.43. La condensation de masse n'a pas modifié l'hypothèse (2.31) de la borne minimale du paramètre de la chaleur t . En effet, t doit toujours satisfaire la condition suivante :

$$h^{3p+2} < t^2. \quad (2.47)$$

Le résidu ψ du résultat principal (section 2.1.1) est fonction de la distance δ , du pas de discrétisation h et du paramètre du noyau Gaussien t qui doit satisfaire l'équation (2.20). Pour finir l'interprétation du spectral clustering non normalisé, il faut revenir aux vecteurs propres de la matrice affinité A , autrement dit, "supprimer" la matrice $\widetilde{M}$ dans l'équation (2.45). Or $\widetilde{M}$ est une matrice diagonale mais si on avait $\widetilde{M} = a\mathbb{I}_n$ alors le résultat serait immédiat. Le i^{eme} élément de $\widetilde{M}$ est donné par : $\widetilde{M}_{i,i} = \text{supp}(\phi_i)/(p+1)$. Demander à avoir $\widetilde{M} \approx a\mathbb{I}_N$ revient à avoir $|\text{supp}(\phi_i)| \approx |\text{supp}(\phi_j)|$, pour tout (i, j) . Ce qui revient encore à demander plus de régularité sur le maillage. En notant Φ la moyenne des tailles sur tous les supports de ϕ_j , cela revient finalement à avoir la relation suivante pour tout j

$$\frac{|\text{supp}(\phi_j)| - \Phi}{\Phi} \rightarrow 0 \text{ quand } h \rightarrow 0.$$

Avec l'hypothèse de régularité définie dans la proposition 2.42, on avait, pour tout $j : |\text{supp}(\phi_j)| \approx C_j h^p$. Ces deux hypothèses ensemble, le support de ϕ_j vérifie donc :

$$|\text{supp}(\phi_j)| = C_0 h^p + O(h^{p+1}). \quad (2.48)$$

Avec cette dernière hypothèse de régularité du maillage, la matrice $\widetilde{M}$ est remplacée par une matrice homogène à une matrice identité et la propriété de clustering peut être énoncée avec uniquement la matrice affinité A comme suit.

Proposition 2.44. *Supposons que les supports des fonctions ϕ_j pour tout j vérifient l'équation (2.48). Soient les $W_{n,i}$ définis par la proposition 2.36. Alors pour $t > 0$ on a :*

$$(4\pi t)^{-\frac{p}{2}} C_0 h^p (A + \mathbb{I}_N) W_{n,i} = e^{-\lambda_{n,i} t} W_{n,i} + O\left(t, h^2, \frac{h^{p+1}}{t}, \frac{h^{3p+2}}{t^2}\right) \quad (2.49)$$

où quand $(h, t) \rightarrow 0$ avec $h^{3p+2}/t^2 \rightarrow 0$, $h^{p+1}/t \rightarrow 0$ et $\delta \rightarrow 0$.

Démonstration. D'après l'hypothèse (2.48) et la définition (2.43) des éléments de $\widetilde{M}$, le résultat est immédiat. $\square$

Avec une condensation de masse et une hypothèse d'homogénéité de maillage, la propriété de spectral clustering est maintenant ramenée sur la matrice affinité A . On a donc montré que sous plusieurs hypothèses sur le paramètre t (ou σ) et une isotropie de la distribution, les vecteurs propres issus de la matrice affinité sont proches des ceux de l'opérateur de la chaleur discrétisé avec conditions de Dirichlet. En effet, lorsque $t/h^2 \rightarrow 0$, les vecteurs propres associés à des valeurs propres proches de 0 correspondent aux fonctions propres discrétisées. Ainsi, les données projetées dans l'espace spectral sont concentrées par classes et séparées.

Les images de la figure 2.9 illustrent le produit entre la matrice affinité A et les fonctions propres discrétisées $V_{1,i}$, $i \in \{1, 2\}$. De la même façon, la figure suivante représente les vecteurs propres de la matrice A qui sont le plus semblable respectivement de la fonction propre discrétisée $V_{1,1}$ et $V_{1,2}$ pour la valeur $t = t_h$ où t_h représente l'heuristique (1.2) développée dans le chapitre 1 tel que

$$t_h = \frac{\sigma_h}{4} = \frac{D_{\max}}{8N^{\frac{1}{p}}}.$$

Pour terminer, on peut alors donner le résultat de clustering sur la matrice A .

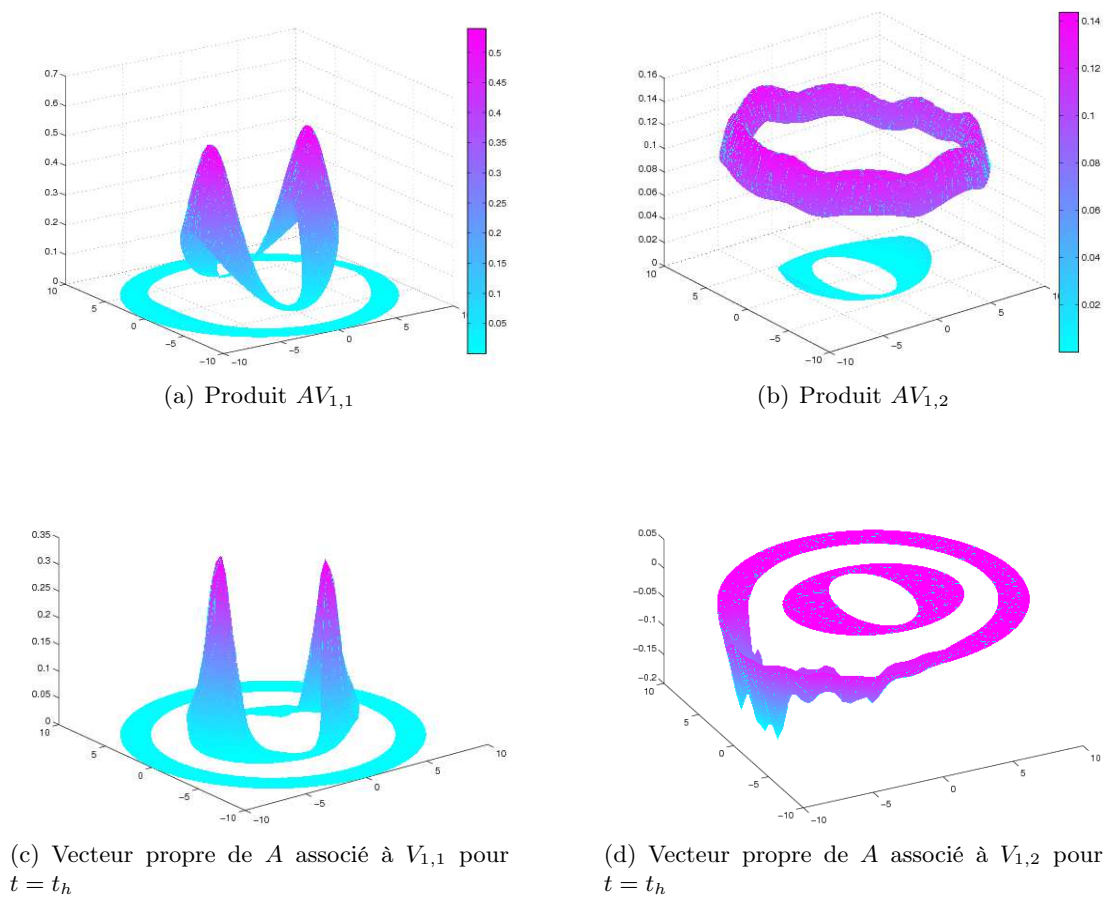


FIGURE 2.9 – Exemple des donuts

Théorème 2.45. (*Propriété de Clustering*) En reprenant les hypothèses de la proposition 2.44, il existe une famille de vecteurs, notés $Z_{n,i}$ pour $i \in \{1, \dots, k\}$ et $n > 0$ avec $\|Z_{n,i}\|_2 = 1$, telle que pour tout $x_r \in S$ et pour tout $i \in \{1, \dots, k\}$ on obtient le résultat suivant :

$$\left[\begin{array}{l} \exists \alpha > 0, \exists n > 0, \forall t > 0 \text{ avec } t, h^2/t \text{ et } h^{3p+1}/t^2 \text{ assez petits,} \\ Z_{n,i} = \operatorname{argmax}_{\{Z_{m,j}, m \in \mathbb{N}, j \in \{1, \dots, k\}\}} |((A + \mathbb{I}_N)_{\cdot r} | Z_{m,j})_2| \\ \text{et } |((A + \mathbb{I}_N)_{\cdot r} | Z_{n,i})_2| > \alpha \end{array} \right] \iff x_r \in \mathcal{O}_i, \quad (2.50)$$

où $(A + \mathbb{I}_N)_{\cdot r}$ correspond à la $r^{\text{ième}}$ colonne de la matrice $A + \mathbb{I}_N$.

Démonstration. On considère les $W_{n,i}$ définis dans la proposition 2.44 et on pose $Z_{n,i} = \frac{W_{n,i}}{\|W_{n,i}\|_2}$. D'après le théorème 2.38, les $W_{n,i}$ satisfont l'argument maximal du théorème. Par contre, ce résultat est énoncé avec la norme définie avec la matrice de masse M . Il faut maintenant comparer la norme avec M et la norme euclidienne. Par définition et en utilisant l'équation (2.44), on obtient pour $w \in L^2(V_h)$ et W les coordonnées de w dans la base, avec les hypothèses de régularité du maillage (2.48) : $(W|W)_M = \int_{\mathcal{O}} w(x)^2 dx = (W|W)_{\widetilde{M}} + O(h) = \sum_i W_i \widetilde{M}_{ii} W_i + O(h) = \sum W_i^2 (C_0 h^p + O(h^{p+1})) + O(h)$ donc $\|W\|_M = \|W\|_2 \sqrt{C_0 h^p} + O(\sqrt{h})$. Donc pour les $W_{n,i}$, $\|W_{n,i}\|_2 = \frac{1}{\sqrt{C_0 h^p}} + O(\frac{1}{\sqrt{h}})$ puis avec $W_{n,i} = Z_{n,i} \|W_{n,i}\|_2$ dans (2.49), il vient :

$$\begin{aligned} C_0 h^p (4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) Z_{n,i} (1 + O(h^{(p-1)/2})) &= e^{-\lambda_{n,i} t} Z_{n,i} (1 + O(h^{(p-1)/2})) \\ &+ O\left(th^{p/2}, h^{2+p/2}, \frac{h^{3p/2+1}}{t}, \frac{h^{3p+2+p/2}}{t^2}\right) \end{aligned} \quad (2.51)$$

Comme les $Z_{n,i}$ sont normés, $\|AZ_{n,i}\|_2 \leq \|A\|_\infty$ avec $\|A\|_\infty = \max_{1 \leq j \leq n} \sum_{i=1}^N A_{ij} \leq 2N$ par définition de l'affinité (1.1). Or N représente le nombre de points dans le volume $\mathcal{O}$ donc avec l'hypothèse de régularité du maillage (2.44), il est égal au rapport entre le volume total et le volume de chaque élément fini soit : $N \approx 1/h^p$. Donc $C_0 h^p AZ_{n,i} (1 + O(h^{(p-1)/2})) = C_0 h^p AZ_{n,i} + O(h^{(p-1)/2})$. Donc

$$C_0 h^p (4\pi t)^{-\frac{p}{2}} (A + \mathbb{I}_N) Z_{n,i} = e^{-\lambda_{n,i} t} Z_{n,i} + O\left(th^{p/2}, h^{(p-1)/2}, \frac{h^{3p/2+1}}{t}, \frac{h^{3p+2+p/2}}{t^2}\right).$$

□

Ce dernier théorème montre donc le fonctionnement de la classification spectrale sans étape de normalisation. Or ces résultats dépendent du pas de discrétisation h et du paramètre de l'affinité t . Les rôles de h et t sont donc étudiés dans la section suivante à travers d'autres expérimentations numériques.

2.5 Discussions

Dans cette section, nous présentons des expérimentations numériques sur des exemples présentant des cas de séparations linéaire et non linéaire. Ainsi les diverses étapes développées dans les chapitres seront représentées. Le choix du paramètre de la chaleur sera évoqué et comparé aux hypothèses issues de la théorie et un exemple de cas limite entre le continu et le discret sera présenté. Ensuite, le passage du discret au continu sera évoqué via la triangulation de Delaunay. L'étape de normalisation n'étant pas prise en compte dans tout ce chapitre, elle fera l'objet d'une brève étude et de discussions. Enfin, nous reviendrons sur des cas limites de validité de l'étude.

2.5.1 Expérimentations numériques

Après l'exemple des donuts, deux cas tests sont maintenant considérés : le premier dont les classes sont séparables par hyperplans et le second dont les classes ne sont pas séparables par hyperplans. L'influence du paramètre t est analysée sur les différences entre les fonctions propres discrétisées de l'opérateur $S_D^{\mathcal{O}}(t)$ et les vecteurs propres de la matrice $(A + \mathbb{I}_N)M$. Les figures 2.10 et respectivement 2.13 représentent, pour $i \in \{1, 2, 3\}$:

- (a) : l'ensemble des données S et le résultat du clustering spectral pour $t = t_h$,
- (b)-(c)-(d) : les fonctions propres $V_{1,i}$ associées à la première valeur propre sur chaque composante connexe,
- (e)-(f)-(g) : le produit de la matrice $(A + \mathbb{I}_N)M$ avec ces fonctions discrétisées $V_{1,i}$,
- (h)-(i)-(j) : le produit de la matrice affinité A avec $V_{1,i}$.

Ces figures montrent que la localité du support des fonctions propres est préservée avec la discrétisation des opérateurs $(A + \mathbb{I}_N)M$ et A . Pour montrer la propriété d'invariance géométrique, nous comparons la corrélation entre les fonctions propres $V_{n,i}$ et les vecteurs propres de $(A + \mathbb{I}_N)M$ qui maximisent le coefficient de projection avec $V_{1,i}$. En d'autres termes, soit $V_{1,i} = \Pi_h \widetilde{v}_{1,i} \in V_h$ et X_l le vecteur propre de $(A + \mathbb{I}_N)M$ tel que : $l = \arg \max_j (V_{1,i} | X_j)$, les corrélations ω et τ respectivement entre les vecteurs X_l et $V_{1,i}$ et entre les vecteurs Y_l et $V_{1,i}$ correspondent alors pour tout $i \in \{1, 2, 3\}$ à :

$$\omega = \frac{|(V_{1,i} | X_l)|}{\|V_{1,i}\|_2 \|X_l\|_2} \text{ et } \tau = \frac{|(V_{1,i} | Y_l)|}{\|V_{1,i}\|_2 \|Y_l\|_2}.$$

Les différences en norme, notées α et β , respectivement entre les vecteurs X_l et $V_{n,i}$ et entre les vecteurs Y_l et $V_{n,i}$ sont définies par, pour tout $i \in \{1, 2, 3\}$:

$$\alpha = \|X_l - V_{1,i}\|_2 \text{ et } \beta = \|Y_l - V_{1,i}\|_2$$

. Comme les matrices $(A + \mathbb{I}_N)M$ et A dépendent du paramètre t , les figures 2.11 et 2.14 représentent, pour $i \in \{1, 2, 3\}$:

- (a)-(c) : l'évolution de la différence en norme, α (respectivement β), entre $V_{1,i}$ avec les vecteurs propres de la matrice $(A + \mathbb{I}_N)M$ (respectivement A) en fonction de t ,
- (b)-(d) : l'évolution de la corrélation, ω (respectivement τ), entre $V_{1,i}$ avec les vecteurs propres de la matrice $(A + \mathbb{I}_N)M$ (respectivement A) en fonction de t .

Lorsque la corrélation τ et la différence en norme β sont étudiées pour la matrice affinité A , des pics apparaissent ponctuellement en fonction de la valeur de t . Ils correspondent à un changement de vecteur propre X_l de la matrice A maximisant le coefficient de projection $(V_{1,i} | X_j)$ pour cette valeur de t .

Le coefficient de projection ω , proche de 1 pour une valeur optimale de t , souligne le fait que les fonctions discrétisées de l'opérateur Laplacien avec conditions de Dirichlet sur un ouvert plus grand Ω incluant une composante connexe $\mathcal{O}_i$ sont proches des vecteurs propres de l'opérateur de la chaleur discrétisé sans condition aux frontières. Ce comportement reste sous la contrainte d'un paramètre t qui requiert d'être choisi dans une plage appropriée. Le même comportement est observé respectivement avec la différence en norme.

L'influence de la discrétisation du maillage est aussi étudiée sur la figure en superposant, pour $h_{max} \in \{0.2, 0.4, 0.6\}$, les courbes de corrélation entre les vecteurs $V_{1,i}$ et Y_l pour chaque composante connexe $i \in \{1, 2, 3\}$. Plus le maillage est raffiné, plus la corrélation est forte. En effet, plus on raffine la maillage, plus on s'approche de la version continue du clustering spectral et les vecteurs propres en sont d'autant plus localisés.

Sur les figures 2.11 et 2.14, la ligne verticale en pointillés indique la valeur du paramètre t de l'heuristique 1.2 développée dans le chapitre 1 :

$$t_h = \frac{\sigma_h}{4} = \frac{D_{\max}}{8N^{\frac{1}{p}}},$$

avec $D_{\max} = \max_{1 \leq i, j \leq N} \|x_i - x_j\|_2$.

Cette valeur est proche de l'optimum pour les exemples 2.10 et 2.13 (a) : cette valeur critique donne une bonne estimation de la distance δ de séparation entre les classes. De plus, les figures 2.10 (a) et 2.13 (a) représentent le résultat du clustering spectral pour cette valeur confirment cette observation.

2.5.2 Choix du paramètre gaussien

Dans une première partie, nous comparons l'heuristique définie dans le chapitre 1 avec les hypothèses sur le paramètre de la chaleur issues de l'interprétation via l'équation de la chaleur et la théorie des éléments finis. Puis, la différence entre le noyau de la chaleur K_h et de la matrice d'affinité gaussienne sera étudiée en fonction du paramètre de la chaleur. Enfin, un cas limite d'étude entre la version continue et la version discrète sera présenté.

Lien avec l'heuristique

Dans l'approximation intérieure de l'ouvert Ω_i par un ouvert $\mathcal{O}_i$, $\delta_i = \inf_{y \in \mathcal{O}_i} \delta_i(y)$ représente la distance de $\mathcal{O}_i$ à Ω_i où $\delta_i(y)$ est la distance d'un point $y \in \mathcal{O}_i$ à $\partial\Omega_i$, la frontière de Ω_i . Avec l'interprétation par les éléments finis, tous les points x_i de Σ appartiennent à l'intérieur de l'ouvert $\mathring{\mathcal{O}}$ et l'ouvert $\mathcal{O}$ est la réunion de tous les éléments finis K de la triangulation $\mathcal{T}_h$, comme le montre la figure 2.16 :

$$\mathcal{O} = \bigcup_{K \in \mathcal{T}_h} K.$$

La distance de séparation $\delta = \inf_i \delta_i$ peut alors être exprimée en fonction du pas de discrétisation h . D'après la définition de la séparation entre Ω et $\mathcal{O}$ et la représentation discrète, la distance δ est homogène à h :

$$\delta \propto h. \quad (2.52)$$

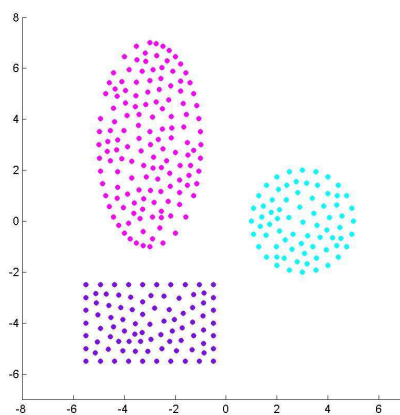
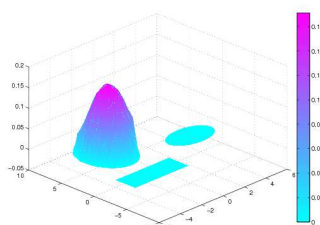
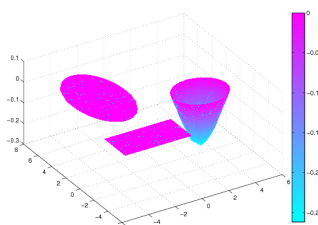
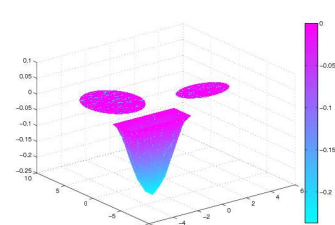
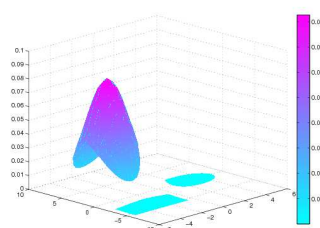
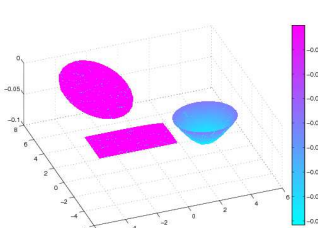
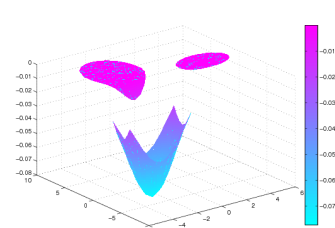
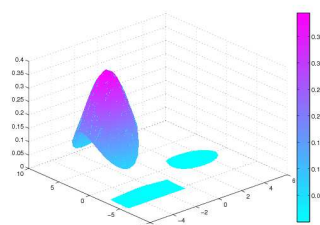
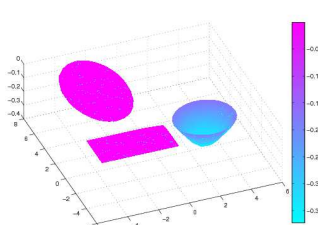
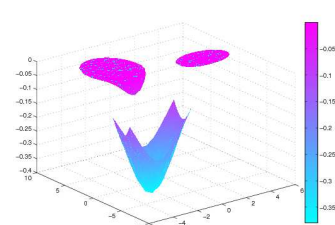
De plus, les résultats démontrés précédemment ont permis de définir deux conditions (2.20) et (2.47) sur la paramètre t :

$$h^{3p+2} \ll t^2 \ll \delta^4.$$

Autrement dit, avec l'équation (2.52), le paramètre t est fonction du pas h et est borné par :

$$h^{3p+2} \ll t^2 \ll h^4.$$

Donc t est une fraction de la distance h^2 .

(a) Exemple1 : $N = 302$ points(b) $V_{1,1}$ (c) $V_{1,2}$ (d) $V_{1,3}$ (e) Produit $(A + \mathbb{I}_N)MV_{1,1}$ (f) Produit $(A + \mathbb{I}_N)MV_{1,2}$ (g) Produit $(A + \mathbb{I}_N)MV_{1,3}$ (h) Produit $AV_{1,1}$ (i) Produit $AV_{1,2}$ (j) Produit $AV_{1,3}$ FIGURE 2.10 – Exemple 1 : Fonction propre discrétisée $V_{1,i}$ et produits $(A + \mathbb{I}_N)MV_{1,i}$, $AV_{1,i}$ pour $i \in \{1, 2, 3\}$

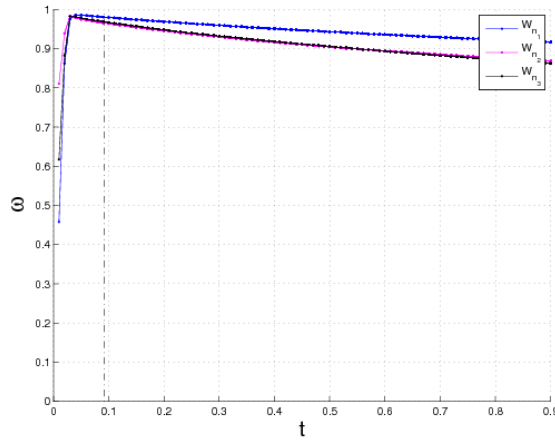
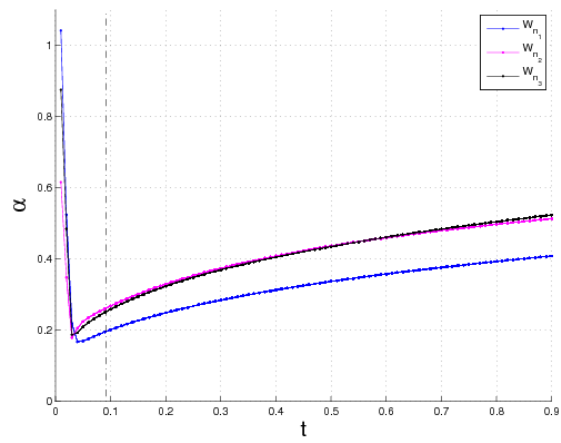
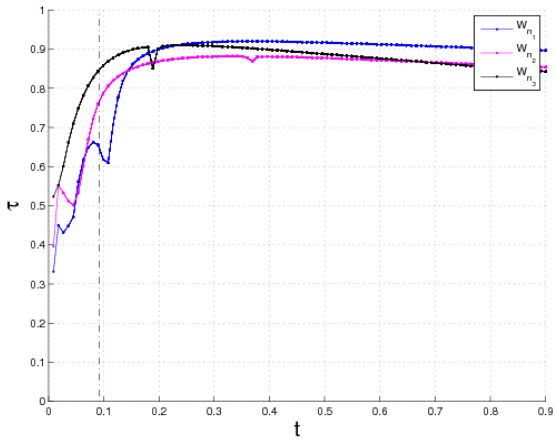
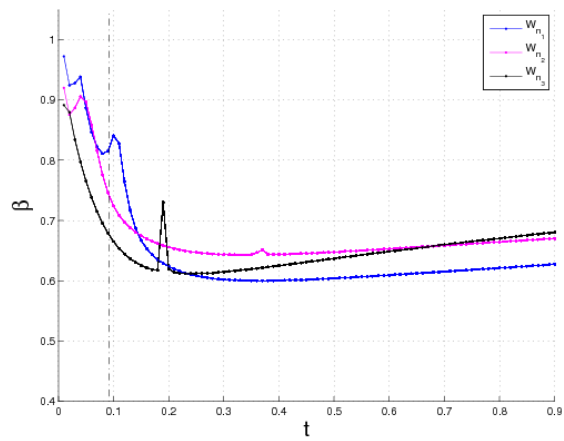
(a) Corrélation ω en fonction de t (b) Norme α en fonction de t (c) Corrélation τ en fonction de t (d) Norme β en fonction de t

FIGURE 2.11 – Exemple 1 : Corrélation et différence en norme entre $V_{1,i}$ et, respectivement, les vecteurs propres X_l de $(A + \mathbb{I}_N)M$ et Y_l de A , pour $i \in \{1, 2, 3\}$

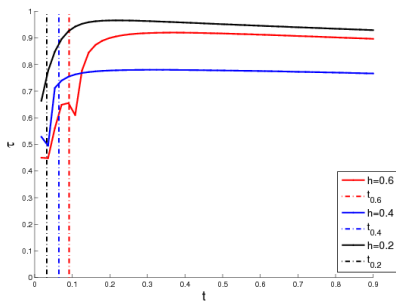
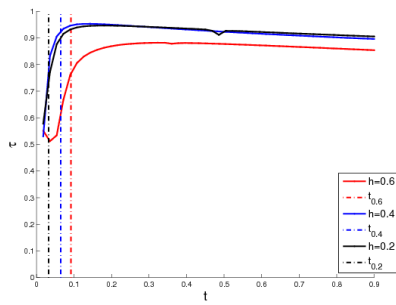
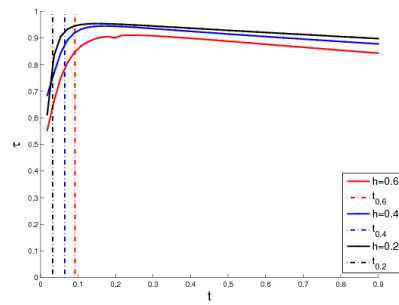
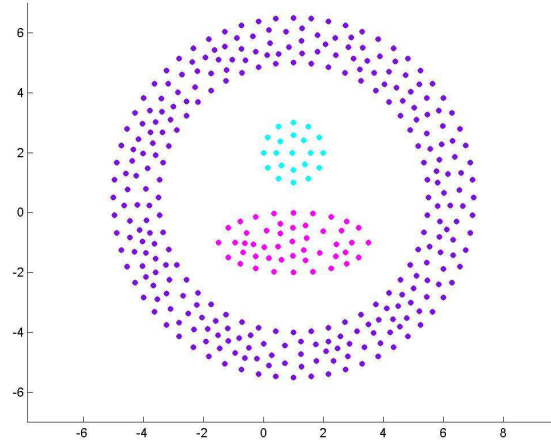
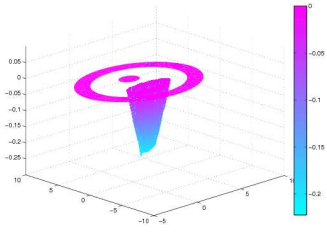
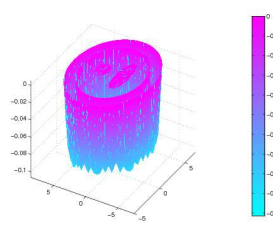
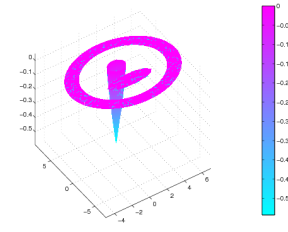
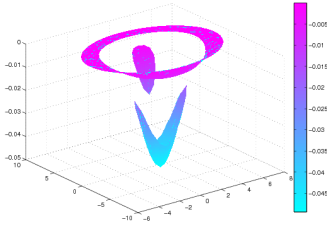
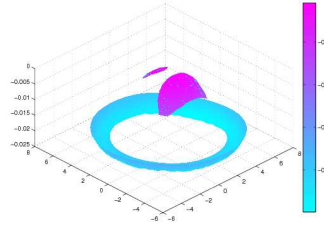
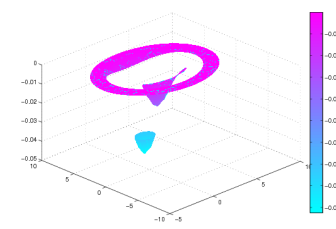
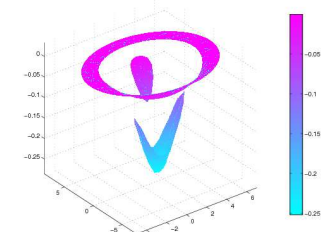
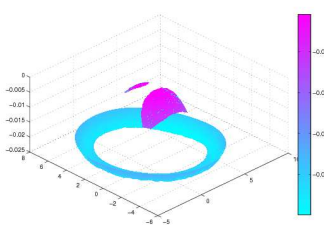
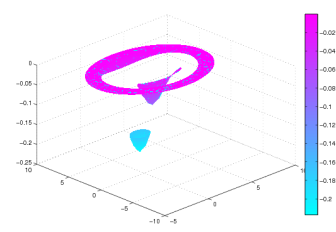
(a) Corrélation τ en fonction de t et du pas h pour $V_{1,1}$ (b) Corrélation τ en fonction de t et du pas h pour $V_{1,2}$ (c) Corrélation τ en fonction de t et du pas h pour $V_{1,3}$

FIGURE 2.12 – Exemple 1 : Influence de la discrétisation sur la corrélation entre $V_{1,i}$ et les vecteurs propres Y_l de A , pour $i \in \{1, 2, 3\}$

(a) Exemple 2 : $N = 368$ points(b) $V_{1,1}$ (c) $V_{1,2}$ (d) $V_{1,3}$ (e) Produit $(A + \mathbb{I}_N)MV_{1,1}$ (f) Produit $(A + \mathbb{I}_N)MV_{1,2}$ (g) Produit $(A + \mathbb{I}_N)MV_{1,3}$ (h) Produit $AV_{1,1}$ (i) Produit $AV_{1,2}$ (j) Produit $AV_{1,3}$ FIGURE 2.13 – Exemple 2 : Fonction propre discrétisée $V_{1,i}$ et produits $(A + \mathbb{I}_N)MV_{1,i}$, $AV_{1,i}$ pour $i \in \{1, 2, 3\}$

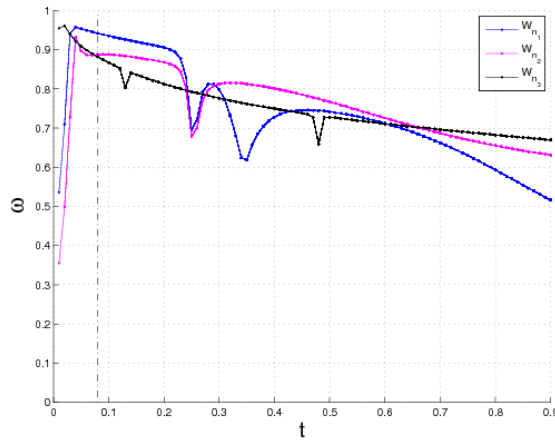
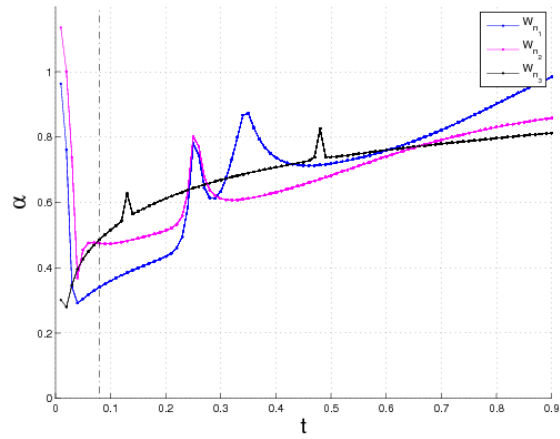
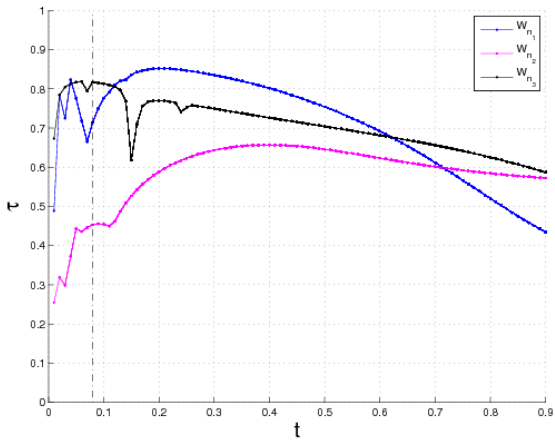
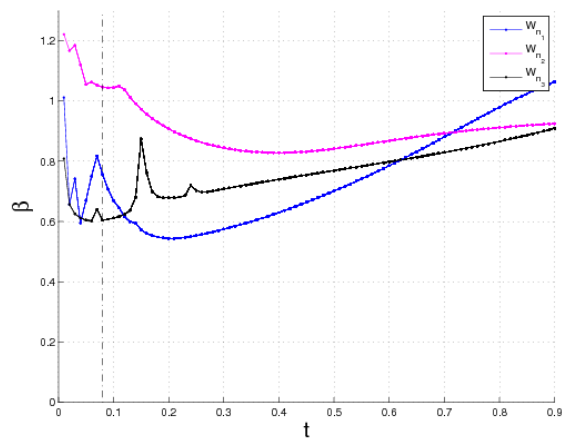
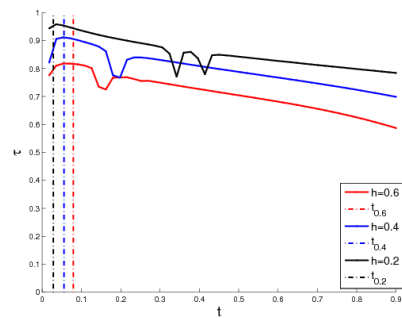
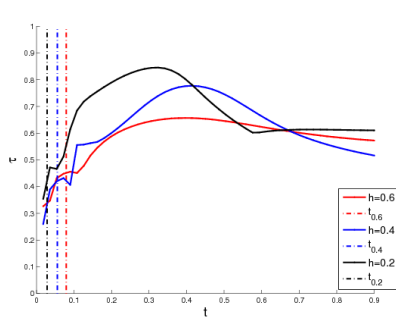
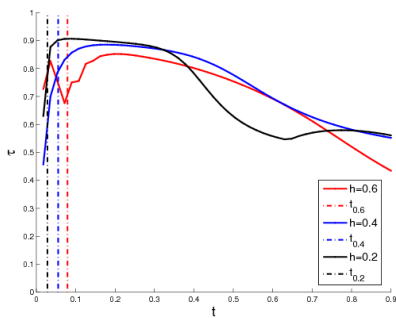
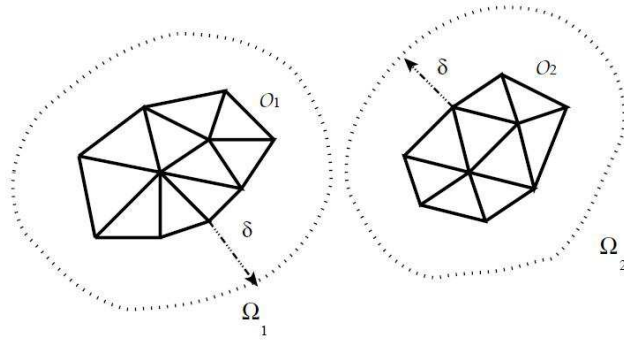
(a) Corrélation ω en fonction de t (b) Norme α en fonction de t (c) Corrélation τ en fonction de t (d) Norme β en fonction de t

FIGURE 2.14 – Exemple 2 : Corrélation et différence en norme entre $V_{1,i}$ et respectivement les vecteurs propres X_l de $(A + \mathbb{I}_N)M$ et Y_l de A , pour $i \in \{1, 2, 3\}$



(a) Corrélation τ en fonction de t et du pas h pour $V_{1,1}$ (b) Corrélation τ en fonction de t et du pas h pour $V_{1,2}$ (c) Corrélation τ en fonction de t et du pas h pour $V_{1,3}$

FIGURE 2.15 – Exemple 2 : Influence de la discrétisation sur la corrélation entre $V_{1,i}$ et les vecteurs propres Y_l de A , pour $i \in \{1, 2, 3\}$

FIGURE 2.16 – Estimation de la distance δ

Revenons à l'estimation heuristique (1.2) du paramètre σ développée dans le chapitre 1. Le terme σ est une fonction de la distance référence σ_0 définie par :

$$\sigma_0 = \frac{D_{max}}{N^{\frac{1}{p}}}.$$

Cette distance représente la distance de séparation dans un cas où la distribution des points est équirépartie c'est-à-dire lorsque les points d'une même composante connexe sont équidistants. Dans ce contexte de maillage et d'éléments finis, cette définition correspond à une estimation du pas de discrétisation moyen. Donc la définition de l'heuristique, c'est-à-dire la distance référence $\frac{\sigma}{2}$, satisfait les conditions (2.20) et (2.47). L'heuristique (1.2) vérifie alors la condition (2.20) pour lesquelles les vecteurs propres seront très localisés sur les composantes connexes.

Différences entre le noyau de la chaleur K_H et la matrice affinité A

Quelques différences apparaissent dans les définitions respectives de la matrice affinité (1.1) et du noyau de la chaleur. Rappelons la formule reliant le noyau de la chaleur K_H et la matrice affinité définie dans l'algorithme 1 (1.1) :

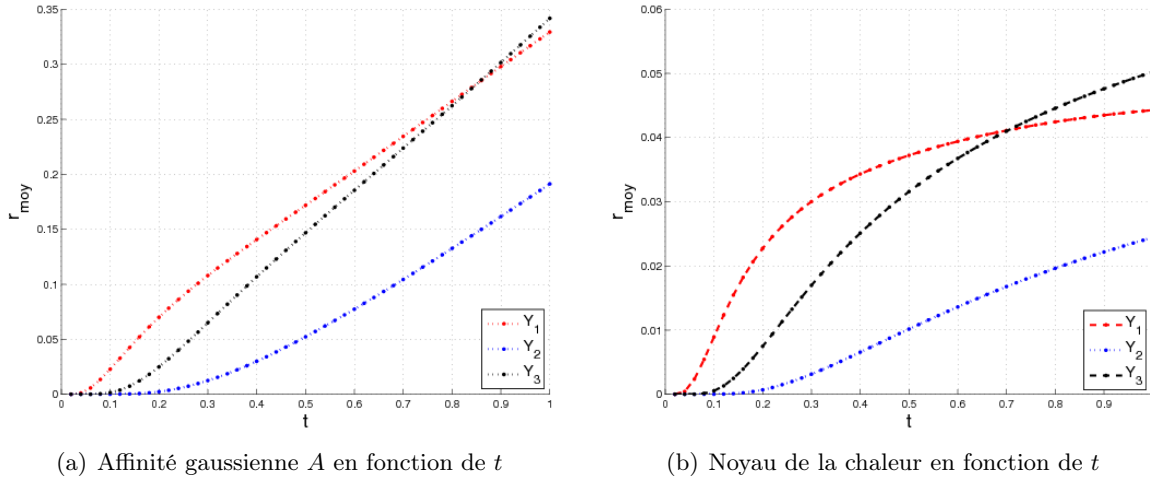
$$K_H \left(\frac{\sigma^2}{2}, x_i, x_j \right) = (2\pi\sigma^2)^{\frac{p}{2}} (A + \mathbb{I}_N)_{ij}.$$

Donc l'affinité Gaussienne utilisée pour l'algorithme 1 n'utilise pas le coefficient de normalisation $(2\pi\sigma^2)^{-p/2}$ et la diagonale de la matrice A est nulle.

La mesure de qualité basée sur les ratios de normes de Frobenius présentée dans le chapitre 1 est utilisée pour comparer les différences entre les noyaux. La moyenne des ratios de normes de Frobenius introduits dans la section 1.3.1 sur chaque cluster est étudiée en fonction du paramètre t .

$$r_{moy} = \frac{1}{3} \sum_{i=1}^3 \sum_{j=1, j \neq i}^3 \frac{\|L_{ij}\|_F}{\|L_{ii}\|_F}.$$

Si r_{moy} est proche de 0, l'affinité inter-classes est alors négligeable par rapport à l'affinité intra-classe. Ce cas s'approche d'un cas idéal de partitionnement des données où l'on considère des clusters bien séparés. Ce critère r_{moy} est donc observé respectivement pour le noyau de la chaleur et l'affinité de

FIGURE 2.17 – Exemple 2 : Mesure de qualité de clustering en fonction de t

clustering spectral et représenté figure 2.17. L'allure des courbes de la figure 2.17 (a) pour l'affinité gaussienne de l'algorithme est plus linéaire que celles de la figure 2.17(b). De plus, l'amplitude de la variation est différente d'un rapport de 10 en raison de l'absence de coefficient de normalisation $(4\pi t)^{-\frac{p}{2}}$ avec $t \leq 1$. Cependant, le support sur lequel r_{moy} est nul est identique dans les deux cas et donc le comportement par classe est identique mais avec une plus forte variation avec la définition de l'affinité gaussienne. Ceci tend à valider le choix de r_{moy} comme critère d'un bon clustering.

Cas limite

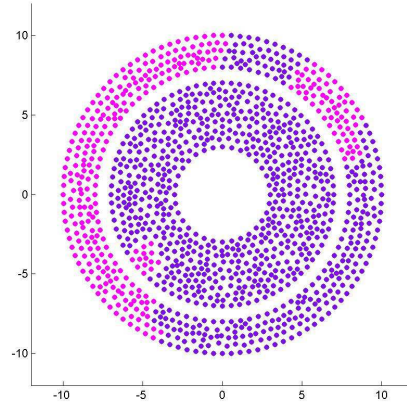
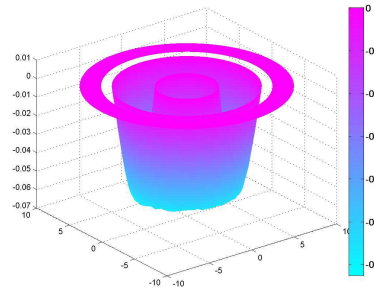
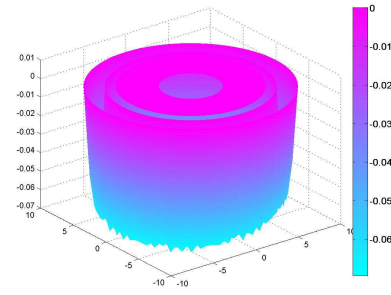
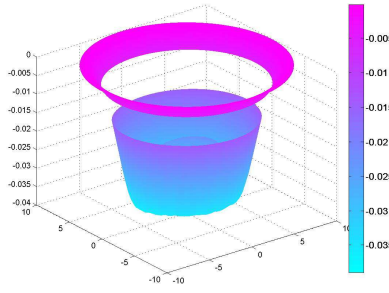
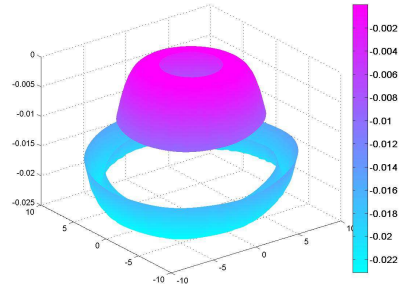
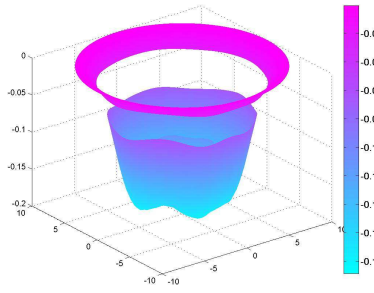
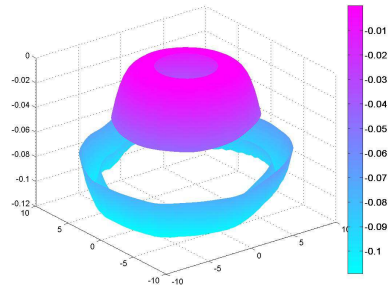
Cependant, des limites peuvent apparaître lors du passage du continu au discret. L'exemple représenté par la figure 2.18 où l'on considère des couronnes larges concentriques montre que le résidu de l'équation (2.49) dépend de la discrétisation et de la séparation des classes. En effet, sur les graphes de la figure 2.19 la valeur $t = t_h$ coïncide avec l'optimum dans le continu mais est très éloigné de l'optimum lors du passage au discret, ce qui est représenté sur la figure 2.18 où le résultat du spectral clustering pour $t = t_h$ est tracé. Ainsi la séparation entre les classes est trop faible pour avoir un résidu négligeable et conserver la propriété géométrique sur les vecteurs propres.

2.5.3 Passage du discret au continu : Triangulation de Delaunay

Cette interprétation de la classification spectrale requiert de revenir à une formulation continue en vue de dégager une propriété de classification pour enfin revenir à un ensemble fini de points via un maillage par éléments finis. Cependant, cette représentation par éléments finis, à savoir définir un maillage pour tout ensemble fini de points, est-elle réaliste ?

Les hypothèses d'isotropie de distribution des données et la régularité, la quasi-uniformité du maillage permettent de définir un maillage pour un ensemble quelconque de points via la triangulation de Delaunay. En effet, les triangulations de Delaunay [47] présentent des propriétés de régularité utiles pour l'approximation par éléments finis.

Définition 2.46. Pour un ensemble de points S de l'espace euclidien $\mathbb{R}^p$ constituant les sommets

(a) Exemple 3 : $N = 368$ points(b) $V_{1,1}$ (c) $V_{1,2}$ (d) Produit $(A + \mathbb{I}_N)MV_{1,1}$ (e) Produit $(A + \mathbb{I}_N)MV_{1,2}$ (f) Produit $AV_{1,1}$ (g) Produit $AV_{1,2}$ FIGURE 2.18 – Exemple 3 : Fonction propre discrétisée $V_{1,i}$ et produits $(A + \mathbb{I}_N)MV_{1,i}$, $AV_{1,i}$ pour $i \in \{1, 2\}$

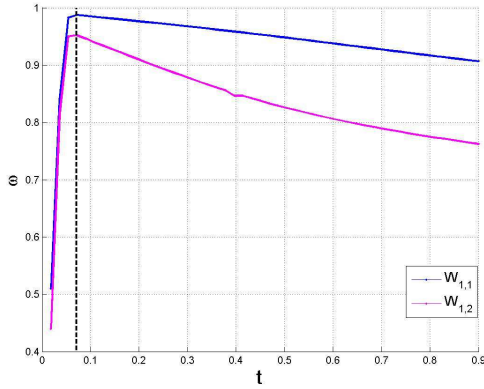
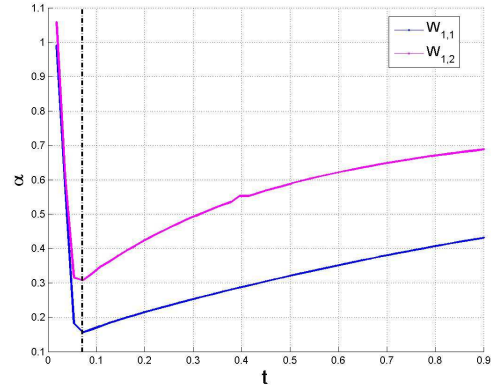
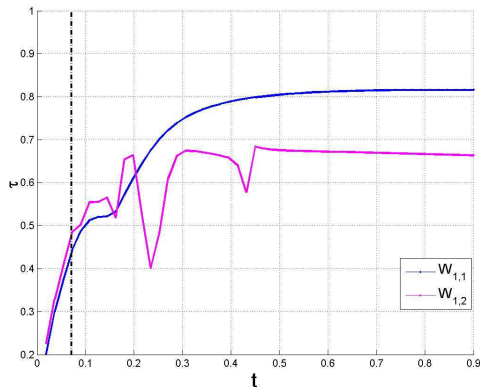
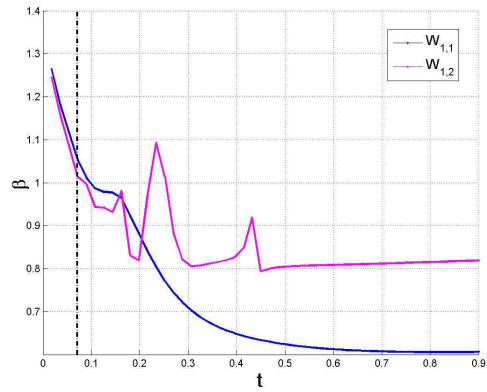
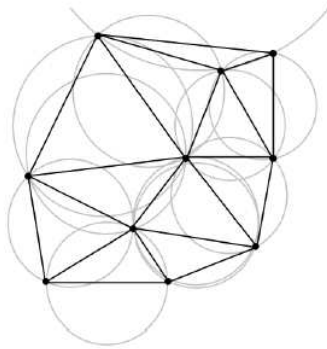
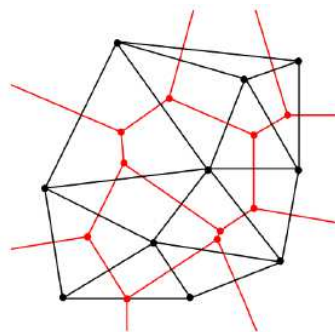
(a) Corrélation ω en fonction de t (b) Norme α en fonction de t (c) Corrélation τ en fonction de t (d) Norme β en fonction de t

FIGURE 2.19 – Exemple 3 : Corrélation et différence en norme entre $V_{1,i}$ et respectivement les vecteurs propres X_l de $(A + \mathbb{I}_N)M$ et Y_l de A , pour $i \in \{1, 2\}$



(a) Triangulation de Delaunay et cercles circonscrits



(b) Triangulation de Delaunay et dual de Voronoï

FIGURE 2.20 – Principe de la triangulation de Delaunay

d'une triangulation $\mathcal{T}_h$, la triangulation $\mathcal{T}_h$ est de Delaunay si aucun sommet de S ne se trouve dans l'intérieur du disque (ou sphère) circonscrit d'un simplexe K de $\mathcal{T}_h$.

Cette propriété de "la sphère ouverte circonscrite vide" est représentée par la figure 2.20 (a). Elle permet de contrôler le rapport entre le rayon de la sphère inscrite au rayon de la sphère circonscrite pour chaque maille. Ainsi l'hypothèse de régularité du maillage (2.25) est vérifiée. Les propriétés de la triangulation sont les suivantes :

1. l'union des simplexes $\bigcup_{K \in \mathcal{T}_h} K$ constitue l'enveloppe convexe des points de S ;
2. S_h est le graphe dual du diagramme de Voronoï associé à S (figure 2.20 (b)).

Il a été démontré [47] que la construction de la triangulation de Delaunay d'un ensemble de points de dimension p est un problème équivalent à celui de la construction de l'enveloppe convexe d'un ensemble de points en dimension $p + 1$. Comme l'enveloppe convexe existe, la triangulation de Delaunay existe aussi.

Proposition 2.47. (*Unicité de la triangulation de Delaunay*) Une triangulation de Delaunay $\mathcal{T}_h$ est unique si les points de S sont en position générale, c'est-à-dire s'il n'y a pas $p + 1$ points dans le même hyperplan ni $p + 2$ points sur la même hypersphère.

La figure 2.21 représente les triangulations de Delaunay appliquées aux noeuds d'un maillage (e)-(f) ou à un ensemble quelconque de points comme les exemples précédemment présentés (figure 2.21 (a) à (d)).

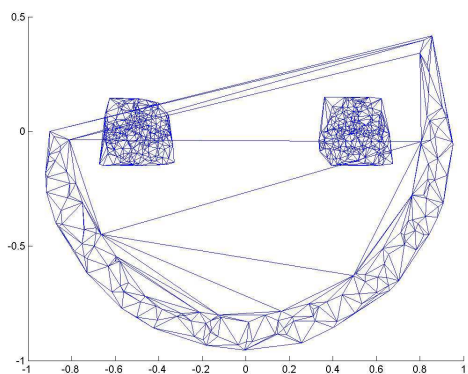
2.5.4 Etape de normalisation

Dans tout ce chapitre, nous avons considéré le spectral clustering sans son étape de normalisation. A l'origine, l'algorithme 1 de Ng, Jordan et Weiss [84] normalise la matrice affinité gaussienne avant d'en extraire les plus grands vecteurs propres cf l'algorithme 3. Cependant quelles modifications cette étape de normalisation peut-elle engendrer par rapport à la théorie précédente ?

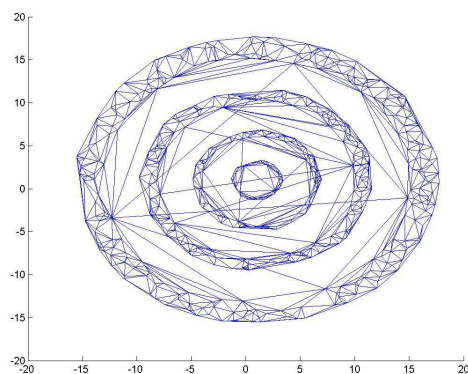
Les articles relatifs aux marches aléatoires [79, 114] interprètent la méthode de spectral clustering comme un problème homogène avec condition de Neumann. Cette interprétation est expliquée par le fait que la marche aléatoire peut seulement sauter d'un point à un autre sans être absorbé et la conservation de la matrice de transition en probabilités aboutit à considérer des conditions de Neumann. Cette étape modifierait donc le comportement de la méthode qui considère un problème aux valeurs propres avec des conditions de Dirichlet (ou conditions d'absorption) et amènerait donc à considérer les fonctions propres d'un problème d'équation de la chaleur avec des conditions de Neumann suivant :

$$(\mathcal{P}_\Omega^N) \begin{cases} \partial_t u - \Delta u = 0 \text{ sur } \mathbb{R}^+ \times \Omega, \\ \partial_n u = 0, \text{ sur } \mathbb{R}^+ \times \partial\Omega, \\ u(t = 0) = f, \text{ sur } \Omega. \end{cases}$$

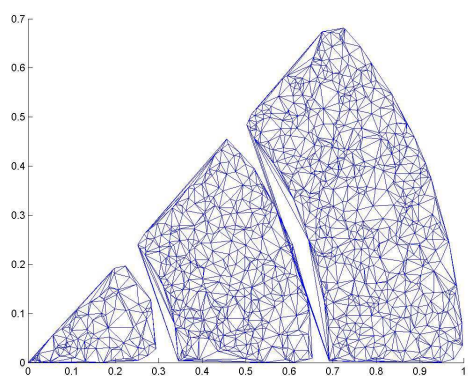
L'avantage de cette nouvelle approche est que l'opérateur de Laplace avec conditions de Neumann reste compact autoadjoint positif et, de plus, ses valeurs propres strictement positives seront toutes inférieures ou égales à 1. En outre, les fonctions égales à 1 sur Ω_i et 0 sur $\Omega \setminus \Omega_i$, pour $i \in \{1, \dots, k\}$ sont fonctions propres associées à la valeur propre 1 de l'opérateur solution de l'équation de la chaleur avec conditions de Neumann. En effet, en considérant l'équation de la chaleur avec conditions de Neumann, $(-\Delta + \partial_t)$ appliquée à ces fonctions donnent la fonction nulle et de plus, la dérivée



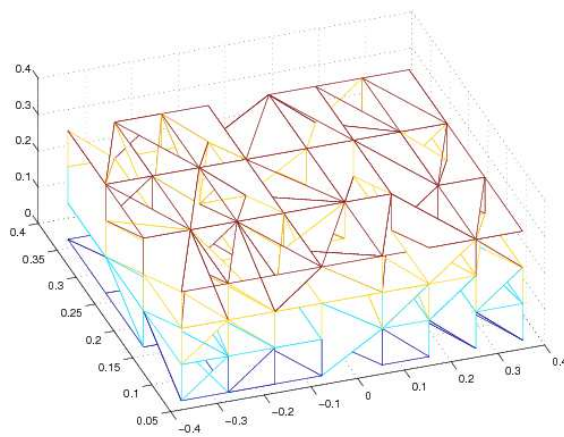
(a) Smiley



(b) Cible



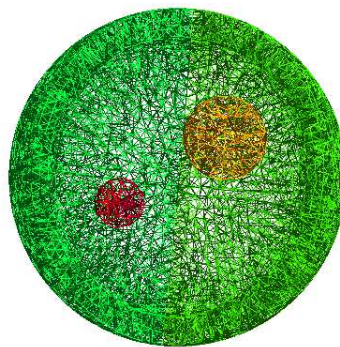
(c) Hypersphère 2D



(d) 6 blocs 3D



(e) 2 sphères tronquées



(f) 3 sphères tronquées

$$\frac{Y}{Z} X$$

$$\frac{Y}{Z} X$$

FIGURE 2.21 – Triangulation de Delaunay sur divers exemples

normale est nulle. Donc considérer ce nouveau problème permet de déterminer exactement les k vecteurs propres qui satisfont les hypothèses du théorème 2.3. Du point de vue de l'algorithme 3, l'étape de normalisation revient à considérer la matrice $D^{-1/2}AD^{-1/2}$ donc les valeurs propres seront inférieures ou égales à 1. Donc les k vecteurs propres satisfaisant les hypothèses du théorème 2.3 correspondent aux k plus grandes valeurs propres. Pour résumer, si la matrice affinité normalisée est proche d'une discrétisation du noyau de Green K_N du problème $(\mathcal{P}_\Omega^N)$ alors on pourra limiter la base Hilbertienne des fonctions propres $(v_n)_n$ de l'opérateur de Laplace avec conditions de Neumann aux k premières fonctions propres associées à la valeur propre 1.

Pour appuyer cette nouvelle approche et pour incorporer l'étape de normalisation, des expérimentations numériques sont effectuées sur les exemples 1 et 2. Sur les figures 2.22 et 2.23 (a)-(b)-(c), les 3 plus petites fonctions propres de l'opérateur de Laplace avec conditions de Neumann aux frontières sont représentées. Et (d)-(e)-(f), les 3 plus grands vecteurs propres de $D^{-1/2}AD^{-1/2}$ associés à la valeur propre 1. Dans les deux cas, l'allure des fonctions propres et des vecteurs propres est similaire et ne réfute pas cette approche de l'étape de normalisation.

Algorithm 3 Algorithme de classification spectrale

Input : Ensemble des données S , Nombre de clusters k

1. Construction de la matrice affinité $A \in \mathbb{R}^{N \times N}$ définie par :

$$A_{ij} = \begin{cases} \exp(-\|x_i - x_j\|^2 / 2\sigma^2) & \text{si } i \neq j, \\ 0 & \text{sinon.} \end{cases}$$

2. Construction de la matrice normalisée $L = D^{-1/2}AD^{-1/2}$ où D matrice diagonale définie par :

$$D_{i,i} = \sum_{j=1}^N A_{ij}.$$

3. Construction de la matrice $X = [X_1 X_2 \dots X_k] \in \mathbb{R}^{N \times k}$ formée à partir des k plus grands vecteurs propres x_i , $i = \{1, \dots, k\}$ de L .
4. Construction de la matrice Y formée en normalisant les lignes de X :

$$Y_{ij} = \frac{X_{ij}}{\left(\sum_j X_{ij}^2\right)^{1/2}}.$$

5. Traiter chaque ligne de Y comme un point de $\mathbb{R}^k$ et les classer en k clusters via la méthode *K-means*.
 6. Assigner le point original x_i au cluster j si et seulement si la ligne i de la matrice Y est assignée au cluster j .
-

2.5.5 Cas limites de validité de l'étude

Des exemples ont été recensés pour montrer les limites du clustering spectral. Nadler et al [80] expliquent que la mesure de similarité, basée sur une notion locale, est insuffisante dans certains cas

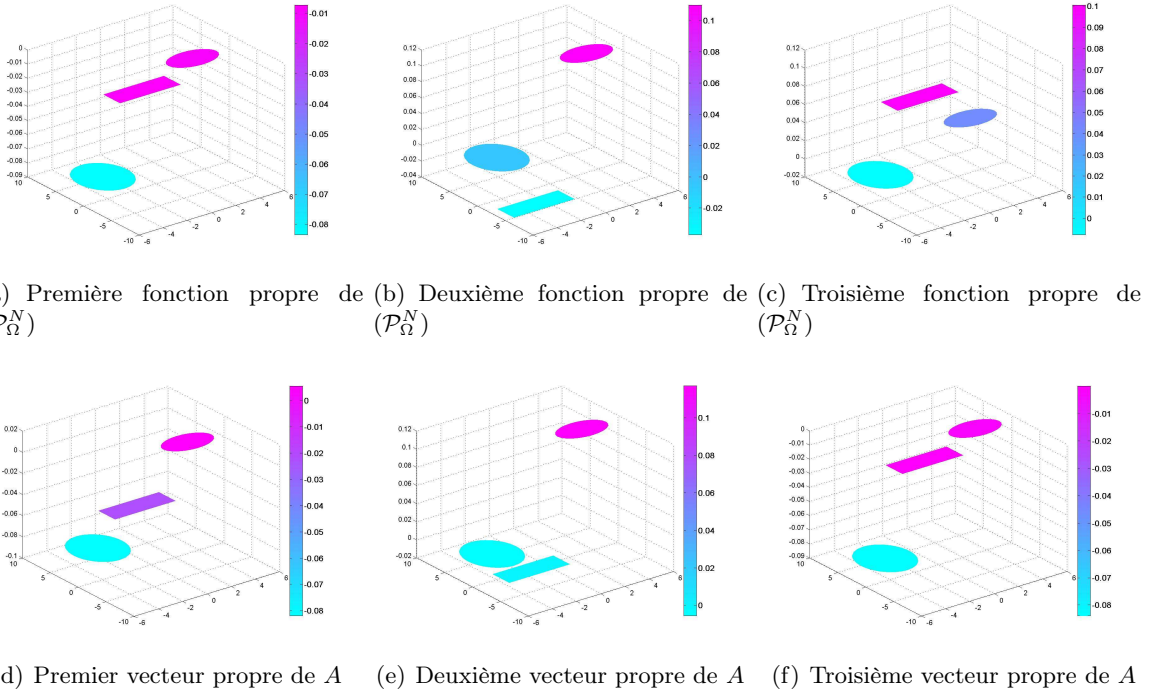


FIGURE 2.22 – Exemple 1 : Vecteurs propres avec étape de normalisation

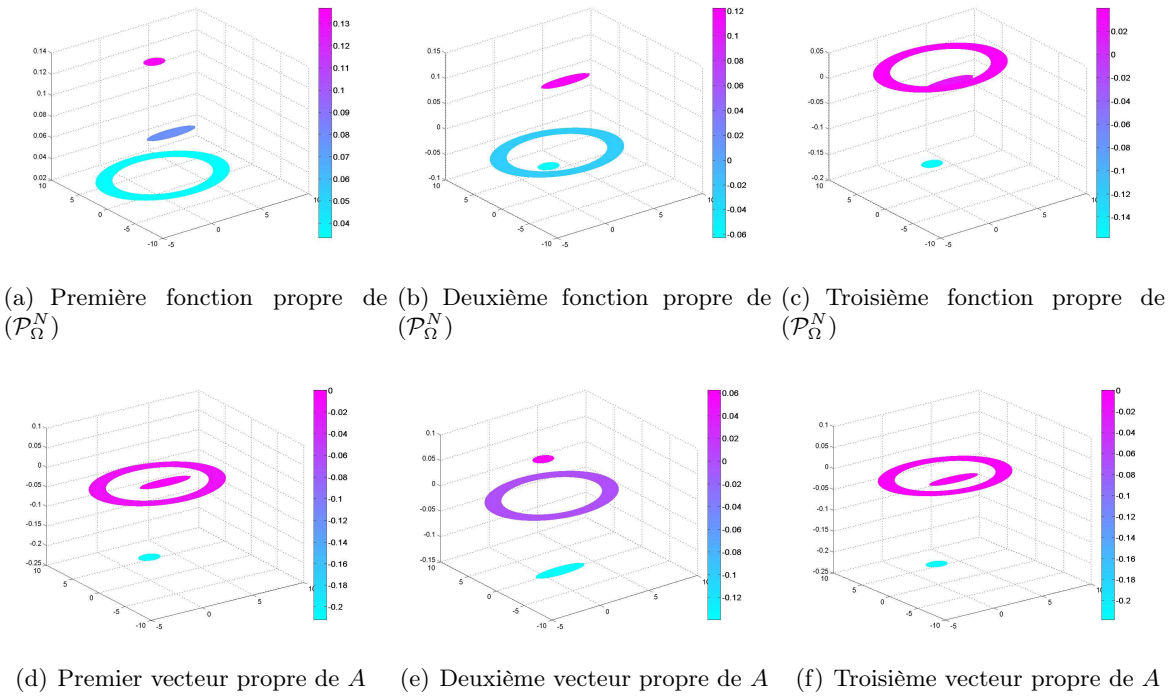


FIGURE 2.23 – Exemple 2 : Vecteurs propres avec étape de normalisation

pour définir une partition globale. Plusieurs exemples montrent ces limites notamment une bille sur une règle ou bien le cas de 3 nuages de points issus de distributions gaussiennes avec des paramètres différents représentés figure 2.24.

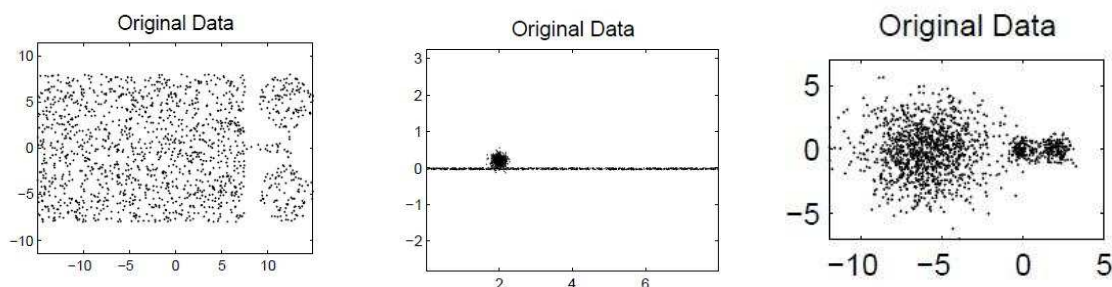


FIGURE 2.24 – Exemples de cas limites [80]

Cette interprétation via l'équation de la chaleur et les éléments finis permettent de donner une explication. En effet, les exemples évoqués ci-dessus présentent des clusters ayant une surface de contact entre eux. Ce genre de configuration ne permet pas d'appliquer des conditions aux frontières entre les clusters. Ainsi, les fonctions propres issues de l'équation de la chaleur ($\mathcal{P}_{\mathbb{R}^p}$) restreinte au support des clusters sont localisées sur une seule composante connexe à support sur les deux clusters. Les exemples traités par Nadler et *al* relèvent donc a priori plutôt de la reconnaissance de forme que de la classification de données.

Chapitre 3

Parallélisation de la classification spectrale

Dans ce chapitre, l’aspect numérique est privilégié. Un des objectifs des méthodes de classification non supervisée est de les appliquer sur de grands ensembles de données. Or l’étape d’extraction des vecteurs propres est très coûteuse car la matrice d’affinité (1.1) est pleine. Ainsi, des limites en terme de coût numérique apparaissent. Plusieurs approches via le calcul parallèle ont été envisagées récemment pour pallier ce défaut. En effet, divers auteurs exploitent des techniques d’Algèbre linéaire comme la méthode de Nyström [44] ou bien des méthodes d’Arnoldi [92, 38] pour une implémentation en parallèle et afin de réduire le coût numérique. Cependant, la matrice affinité entière est construite et utilisée dans l’algorithme. De plus, le problème du choix préalable du nombre de classes k reste ouvert.

Dans ce chapitre, les propriétés théoriques développées au chapitre précédent sont exploitées pour proposer deux nouvelles stratégies de clustering spectral par décompositions en sous-domaines. En appliquant indépendamment l’algorithme sur des sous-domaines particuliers, on montre que la réunion des partitions locales représente une partition globale ou une sous-partition de l’ensemble des points. Nous proposons d’appliquer le clustering spectral sur des sous-domaines directement en divisant l’ensemble des données en sous-ensembles via leurs coordonnées géométriques. Avec un paramètre global d’affinité gaussienne (1.1) et une méthode pour déterminer le nombre de clusters, chaque processus applique indépendamment l’algorithme clustering spectral sur des sous ensembles de données et fournit une partition locale. Une étape de regroupement assure la connexion entre les sous-ensembles de points et détermine une partition globale à partir des locales. Après une première phase d’expérimentations numériques, une alternative sera proposée puis testée sur des cas géométriques et des cas de segmentation d’images. Le potentiel de parallélisme de l’algorithme tout comme son comportement numérique et ses limitations seront étudiés.

3.1 Classification spectrale parallèle : justification

Il a été démontré [11, 12] que l’affinité gaussienne (1.1) peut être interprétée comme une discrétisation du noyau de la chaleur. En particulier, d’après le précédent chapitre, nous avons montré que cette matrice est une représentation discrète d’un opérateur de la chaleur défini sur un domaine, approprié de $\mathbb{R}^p$. Grâce aux propriétés spectrales de l’équation de la chaleur, les vecteurs propres de cette matrice sont asymptotiquement une représentation discrète de fonctions propres L^2 dont les supports sont inclus sur une seule composante connexe à la fois.

Appliquer le clustering spectral revient donc à projeter sur le support de ces fonctions propres

particulières. Ainsi, le clustering spectral est appliqué sur des sous-domaines pour identifier les composantes connexes. Les sous-domaines sont définis de façon directe en divisant l'ensemble des données ponctuelles S en se servant de leurs coordonnées géométriques. Un partitionnement peut alors être extrait indépendamment et en parallèle sur chaque sous-ensemble. Puis, en vue d'une étape de regroupement des partitions locales issues des différents sous-domaines, le clustering spectral est appliqué sur un ensemble spécifique de données correspondant à une bande d'intersection le long des frontières du découpage géométrique des données. La connexion des clusters dont le support s'étend sur plusieurs sous-domaines sera assurée via une relation d'équivalence [86] : $\forall x_{i_1}, x_{i_2}, x_{i_3} \in S$,

$$\text{Si } x_{i_1}, x_{i_2} \in C^1 \text{ et } x_{i_2}, x_{i_3} \in C^2 \text{ alors } C^1 \cup C^2 \subset P \text{ et } x_{i_1}, x_{i_2}, x_{i_3} \in P \quad (3.1)$$

où S est l'ensemble des données, C^1 et C^2 deux clusters distincts et P un cluster plus grand dans lequel sont inclus C^1 et C^2 .

3.1.1 Etude dans le continu

Considérons un ouvert Ω formé de n composantes connexes distinctes Ω_i tel que :

$$\Omega = \bigcup_{i=1}^n \Omega_i \text{ avec } \Omega_i \cap \Omega_j = \emptyset, \forall i \neq j.$$

En reprenant l'hypothèse développée au chapitre 2, on établit l'hypothèse suivante :

Hypothèse 3.1. *Les ouverts Ω_i sont écartés d'une distance minimale notée ϵ telle que :*

$$\inf_{i \neq j} d(\Omega_i, \Omega_j) \geq \epsilon > 0. \quad (3.2)$$

On se donne un pavage par des fermés C_i de l'ouvert $\Omega \subset \bigcup_{i=0}^q C_i$ comme le représente la figure 3.1. Typiquement, les C_i représentent les sous-domaines dans la résolution parallèle. On ne suppose pas que $C_i \cap C_j = \emptyset$ pour $i \neq j$.

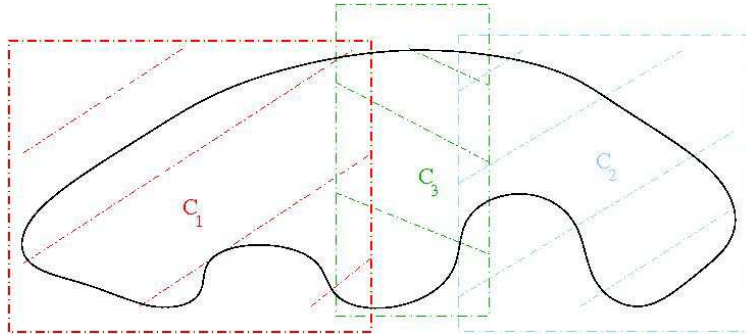
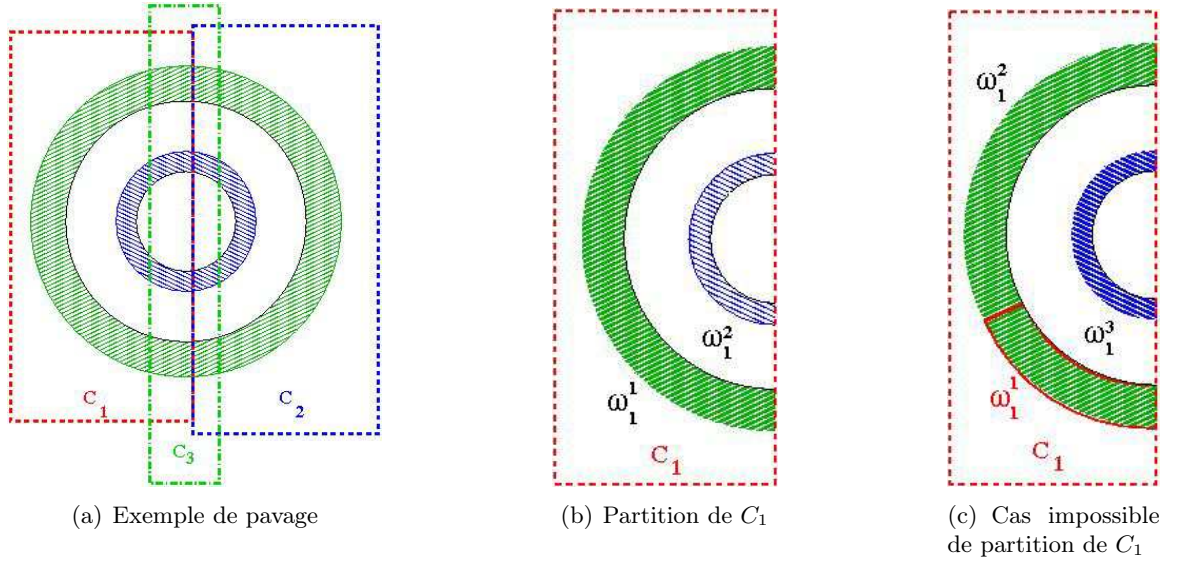


FIGURE 3.1 – Exemple de pavage d'un ouvert Ω

On note $(\omega_i^k)_{k=1..n_i}$ les clusters issus du pavé C_i . Les ω_i^k sont pris comme des ensembles connexes et fermés pour la topologie induite sur Ω . La partition du pavage $C_i \cap \Omega$ est notée : $\mathcal{P}_{C_i} = \{\omega_i^k, k \in \{1, \dots, n_i\}\}$ et on notera $\mathcal{P}_C = \{\omega_i^k, i, k\}$. Le lien entre les clusters de deux pavés est défini par la relation d'équivalence suivante :

FIGURE 3.2 – Exemple de $\omega_1^1 \mathcal{R}^{P_C} \omega_2^1$

Proposition 3.2. Soit $\mathcal{O} = \{\mathcal{O}_i^k, i \in [1, n], k \in [1, n_i]\}$ une famille de sous-ensembles $\mathcal{O}_i^k$ de Ω telle que $\mathcal{O}_i^k \subset C_i$, pour tout $k \in [1, n_i]$ (n_i indiquant le nombre de sous-ensembles considérés dans chaque pavé C_i). On définit $\mathcal{R}^{\mathcal{O}}$ la relation suivante entre deux sous-ensembles $\mathcal{O}_i^k$ et $\mathcal{O}_j^l$ respectivement aux deux pavés C_i et C_j :

$$\mathcal{O}_i^k \mathcal{R}^{\mathcal{O}} \mathcal{O}_j^l \iff \left[\begin{array}{l} \exists r \in \mathbb{N}, \exists \{i_1, \dots, i_r\} \subset [1, n], \exists (k_1, \dots, k_r) \in \Pi_{t=1}^r [1, n_{i_t}] \\ \text{tel que } \forall m \in \{1, \dots, r-1\}, \mathcal{O}_{i_m}^{k_m} \cap \mathcal{O}_{i_{m+1}}^{k_{m+1}} \neq \emptyset \text{ avec } \begin{cases} \mathcal{O}_{i_1}^{k_1} = \mathcal{O}_i^k \\ \mathcal{O}_{i_r}^{k_r} = \mathcal{O}_j^l \end{cases} \end{array} \right] \quad (3.3)$$

Alors $\mathcal{R}^{\mathcal{O}}$ est une relation d'équivalence sur $\mathcal{O}$. On notera $\mathcal{O}/\mathcal{R}^{\mathcal{O}}$ l'ensemble des classes d'équivalence sur $\mathcal{O}$.

Démonstration. Cette relation est évidemment réflexive et symétrique.

Montrons la transitivité de $\mathcal{R}^{\mathcal{O}}$ c'est-à-dire montrons que si $\mathcal{O}_i^k \mathcal{R}^{\mathcal{O}} \mathcal{O}_j^l$ et $\mathcal{O}_j^l \mathcal{R}^{\mathcal{O}} \mathcal{O}_r^s$ alors $\mathcal{O}_i^k \mathcal{R}^{\mathcal{O}} \mathcal{O}_r^s$ avec $i \neq j$ et $k \neq l$. Soient $\mathcal{O}_i^k \mathcal{R}^{\mathcal{O}} \mathcal{O}_j^l$ et $\mathcal{O}_j^l \mathcal{R}^{\mathcal{O}} \mathcal{O}_r^s$ alors il existe deux entiers a_1 et b_1 tels que :

$$\mathcal{O}_i^k \mathcal{R}^{\mathcal{O}} \mathcal{O}_j^l \iff \left[\begin{array}{l} \exists a_1 \in \mathbb{N}, \exists \{i_1, \dots, i_{a_1}\} \subset [1, n], \exists (k_1, \dots, k_{a_1}) \in \Pi_{t=1}^{a_1} [1, n_{i_t}] \text{ tel que } \forall m \in [1, a_1 - 1], \\ \mathcal{O}_{i_m}^{k_m} \cap \mathcal{O}_{i_{m+1}}^{k_{m+1}} \neq \emptyset \text{ avec } \mathcal{O}_{i_1}^{k_1} = \mathcal{O}_i^k \text{ et } \mathcal{O}_{i_{a_1}}^{k_{a_1}} = \mathcal{O}_j^l \end{array} \right].$$

$$\mathcal{O}_j^l \mathcal{R}^{\mathcal{O}} \mathcal{O}_r^s \iff \left[\begin{array}{l} \exists b_1 \in \mathbb{N}, \exists \{j_1, \dots, j_{b_1}\} \subset [1, n], \exists (l_1, \dots, l_{b_1}) \in \Pi_{t=1}^{b_1} [1, n_{j_t}] \text{ tel que } \forall m' \in [1, b_1 - 1], \\ \mathcal{O}_{j_{m'}}^{l_{m'}} \cap \mathcal{O}_{j_{m'+1}}^{l_{m'+1}} \neq \emptyset \text{ avec } \mathcal{O}_{j_1}^{l_1} = \mathcal{O}_j^l \text{ et } \mathcal{O}_{j_{b_1}}^{l_{b_1}} = \mathcal{O}_r^s \end{array} \right].$$

Soient les listes d'indices $\mathcal{I}$ et $\mathcal{I}'$ avec $\text{Card}(\mathcal{I}) = \text{Card}(\mathcal{I}') = q = a_1 + b_1 - 1$, définies par $\mathcal{I} = \{i_1, \dots, i_{a_1-1}, i_{a_1} = j_1, \dots, j_{b_1}\} := \{\mathcal{I}_1, \dots, \mathcal{I}_q\}$ et $\mathcal{I}' = \{k_1, \dots, k_{a_1-1}, i_{a_1} = l_1, \dots, l_{b_1}\} := \{\mathcal{I}'_1, \dots, \mathcal{I}'_q\}$. Alors $\forall m \in \{1, \dots, q-1\}$, $\mathcal{O}_{\mathcal{I}_m}^{\mathcal{I}'_m} \cap \mathcal{O}_{\mathcal{I}_{m+1}}^{\mathcal{I}'_{m+1}} \neq \emptyset$ avec $\mathcal{O}_{\mathcal{I}_1}^{\mathcal{I}'_1} = \mathcal{O}_i^k$ et $\mathcal{O}_{\mathcal{I}_q}^{\mathcal{I}'_q} = \mathcal{O}_r^s$. Donc $\mathcal{O}_i^k \mathcal{R}^{\mathcal{O}} \mathcal{O}_r^s$ et la relation $\mathcal{R}^{\mathcal{O}}$ est transitive et est bien une relation d'équivalence. $\square$

Cette définition est l'équivalent de la relation d'équivalence des ϵ -chaînes [86]. Les ouverts Ω_i qui sont à support sur plusieurs pavés C_i sont reconstitués à l'aide de la proposition suivante.

Théorème 3.3. *Supposons que les éléments de la partition du pavé $C_i \cap \Omega = \bigcup_k \omega_i^k$ pour chaque $i \in [1, q]$ sont exactement les composantes connexes de $C_i \cap \Omega$.*

Alors pour tout $\bar{A} \in \mathcal{P}_C / \mathcal{R}^{\mathcal{P}^C}$, il existe $i \in \{1, \dots, n\}$ tel que $\Omega_i = \bigcup_{\omega_i^k \in \bar{A}} \omega_i^k$.

Réciproquement, pour tout $i \in \{1, \dots, n\}$, il existe $\bar{A} \in \mathcal{P}_C / \mathcal{R}^{\mathcal{P}^C}$ tel que $\bigcup_{\omega_i^k \in \bar{A}} \omega_i^k = \Omega_i$.

Pour établir la preuve de la proposition 3.3, il faut d'abord introduire un lemme technique.

Lemme 3.4. *Soient $i \in [1, q]$ et $\bar{A} \in \mathcal{P}_{C_i} / \mathcal{R}^{\mathcal{P}^C}$. Alors pour tout $\omega \in \bar{A}$ et $p \in [1, n]$, on a :*

$$\omega \in \bar{A}, \omega \cap \Omega_p \neq \emptyset \implies \forall \omega' \in \bar{A}, \omega' \subset \Omega_p.$$

Démonstration. Soient $i \in [1, q]$ tel que $\bar{A} \in \mathcal{P}_{C_i} / \mathcal{R}^{\mathcal{P}^C}$ et $\omega \in \bar{A}$ avec $\Omega_p \cap \omega \neq \emptyset$. Par l'hypothèse (3.1), ω étant connexe, on a $\omega \subset \Omega_p$. Soit maintenant $\omega' \in \bar{A}$. Par la définition de $\mathcal{R}^{\mathcal{P}^C}$, $\exists r \in \mathbb{N}, \exists \{i_1, \dots, i_r\} \subset [1, n], \exists (k_1, \dots, k_r) \in \prod_{t=1}^r [1, n_{i_t}]$ tel que $\forall m \in [1, r-1], \omega_{i_m}^{k_m} \cap \omega_{i_{m+1}}^{k_{m+1}} \neq \emptyset$ avec $\omega_{i_1}^{k_1} = \omega$ et $\omega_{i_r}^{k_r} = \omega'$. Comme $\omega_{i_1}^{k_1} \cap \omega_{i_2}^{k_2} \subset \omega_{i_1}^{k_1} = \omega \subset \Omega_p$, on en déduit que $\omega_{i_2}^{k_2} \cap \Omega_p \neq \emptyset$ et donc, comme précédemment que $\omega_{i_2}^{k_2} \subset \Omega_p$. Par une récurrence directe, on montre donc que pour tout $m \in [1, r-1]$, $\omega_{i_m}^{k_m} \subset \Omega_p$ et donc $\omega' \subset \Omega_p$. Donc pour tout $\omega' \in \bar{A}$, on a $\omega' \subset \Omega_p$. $\square$

Ce lemme montre que toute classe d'équivalence est incluse dans une seule composante connexe Ω_p . Ainsi le théorème 3.3 montre qu'elle correspond exactement à un pavage de la composante connexe.

Démonstration du théorème 3.3. Soit $\bar{A} \in \Omega / \mathcal{R}^{\mathcal{P}^C}$. On note $A = \bigcup_{\omega_i^k \in \bar{A}} \omega_i^k$. Par le lemme 3.4, on sait déjà qu'il existe p tel que $A \subset \Omega_p$. Supposons que $A \neq \Omega_p$, alors il existe $x \in \Omega_p \setminus A$. Par le clustering effectué sur Ω , il existe un ω_i^k tel que $x \in \omega_i^k$. De plus, ω_i^k et Ω_p étant connexes avec $\Omega_p \cap \omega_i^k \neq \emptyset$, car $x \in \Omega_p \cap \omega_i^k$, on a donc $\omega_i^k \subset \Omega_p$. Soient $\mathcal{B}$ l'ensemble défini par $\mathcal{B} = \{\omega_i^k \text{ tel que } \omega_i^k \cap A = \emptyset \text{ et } \omega_i^k \cap \Omega_p \neq \emptyset\}$ et l'ouvert $B = \bigcup_{\omega_i^k \in \mathcal{B}} \omega_i^k$, $B \neq \emptyset$. Alors $B \subset \Omega_p$ et on a $B \cap A = \emptyset$ et $\Omega_p = A \cup B$. Or A et B sont fermés pour la topologie induite sur Ω_p et disjoints ce qui est impossible car Ω_p est connexe. Donc $\Omega_p = A$.

Réciproquement, soit une des composantes connexes Ω_p et soit $x \in \Omega_p$. Alors par recouvrement, il existe un ensemble ω_i^k tel que $x \in \omega_i^k$ et donc par ce qui précède $\Omega_p = \bigcup_{\omega \in \omega_i^k} \omega$. $\square$

Donc à partir d'une décomposition en sous-domaines d'un ensemble de données S , les différentes partitions issues des sous-domaines seront connectées grâce à la relation d'équivalence et au résultat du théorème 3.3.

3.1.2 Passage au discret

Soit S l'ensemble des points de l'ouvert $\Omega = \bigcup_{i=1}^n \Omega_i$ et soit S_i l'ensemble des points sur chaque pavé C_i pour $i \in \{1, \dots, q\}$, $S_i = C_i \cap S, \forall i \in \{1, \dots, q\}$. On note S_i^k les éléments de la partition de S_i et $\mathcal{S} = \{S_i^k, i \in [1, q], k \in [1, n_i]\}$. On suppose que pour tout k et tout $i \in [1, q]$ $S_i^k \neq \emptyset$.

Dès maintenant et dans toute la suite, on notera par $(\omega_i^k)_k$ pour chaque $i \in [1, q]$ les composantes connexes de $C_i \cap \Omega$. D'après le théorème 3.3, on sait que à partir des $(\omega_i^k)_{i,k}$ et de la relation $\mathcal{R}^{\mathcal{P}^C}$, on peut reconstituer les composantes connexes Ω_i . Maintenant, on propose un équivalent du théorème 3.3 pour permettre de regrouper les ensembles de points S_i^k à l'aide de la relation $\mathcal{R}^{\mathcal{S}}$. En outre, si l'ouvert $C_i \cap \Omega$ induit un clustering sur S_i compatible au sens de la définition 2.1, appelé aussi clustering idéal, alors pour chaque $i \in [1, q]$, les éléments S_i^k de la partition du pavé S_i sont exactement les points appartenant aux composantes connexes de $C_i \cap \Omega$, c'est-à-dire $S_i^k = \omega_i^k \cap S_i$.

Théorème 3.5. *Supposons que Ω induit un n -clustering sur S compatible et que de plus, si $\omega_i^k \cap \omega_j^l \neq \emptyset$ pour $i, j \in [1, q]$ alors il existe un point $x \in S_i^k \cap S_j^l$.*

Alors pour tout $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$, il existe $i \in \{1, \dots, n\}$ tel que $\Omega_i \cap S = \bigcup_{S_i^k \in \bar{S}_A} S_i^k$.

Réciproquement, pour tout $i \in \{1, \dots, n\}$, il existe $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$ tel que $\bigcup_{S_i^k \in \bar{S}_A} S_i^k = \Omega_i \cap S$.

Pour établir ce résultat, il faut d'abord montrer le lemme suivant.

Lemme 3.6. *Soit $i \in [1, q]$ et soit $\bar{A} \in \mathcal{P}_C/\mathcal{R}^{\mathcal{P}_C}$. Avec les hypothèses du théorème 3.5, on a les résultats suivants :*

1. *Soient $A, B \in \bar{A}$. Alors $(A \cap S)/\mathcal{R}^S = (B \cap S)/\mathcal{R}^S$;*
2. *$\omega_i^k \in \bar{A} \iff S_i^k \in \bar{S}_A$ avec $\bar{S}_A = (A \cap S)/\mathcal{R}^S$ pour $A \in \bar{A}$ quelconque.*

Démonstration.

1. Soient $A, B \in \bar{A}$ quelconques. Donc $\exists r \in \mathbb{N}, \exists \{i_1, \dots, i_r\} \subset [1, n], \exists (k_1, \dots, k_r) \in \prod_{t=1}^r [1, n_{it}]$ tel que $\forall m \{1, \dots, r\}, \omega_{i_m}^{k_m} \cap \omega_{i_{m+1}}^{k_{m+1}} \neq \emptyset$ avec $\omega_{i_1}^{k_1} = A$ et $\omega_{i_r}^{k_r} = B$. Par intersection avec S et par l'hypothèse du théorème 3.5, $\omega_{i_m}^{k_m} \cap \omega_{i_{m+1}}^{k_{m+1}} \neq \emptyset$ implique que $S \cap \omega_{i_m}^{k_m} \cap \omega_{i_{m+1}}^{k_{m+1}} = S_{i_m}^{k_m} \cap S_{i_{m+1}}^{k_{m+1}} \neq \emptyset$. On en déduit donc que $(A \cap S)/\mathcal{R}^S = (B \cap S)/\mathcal{R}^S$.
2. Soient $A \in \bar{A}$ et $\bar{S}_A = (A \cap S)/\mathcal{R}^S$. D'après 1., la définition de $\bar{S}_A$ est correcte car elle ne dépend pas du choix de $A \in \bar{A}$. Si $\omega_i^k \in \bar{A}$ alors d'après 1., $(\omega_i^k \cap S) = S_i^k \in \bar{S}_A$. Réciproquement, si $S_i^k \in \bar{S}_A$ alors $\exists r \in \mathbb{N}, \exists \{i_1, \dots, i_r\} \subset [1, n], \exists (k_1, \dots, k_r) \in \prod_{t=1}^r [1, n_{it}]$ tel que $\forall m \{1, \dots, r\}, S_{i_m}^{k_m} \cap S_{i_{m+1}}^{k_{m+1}} \neq \emptyset$ avec $S_{i_1}^{k_1} = S_i^k = \omega_i^k \cap S$ et $\omega_{i_r}^{k_r} = A \cap S$. Or $\emptyset \neq S_{i_m}^{k_m} \cap S_{i_{m+1}}^{k_{m+1}} = S \cap \omega_{i_m}^{k_m} \cap \omega_{i_{m+1}}^{k_{m+1}} \subset \omega_{i_m}^{k_m} \cap \omega_{i_{m+1}}^{k_{m+1}}$ d'où $\omega_i^k \mathcal{R}^{\mathcal{P}_C} A$. □

Démonstration du théorème 3.5. Soit $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$. Soit $A \in \bar{S}_A$ et on note $\bar{A} = A/\mathcal{R}^{\mathcal{P}_C}$. D'après le théorème 3.3, pour un certain $i \in [1, q]$, on a $\bigcup_{\omega_i^k \in \bar{A}} \omega_i^k = \Omega_i$. Comme $\omega_i^k \cap S = S_i^k$, par le lemme 3.6 avec $\bar{S}_A = (A \cap S)/\mathcal{R}^S$ pour $A \in \bar{A}$, on en déduit que $\Omega_i \cap S = \bigcup_{\omega_i^k \in \bar{A}} \omega_i^k \cap S = \bigcup_{\omega_i^k \in \bar{A}} S_i^k = \bigcup_{S_i^k \in \bar{S}_A} S_i^k$.

Réciproquement, soit $i \in [1, q]$, par le théorème 3.3 il existe $\bar{A} = A/\mathcal{R}^{\mathcal{P}_C}$ tel que $\bigcup_{\omega_i^k \in \bar{A}} \omega_i^k = \Omega_i$. En posant $\bar{S}_A = (A \cap S)/\mathcal{R}^S$, on a : $\bigcup_{S_i^k \in \bar{S}_A} S_i^k = \bigcup_{\omega_i^k \in \bar{A}} S_i^k = \bigcup_{\omega_i^k \in \bar{A}} \omega_i^k \cap S = \Omega_i \cap S$. □

Le théorème précédent valide donc la méthode de décomposition en sous-domaines. En effet, une composante connexe à support sur plusieurs pavés du découpage est reconstituée via la relation d'équivalence $\mathcal{R}^S$ à condition que l'hypothèse du théorème 3.3 soit vérifiée. Ainsi les partitions issues des différents sous-domaines définissent une partition globale comme étant l'ensemble des points de S couvrant chaque composante connexe Ω_i de l'ensemble Ω , c'est-à-dire que chaque cluster correspond à $S \cap \Omega_i$.

Ce résultat est valide à condition de satisfaire l'hypothèse du théorème 3.5 d'intersection non vide entre les clusters : $\omega_i^k \cap \omega_j^l \neq \emptyset$ pour $i, j \in [1, q]$.

Cependant deux problèmes restent ouverts avant de présenter une stratégie de parallélisation. Le premier concerne le paramètre de l'affinité gaussienne qui doit être fixée de manière automatique. Pour cela, les résultats théoriques et heuristiques des chapitres 1 et 2 permettent de définir un choix en conservant la propriété géométrique de clustering. Le deuxième problème concerne le choix du nombre de classes k . En effet, la décomposition en sous-domaines implique que le nombre de clusters

k ne devra pas être fixé en amont puisque la distribution des points variant entre les domaines, ce nombre variera d'un sous-domaine à un autre. Par conséquent, une heuristique doit être définie pour automatiser le choix sur chaque sous domaine sans information a priori.

3.1.3 Définition du nombre de classes k

Le problème du choix du nombre de classes est un problème général pour les algorithmes de classification non supervisée. De nombreuses méthodes existent. Certaines sont basées sur des estimateurs issus de la vraisemblance entre les données [45]. Dans le cadre non supervisé où nous disposons de peu d'informations, de nombreux critères ont été définis et sont principalement divisés, d'une part, en critères internes comme la mesure de similarité définie en utilisant des métriques différentes (euclidienne, Hartigan...) [97], et d'autre part en critères externes soit basés sur des modèles physiques [116, 42] soit sur la mesure de l'écart statistique [15], la stabilité de la partition [15, 14] ou encore la prédiction de la partition [36]. Concernant l'interprétation du spectral clustering via la théorie de perturbation matricielle, une heuristique sur l'écart entre les valeurs propres peut être définie. Cependant ces heuristiques souffrent d'une ou plusieurs des limitations suivantes : choix de la métrique, sensibilité des fonctions coûts, coût numérique des estimations.

La problématique du choix de k est d'autant plus difficile à résoudre car ce nombre peut varier d'un sous-domaine à un autre dans une stratégie de décomposition en sous-domaines. Pour ce faire, nous considérons pour chaque sous-domaine la mesure de qualité basée sur les ratio de normes de Frobenius, présentée au chapitre 1, pour évaluer le nombre de classes.

Heuristique 3.7. Soit n_k un nombre limite de classes à chercher. Après avoir appliqué le spectral clustering pour un nombre de cluster $k' \in [2, n_k]$, la matrice affinité A définie par (1.1) est réordonnée par classe. On obtient la matrice par bloc, notée L , telle que les blocs hors diagonaux représentent les affinités entre les classes et les blocs diagonaux l'affinité intra-classe. Les ratios entre les normes de Frobenius des blocs diagonaux et ceux hors-diagonaux sont évalués pour $i, j \in [1, k']$ et $i \neq j$:

$$r_{ij} = \frac{\|L^{(ij)}\|_F}{\|L^{(ii)}\|_F},$$

avec $\|L^{(ij)}\|_F = \left(\sum_{l=1}^{N_i} \sum_{m=1}^{N_j} |L_{lm}^{(ij)}|^2 \right)^{\frac{1}{2}}$ et où N_i et N_j sont les dimensions du bloc (ij) de L .

Soit $\eta_{k'}$ le ratio moyen des r_{ij} pour une valeur $k' \in [1, n_k]$ défini par :

$$\eta_{k'} = \frac{2}{k'(k'-1)} \sum_{\substack{i=1 \\ j=i+1}}^{k'} r_{ij}. \quad (3.4)$$

Alors le nombre de classe k satisfait la condition suivante, pour tout $k' \in [2, n_k]$:

$$k = \arg \min_{k' \in [2, n_k]} \eta_{k'}. \quad (3.5)$$

Par définition, le nombre approprié de classes k correspond à une situation où des points qui appartiennent à des classes différents aient le moins d'affinité entre eux et, dans le cas contraire, une forte affinité entre eux s'ils appartiennent au même classe. Parmi diverses valeurs de k , le nombre de classe final est défini de telle sorte que l'affinité soit la plus faible entre classes et la plus forte au sein des classes. L'équation (3.5) donne un ratio moyen de l'affinité entre les classes. Si le ratio η_k est proche de 0 alors la matrice affinité réordonnée par classe a une structure bloc-diagonale,

proche du cas idéal et donc le partitionnement obtenu est meilleur.

Sur la figure 3.3, le ratio $\eta_k = \frac{2}{k(k-1)} \sum_{j=i+1}^k r_{ij}$ en fonction du nombre de classes k est représenté sur deux exemples représentant des classes aux densités variées. Rappelons que plus le ratio η est proche de 0, plus la matrice réordonnée a une structure bloc diagonale. Ainsi, on obtient une partition où les classes sont séparées par une distance suffisante assurant des vecteurs propres constants par morceaux. Ainsi, les données projetées dans l'espace spectral sont concentrées par classes et séparées. Dans les deux exemples, le minimum est atteint pour la valeur de k adéquate.

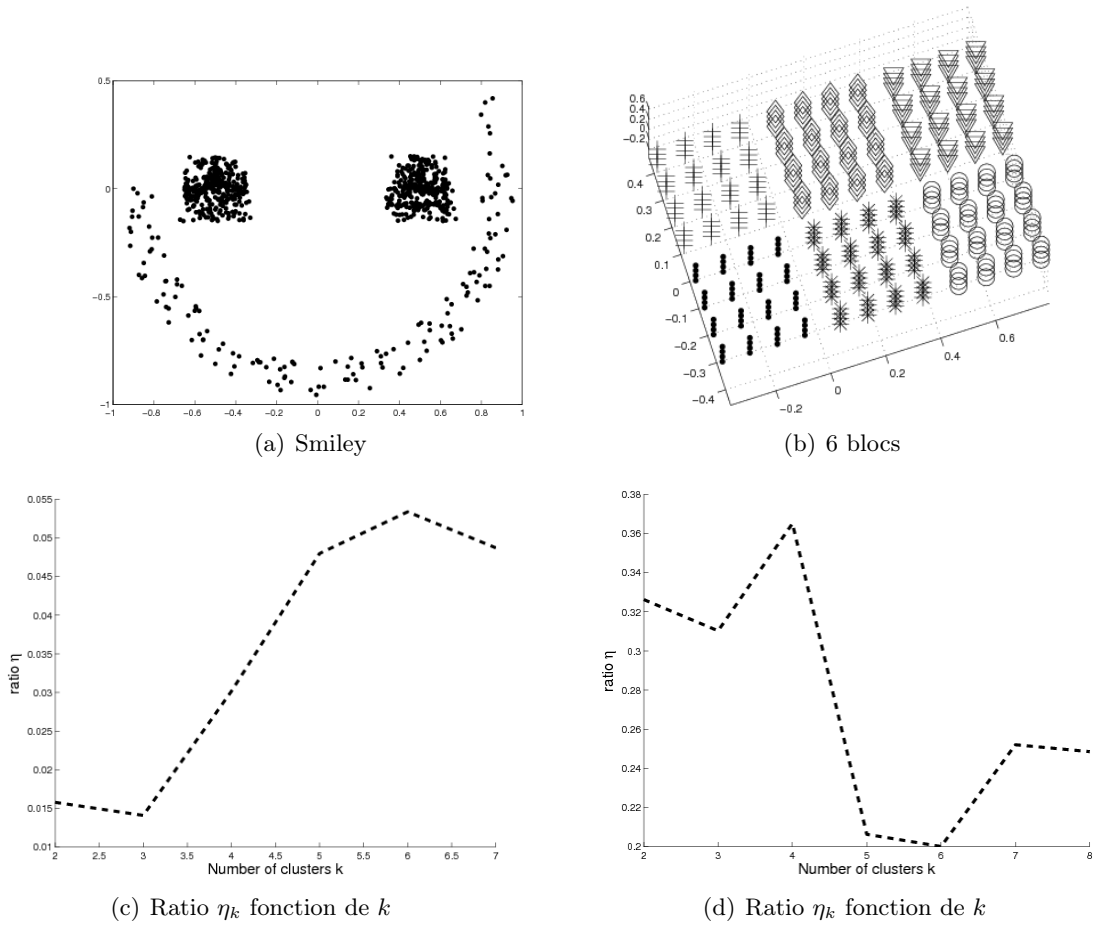


FIGURE 3.3 – Exemples pour le choix du nombre de cluster k

Par contre, diviser l'ensemble entier des données en sous-ensembles peut aboutir à des sous-domaines ne contenant qu'une seule classe. Donc si le nombre de classes k satisfaisant l'équation (3.5) est égal à 2 sur un sous-domaine, le numérateur du ratio η_2 est comparé à son dénominateur. On définit un seuil β qui représentera la proportion pour laquelle l'affinité entre les classes est suffisante pour déclarer qu'il existe 2 classes. Donc, si le ratio η_2 est plus grand que la valeur de β , la valeur de k est fixée à 1 au lieu de 2.

D'un point de vue numérique, cette heuristique exploite les éléments issus des premières étapes de l'algorithme sans avoir à les reprendre : elle revient à concaténer des vecteurs propres, à appliquer le k -means pour obtenir la partition dans l'espace de projection spectral, et à utiliser directement

la matrice affinité en la réordonnant.

3.2 Classification spectrale parallèle avec interface

Dans cette section, nous présentons une première stratégie de décomposition en sous-domaine du spectral clustering. Les diverses étapes de la stratégie parallèle seront étudiées puis la cohérence de la méthode sera vérifiée et enfin des exemples géométriques avec différents maillages seront testés et étudiés. Dans toute la suite, nous considérons l'ensemble de points $S = \{x_i\}_{i=1..N}$.

3.2.1 Implémentation : composantes de l'algorithme

Dans cette section, les différentes composantes de la stratégie parallèle de décomposition en sous-domaines, représentées sur la figure 3.4, sont détaillées.

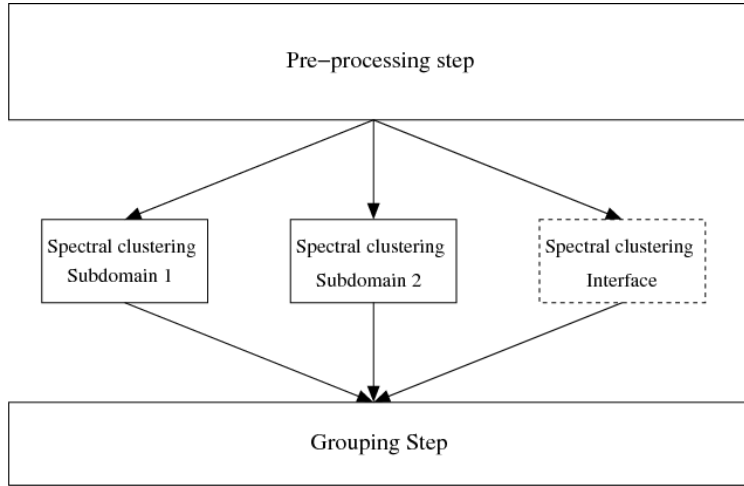


FIGURE 3.4 – Principe du spectral clustering parallèle divisé en $q = 2$ sous-domaines

Etape de prétraitement : partitionner S en q sous-domaines

Incluons toutes les données dans une boîte de côté l_i pour la i^{me} dimension, $i = \{1, \dots, p\}$ où :

$$l_i = \max_{1 \leq i_1, i_2 \leq N} |x_{i_1}(i) - x_{i_2}(i)|, \forall i \in \{1, \dots, p\}. \quad (3.6)$$

En prenant la longueur maximale sur chaque dimension, la boîte est divisée en q sous-boîtes où $q = \prod_{i=1}^p q_i$ et q_i correspond au nombre de subdivisions sur la i -ème dimension. Alors le paramètre d'affinité σ est calculé comme indiqué dans l'heuristique (1.2). Le nombre de processus, noté $nbproc$, est égal à $nbproc = q + 1$.

Décomposition par domaine : interface et sous-domaines

1. Interface

L'interface inclut tous les points qui se trouvent à une distance maximale de γ le long des frontières des découpes. Cette interface a pour fonction de permettre de reconnecter ensemble les clusters dont les points sont sur plusieurs sous-domaines. Prendre une largeur de bande

$\gamma = 3\sigma$ fonction de σ (distance de référence dans le cas d'une distribution uniforme), permet de regrouper les points pour connecter les partitions locales.

Comme l'ensemble des données de l'interface ne couvre pas le même volume que les autres sous-domaines pavés, l'hypothèse d'isotropie n'est pas satisfaite et un paramètre particulier, noté σ^* , d'affinité Gaussienne doit être considéré sur l'interface. La même idée que celle qui nous a servi pour définir l'affinité 1.2 dans le cas d'une distribution isotrope est reprise mais en adaptant le volume adéquat sur l'interface :

$$\sigma^* = \frac{Vol(interface)}{N_{interface}^{\frac{1}{p}}}$$

où $Vol(interface)$ représente le volume réel de l'interface et $N_{interface}$ le nombre de points dans l'interface. Le volume de l'interface est fonction de la largeur γ , du nombre de découpages q et des longueurs $l_1, \dots, l_p$ des boîtes sur chaque dimension comme suit la formule :

$$Vol(interface) = \sum_{i=1}^p (q_i - 1) \gamma^{p-1} l_i - \gamma^p \prod_{i=1}^p (q_i - 1). \quad (3.7)$$

2. Sous-domaines

Chaque processus de 1 à $nbproc$ a un sous-ensemble de données S_i , $i = 1 \dots nbproc$ dont les coordonnées sont incluses dans une sous-boîte géométrique. Le paramètre d'affinité Gaussienne est défini par l'équation suivante (1.2) avec l'ensemble initial S de données :

$$\sigma = \frac{D_{\max}}{N^{\frac{1}{p}}}.$$

Le paramètre est volontairement conservé global et basé sur l'ensemble entier des données pour garder une distance référence, c'est-à-dire un seuil pour la partition globale. Garder le caractère global du paramètre permet de garder des partitions locales en accord avec la partition globale et avec une stratégie de décomposition en sous-domaines.

Classification spectrale sur les sous-domaines

Outre le choix du paramètre spécifique pour l'interface, quelques éléments de l'algorithme doivent être précisés.

1. Nombre de clusters k

Un nombre limite de clusters $nblimit$ est fixé a priori (par défaut, $nblimit = 20$), et correspond le nombre de vecteurs propres que l'on va extraire sur chaque sous-problème. Pour chaque sous-domaine, le nombre de classes k est choisi de telle sorte qu'il satisfasse l'équation (3.3).

2. Espace de projection spectrale

La méthode k -means est utilisée pour définir les classes dans l'espace de projection spectrale. Il existe de nombreuses études sur l'initialisation des centroïdes pour éviter les minima locaux via, entre autres, des estimateurs basés sur la distribution des données, des propriétés géométriques, de la théorie des graphes [4, 19, 66, 57].

L'étude théorique de la méthode de spectral clustering a permis de définir une propriété de clustering fonction du paramètre σ permettant de prouver la séparabilité entre les données

dans l'espace de projection. Dans cette partie, nous ne nous intéressons qu'à l'initialisation des k centroïds, une étape difficile à définir et nous exploitons les propriétés des données dans cet espace. Etant donné que les points sont projetés dans une p -hypersphère unité et qu'ils sont concentrés autour de k points (cf figure 1.2) lorsque les données sont bien séparées, les k centroïds doivent être choisis de manière adéquate pour faire converger l'algorithme. Parmi les points projetés, les centroïds sont choisis tels qu'ils soient les plus éloignés les uns des autres le long d'une même direction. En d'autres termes, soit Y_j un centroïd de l'espace de projection spectral, le centroïd Y_m vérifie :

$$m = \arg \max_j \|Y_i - Y_j\|_\infty.$$

Etape de regroupement

La partition finale est formée du regroupement des partitions issues des $nbproc - 1$ analyses indépendantes du clustering spectral sur les $nbprocs - 1$ sous-ensembles de S . Le regroupement est réalisé à l'aide de la partition de l'interface et de la relation d'équivalence (3.3). Si un point appartient à deux classes différentes, ces deux classes sont assemblées en un plus grand. En sortie de la méthode en parallèle, une partition de l'ensemble S des données et un nombre final de classes k sont obtenus.

Détails liés à la programmation

Quelques détails concernant plus spécifiquement la programmation de la méthode de spectral clustering parallèle doivent être précisés.

1. Calcul du spectre de la matrice affinité

Des routines classiques de la librairie LAPACK [6] sont utilisées pour calculer les valeurs propres et les vecteurs propres de la matrice affinité normalisée relative à chaque sous-ensemble de points.

2. Parallélisation

La parallélisation est effectuée à l'aide de la librairie MPI (Message Passing Interface) qui permet aux $nbproc$ processus de s'exécuter de manière autonome durant leur étape de clustering. L'étape de découpage des données est menée par le premier processus qui transmet ensuite les points de chaque sous-domaine au processus correspondant afin que celui-ci puisse démarrer son clustering indépendamment des autres processus. L'étape de regroupement final des clusters, quant à elle, est assurée par le premier processus lorsque toutes les classifications ont été effectuées.

Un exemple du fonctionnement de la méthode est représenté avec le cas d'une cible divisée en $q = 4$ sous-boîtes sur la figure 3.5. Sur la gauche, le résultat du clustering spectral est affiché pour l'interface. Chaque couleur correspond à une classe. L'interface comprend 13 classes alors que le résultat du clustering sur les 4 sous-domaines, représenté sur la droite de la figure 3.5, a défini 4 classes par sous-domaine. Le résultat global est $k = 4$ classes : une classe représentant un cercle.

3.2.2 Justification de la cohérence de la méthode

Dans cette section, la cohérence de la méthode précédemment introduite est étudiée. Notons par $\hat{S}_k = C_k \cap S$ l'ensemble de points pour le sous-domaine k , $k \in [1, q + 1]$. L'objectif est de

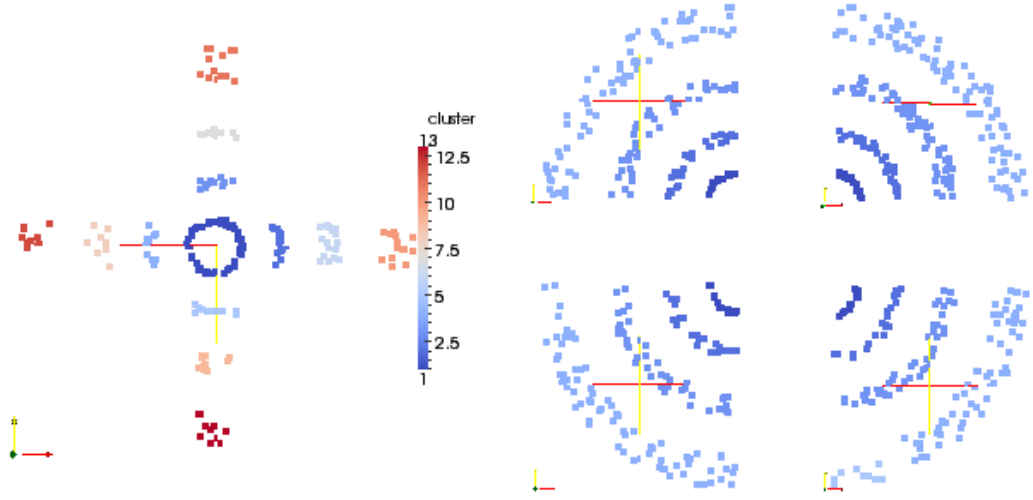


FIGURE 3.5 – Exemple cible : Résultat du spectral clustering pour l'interface et les sous-domaines

montrer qu'à partir d'un clustering idéal de l'ensemble des points $\hat{S}_k$ du sous-domaine k , l'étape de regroupement permet reconstituer les ensembles $S_i = S \cap \Omega_i$. La difficulté est de montrer que le clustering sur les pavés C_k , $k \in \llbracket 1, q \rrbracket$ défini par la stratégie de découpe avec interface permet de définir les classes ω_i^k au sens de l'hypothèse du théorème 3.3 de sorte à obtenir le clustering sur les points S au sens du théorème 3.5. La difficulté pour pouvoir appliquer le théorème 3.5 est de satisfaire l'hypothèse que $\omega_i^k \cap \omega_j^l \neq \emptyset$ pour $i \in \llbracket 1, q+1 \rrbracket$ implique l'existence d'un point $x \in S_i^k \cap S_j^l$. Toutefois, le clustering est réalisé sur les points de S et donc il existe une infinité de façons de définir les pavés fermés C'_k de manière à satisfaire $\hat{S}_k = S \cap C_k = S \cap C'_k$. L'objectif de cette partie est de définir un tel ensemble de pavés C'_k "convenables" pour $k \in \llbracket 1, q+1 \rrbracket$.

Soient $C_1, \dots, C_q$ les q sous-domaines qui permettent le partitionnement géométrique de l'ensemble des données S et soit C_{q+1} le pavé de l'interface à distance γ pour chaque frontière ∂C_j des pavés C_j soit : $C_{q+1} = \{x \in \mathbb{R}^p, \exists j \in \llbracket 1, q \rrbracket, d(x, \partial C_j) < \gamma\}$. De plus, conformément à la stratégie présentée, on suppose que les points appartenant à deux pavés C_i et C_j pour $i \neq j$ sont sur la frontière de ∂C_i et ∂C_j soit : $\text{Int}(C_i) \cap \text{Int}(C_j) = \emptyset, \forall i, j \in \llbracket 1, q \rrbracket$.

Définition 3.8. Soit d_c le maximum de la distance minimale entre un point $x_j \in C_j$ et son plus proche voisin x_k dans un pavé C_k avec $j \neq k$ sur une même composante connexe Ω_i , soit encore :

$$d_c = \max_{\Omega_i} d_{c_i} \quad \text{avec} \quad d_{c_i} = \min_{\substack{x_j \in \Omega_i \cap C_j \cap S, \\ x_k \in \Omega_i \cap C_k \cap S, \\ j \neq k}} d(x_j, x_k). \quad (3.8)$$

Par convention, la distance d_{c_i} sur un ensemble vide sera égale à 0.

On rappelle que $(\omega_i^k)_{k=1..n_i}$ sont les clusters issus du pavé C_i . Soient maintenant les ensembles C'_k pour $k \in \llbracket 1, q+1 \rrbracket$ construits de la manière suivante : on commence avec $C'_i = C_i$ pour $i \in \llbracket 1, q+1 \rrbracket$ puis on applique la règle suivante,

1. si $\omega_i^k \cap \omega_{q+1}^l \neq \emptyset$ et $\omega_i^k \cap \omega_{q+1}^l \cap S = \emptyset$ pour $i \neq q+1$ alors on pose $C'_{q+1} := C'_{q+1} \setminus (\omega_i^k \cap \omega_{q+1}^l)$;

2. si $\omega_i^k \cap \omega_j^l \neq \emptyset$ et $\omega_i^k \cap \omega_j^l \cap S = \emptyset$ pour $i \neq j$, $i \neq q+1$ et $j \neq q+1$ alors on pose $C'_i := C'_i \setminus \Gamma_i^\epsilon$ où $\Gamma_i^\epsilon = \omega_i^k \cap \{x \in \Omega, d(x, \partial\omega_j^l) < \epsilon\}$ pour un $\epsilon > 0$ donné tel que $\epsilon < d_c$ et $\Gamma_i^\epsilon \cap S = \emptyset$. De même, on pose $C'_j := C'_j \setminus \Gamma_j^\epsilon$ pour un ϵ vérifiant les mêmes hypothèses que précédemment.

Ces deux règles sont représentées sur la figure 3.6.

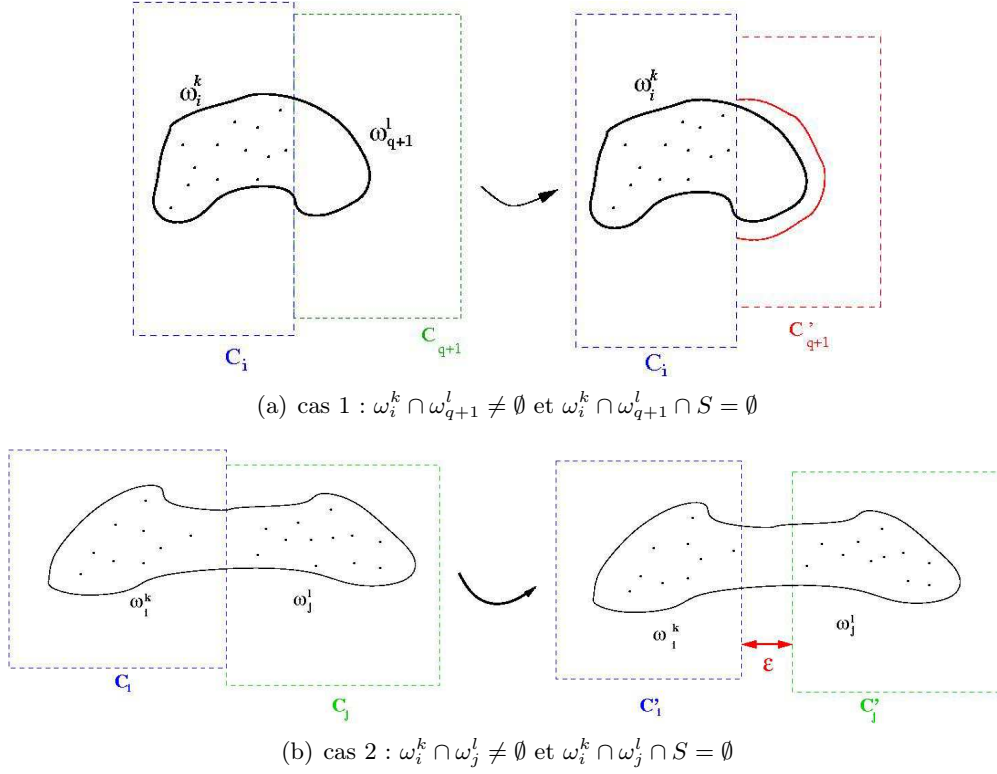


FIGURE 3.6 – Règles pour l'absence de points dans une zone d'intersection

Lemme 3.9. Les C'_k pour $k \in [1, q+1]$ vérifient les propriétés suivantes :

1. C'_k est fermé,
2. $\Omega \subset \bigcup_{k \in [1, q+1]} C'_k$,
3. $C'_k \cap S = \hat{S}_k$.

Démonstration.

1. Par construction des C'_k , ce sont des ensembles fermés (les C_k) auxquels on retranche, dans les deux cas, une intersection d'ouverts. Ils restent donc fermés.
2. Par hypothèse $\Omega \subset \bigcup_{k \in [1, q+1]} C_k$. Il suffit de montrer que les éléments retirés dans les cas 1 et 2 sont toujours présents dans la réunion des C'_k . Pour le cas 1, on retranche $\omega_i^k \cap \omega_{q+1}^l$ de C'_{q+1} mais reste $\omega_i^k \cap \omega_{q+1}^l \subset C'_i \subset \bigcup_{k \in [1, q+1]} C'_k$. Pour le cas 2, on retranche Γ_i^ϵ et Γ_j^ϵ à C'_i et C'_j respectivement, mais d'après l'hypothèse $\epsilon < d_c$, Γ_i^ϵ et Γ_j^ϵ restent inclus dans $C'_{q+1} \subset \bigcup_{k \in [1, q+1]} C'_k$.
3. Par hypothèse, dans les cas 1 et 2, les ensembles retranchés sont d'intersection vide avec l'ensemble S donc l'ensemble $C_k \cap S = \hat{S}_k = C'_k \cap S$ est préservé.

□

On note par $((\omega')_i^k)_{k=1..n_i}$ les clusters issus du pavé C'_i . Par le lemme 3.9, les C'_k forment un ensemble de pavés "convenables" et donc par le théorème 3.3, permettent de retrouver les ω_i^k .

Théorème 3.10. *Supposons que la classification de S est idéale sur chaque $\hat{S}_i$, $i \in [1, q+1]$ et que de plus, $\gamma > d_c$. Alors pour tout $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$, il existe $i \in \{1, ..n\}$ tel que $\Omega_i \cap S = \bigcup_{S_i^k \in \bar{S}_A} S_i^k$. Réciproquement, pour tout $i \in \{1, .., n\}$, il existe $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$ tel que $\bigcup_{S_i^k \in \bar{S}_A} S_i^k = \Omega_i \cap S$.*

Démonstration. Il suffit de montrer que $(\omega')_i^k \cap (\omega')_j^l \neq \emptyset$ pour $i \in [1, q]$ alors il existe un point $x \in S_i^k \cap S_j^l$, le résultat se déduira du théorème 3.5. Or par la définition des C'_k et le lemme 3.9, on a $(\omega')_i^k \cap (\omega')_j^l \cap S = \emptyset$ implique $(\omega')_i^k \cap (\omega')_j^l = \emptyset$. □

Donc, si sur chaque sous-domaine la classification est idéale, alors l'étape de regroupement de la stratégie par interface permet de reconstituer les clusters $\Omega_i \cap S$ via la relation d'équivalence $\mathcal{R}^S$.

Choix du paramètre γ

Pour regrouper des clusters à support sur plusieurs pavés, il faut vérifier $\gamma > d_c$ d'après le théorème 3.10 où d_c représente le maximum de la distance minimale entre un point x_j et son plus proche voisin sur un sous-domaine différent mais au sein de la même composante connexe. Par définition de l'ensemble des points correspondant à l'interface, la largeur de la bande γ est fixée à 3σ où σ représente la distance de référence dans le cas d'une distribution uniforme c'est-à-dire lorsque les points sont supposés tous équidistants les uns des autres. Nous faisons l'hypothèse pour la définition des composantes connexes Ω_k qu'aucune d'entre elles ne peut être subdivisée en deux ouverts recouvrant le même sous-ensemble de points de S et à distance supérieure à γ . Avec cette valeur $\gamma = 3\sigma$, des composantes connexes à support sur plusieurs sous-domaines ont nécessairement des points dans la zone d'intersection. Par conséquent, pour une distribution donnée, trois différentes configurations concernant la densité de clusters sont envisagées :

- si les clusters sont compacts alors la distance entre les points d'un même clusters est strictement inférieure à la distance σ ,
- si les clusters sont tous de faible densité alors la distance entre les points d'un même cluster est de l'ordre de σ ,
- si un cluster a une densité variable avec des zones plus denses que d'autres alors d'après l'hypothèse la séparation entre les zones condensées et moins condensées est inférieure à la séparation entre deux clusters et donc inférieure à 3σ .

Donc cette heuristique pour le choix $\gamma = 3\sigma$ convient dans tous ces cas et satisfait l'inégalité $\gamma > d_c$ du théorème 3.10.

3.2.3 Expérimentations parallèles

La méthode de spectral clustering parallèle doit être testée sur des exemples de grande taille représentant des cas de clustering non triviaux. Pour cela, deux exemples géométriques 3D issus d'un maillage de Delaunay où les données par processeurs sont respectivement équilibrées et non équilibrées sont considérés. Ainsi, en raffinant le maillage, le comportement de la stratégie parallèle sera étudié au fur et à mesure. Les expérimentations numériques sont réalisées sur le supercalculateur Hyperion¹ du CICT (Centre Inter-universitaire de Calcul de Toulouse). Avec ses 352 noeuds

1. <http://www.calmip.cict.fr/spip/spip.php?rubrique90>

N	Nombre de processeurs	Nombre maximal de données par processeur	Temps total (sec)	% du temps total pour le clustering spectral
1905	1	-	21.69	99.58
	3	964	2.75	94.9
	5	1017 (int.)	3.25	96.92
	9	1396 (int.)	8.40	97.98
3560	1	-	127.27	99.85
	3	1020	18.35	98.75
	5	1889 (int.)	21.54	98.84
	9	2618 (int.)	55.13	99.17
10717	1	-	3626.59	99.97
	3	5589	546.93	99.79
	5	5865 (int.)	641.13	99.78
	9	8041 (int.)	1590.45	99.87

TABLE 3.1 – 2 sphères tronquées : cas avec interface

(8 cores/noeud), ce système présente une puissance de 33 Téraflops. Chaque noeud a une mémoire dédiée de 4.5 Go pour chaque core avec un total de 32GB de mémoire disponible sur le noeud.

Exemples géométriques

Le premier exemple représente deux sphères tronquées non-concentriques non-séparables par hyperplan. Cet exemple présente une distribution relativement équilibrée par sous-domaines pour un découpage donné, comme le montre le cas d'un découpage en $q = 2$ sous-domaines représenté sur la figure 3.7. Le résultat du spectral clustering sur les 2 sous-domaines et l'interface est respectivement affiché sur les figures 3.7 (a)-(b) et (c). Le résultat final, après l'étape de regroupement, représenté sur la figure (d), comprend 3 clusters : chaque classe représente une sphère tronquée. Afin d'étudier la stratégie de parallélisation, cet exemple géométrique est découpé en $q = \{2, 4, 8\}$ sous-domaines et testé sur des ensembles de points comprenant de $N = 1905$ à $N = 10717$ points. Les temps pour chaque étape de la méthode parallèle sont mesurés. Le tableau 3.1 indique suivant l'ensemble de données étudié et la découpe en sous-domaine appliquée, le nombre de données dans l'interface, le temps total d'exécution de la méthode et le pourcentage du temps total consacré à l'étape de spectral clustering sur les sous-domaines et l'interface.

De plus, sur la figure 3.8, sont représentés en fonction du nombre de processeurs :

- le speed-up : ratio entre le temps parallèle et le temps séquentiel où le temps parallèle est le temps total pour une décomposition de S en $nbprocs$ et le temps séquentiel correspond le temps pour $nbproc = 1$;
- l'efficacité : ratio entre speed-up et le nombre de processeurs.

Ces deux mesures permettent de juger du gain dans la stratégie de parallélisation en fonction du nombre de processeurs et du raffinement du maillage. Pour chaque découpe et chaque ensemble de points considérés, le nombre final de clusters est 2 et les points associés chaque cluster correspond à l'ensemble des point d'une sphère tronquée. De ces résultats, nous pouvons donner quelques

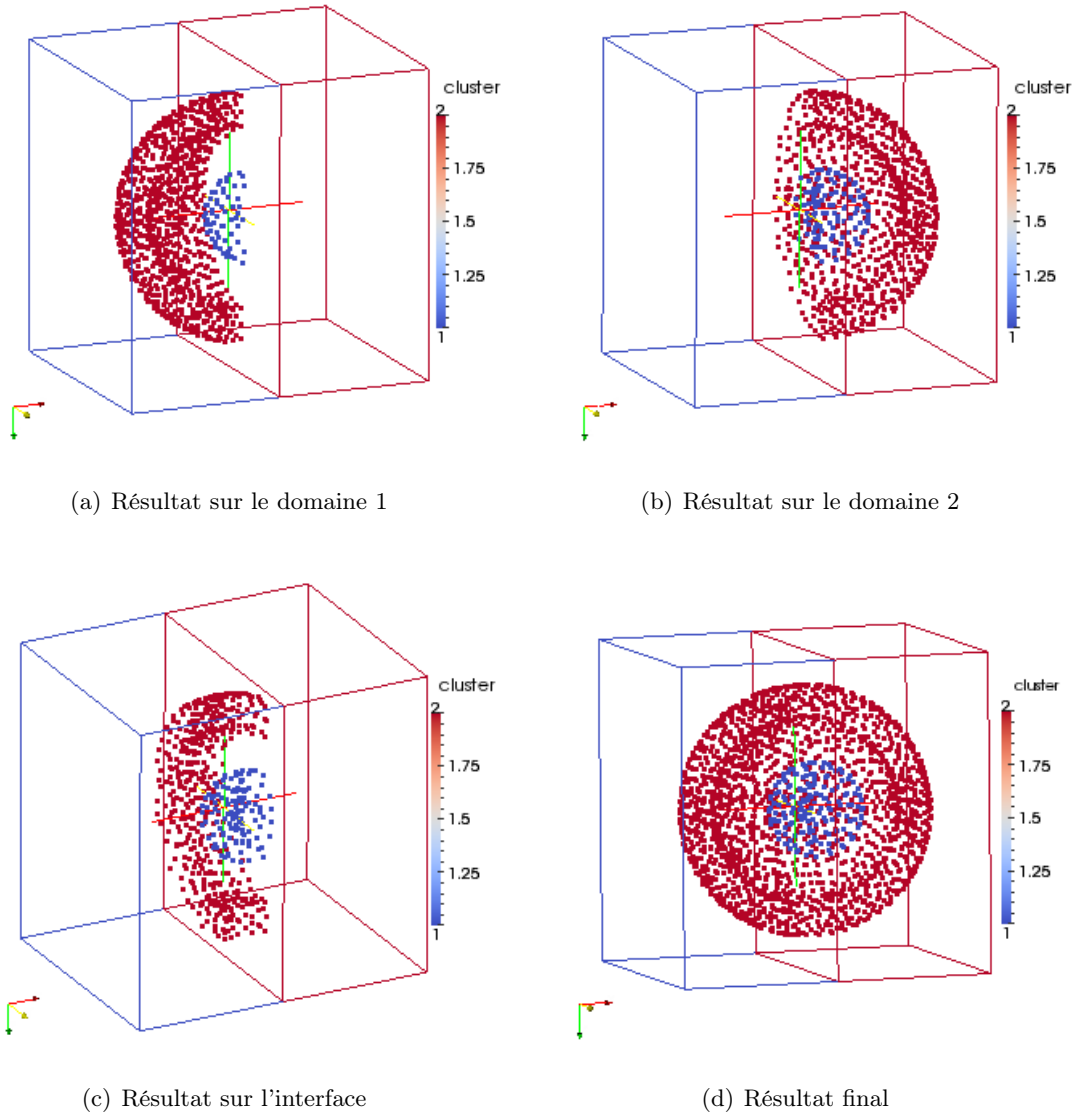


FIGURE 3.7 – Résultats du spectral clustering parallèle avec interface sur l'exemple des 2 sphères découpées en 2 sous-domaines

observations :

- la stratégie parallèle est plus performante que de considérer un seul processeur ;
- la majeure partie du temps est consacrée au spectral clustering sur les sous-domaines et l'interface ;
- l'interface concentre le maximum de points pour $q = 4$ et $q = 8$;

D'après la figure 3.8, La stratégie de parallélisation reste performante mais son efficacité diminue au fur et à mesure que l'on augmente la décomposition en sous-domaines.

Un autre exemple présentant une distribution des points moins uniforme est considéré. En effet, le clustering spectral en parallèle est testé sur un cas 3D géométrique représentant deux sphères non-concentriques tronquées incluses dans une troisième comme le montre la figure 3.9. Sur la droite,

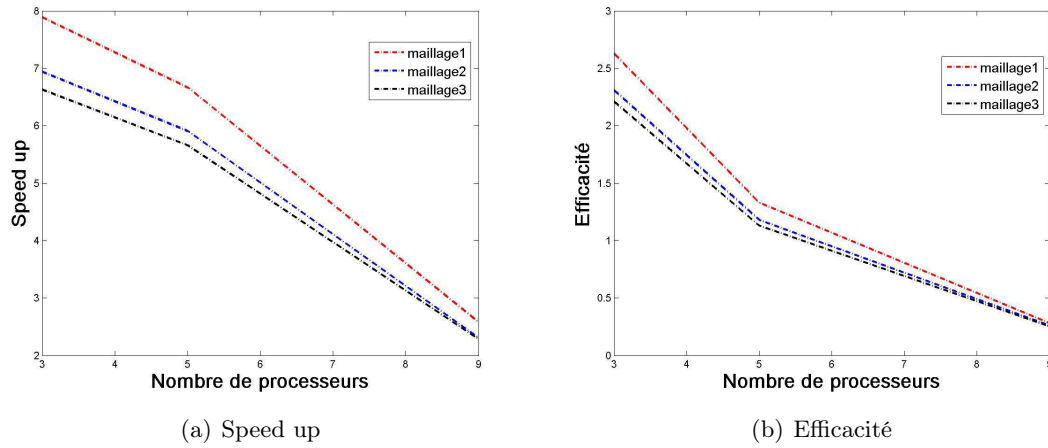
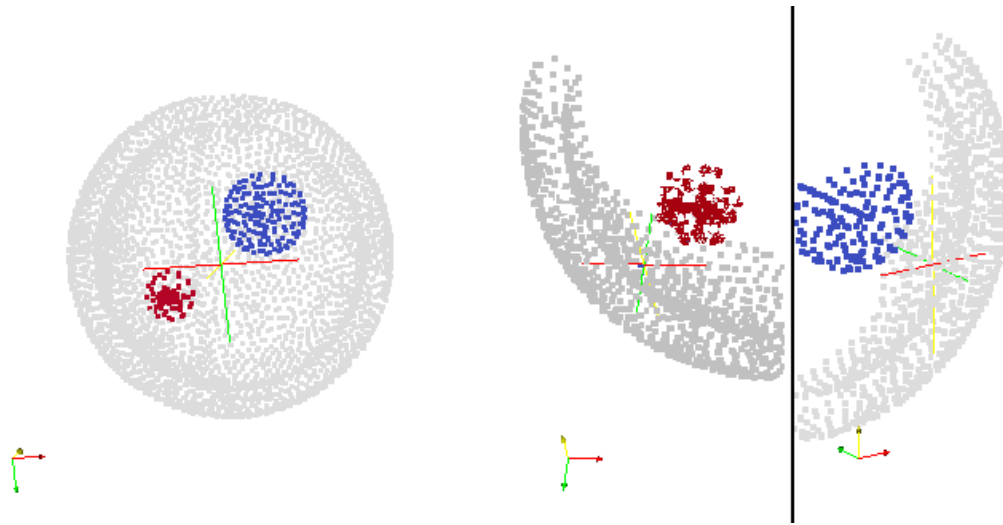


FIGURE 3.8 – 2 sphères tronquées : Speed up et efficacité

figurent des zooms autour de chaque sphère tronquée pour montrer la proximité entre les deux petites sphères et la sphère englobante. La figure 3.10 présente le résultat du spectral clustering sur

FIGURE 3.9 – Exemple géométrique (gauche) et zooms (droite) : $N = 4361$

les deux sous-domaines (a)-(b), celui sur l'interface (c) et le résultat final (d). Chaque sous-domaine comprend 2 clusters et il en est de même pour l'interface.

L'exemple géométrique est testé pour un nombre de points allant de $N = 4361$ à $N = 15247$. Pour ces tests, le domaine est divisé successivement en $q = \{1, 4, 8, 12\}$ boîtes. Les temps pour chaque étape du clustering spectral parallèle sont mesurés. Sur le tableau 3.2, le temps total et le pourcentage de temps consacré à l'algorithme de clustering spectral sur les sous-domaines sont notés pour chaque taille de problème et pour chaque distribution de points sur l'interface.

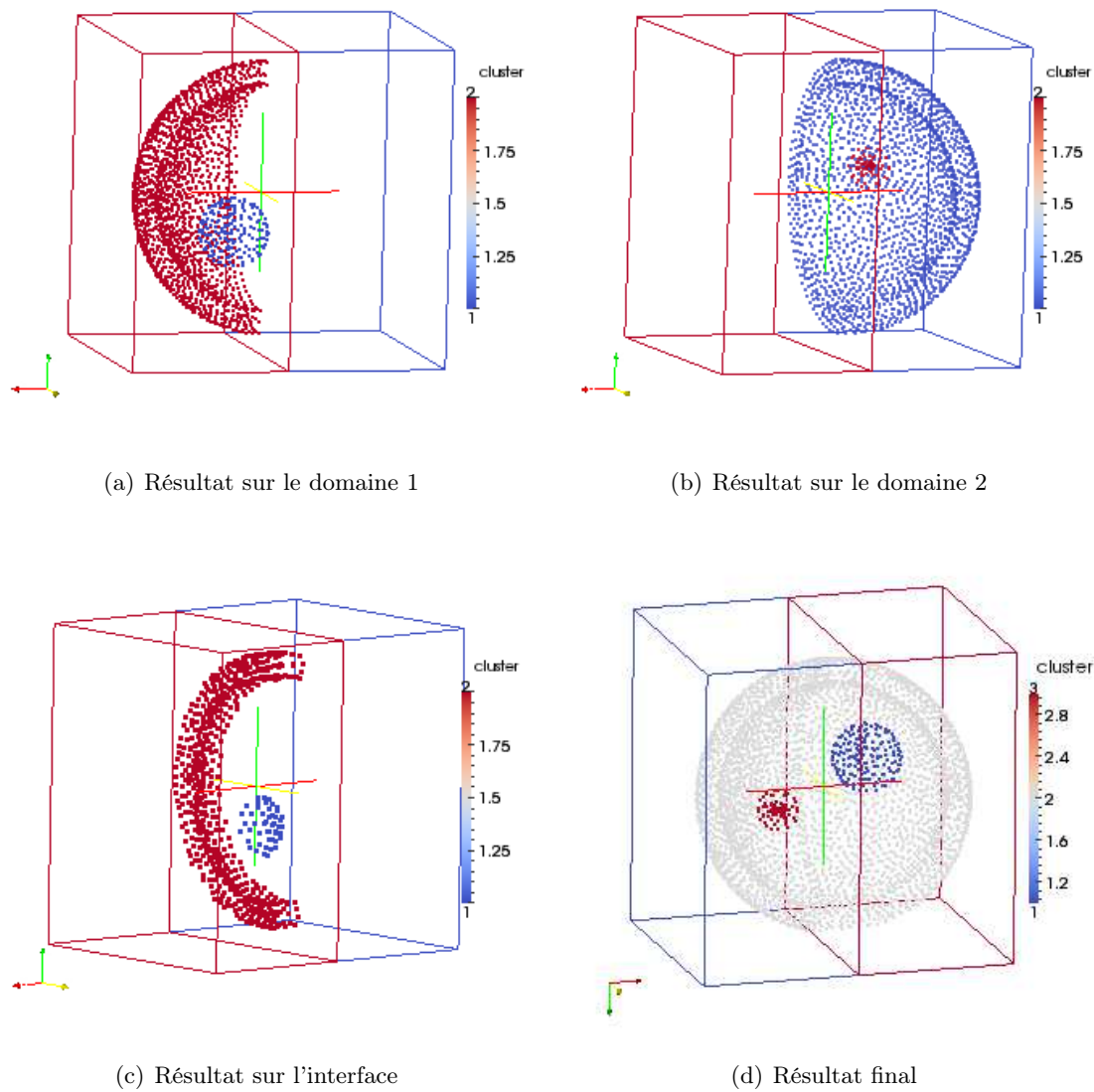


FIGURE 3.10 – Résultats du spectral clustering parallèle sur l'exemple des 3 sphères découpées en 2 sous-domaines

N	Nombre de processeurs	Nombre maximal de données par processeur	Temps total (sec)	% du temps total pour le clustering spectral
4361	1	-	251.12	99.9
	5	1596	13.19	97.4
	9	2131 (int.)	30.6	98.3
	13	2559 (int.)	54.36	98.8
9700	1	-	2716.34	99.9
	5	3601 (int.)	156.04	98.4
	9	4868 (int.)	357.42	99.4
	13	5738 (int.)	610.89	99.7
15247	1	-	11348.43	-
	3	5532 (int.)	549.82	99.5
	9	7531 (int.)	1259.86	99.8
	13	8950 (int.)	2177.16	99.8

TABLE 3.2 – 3 sphères tronquées : cas avec interface

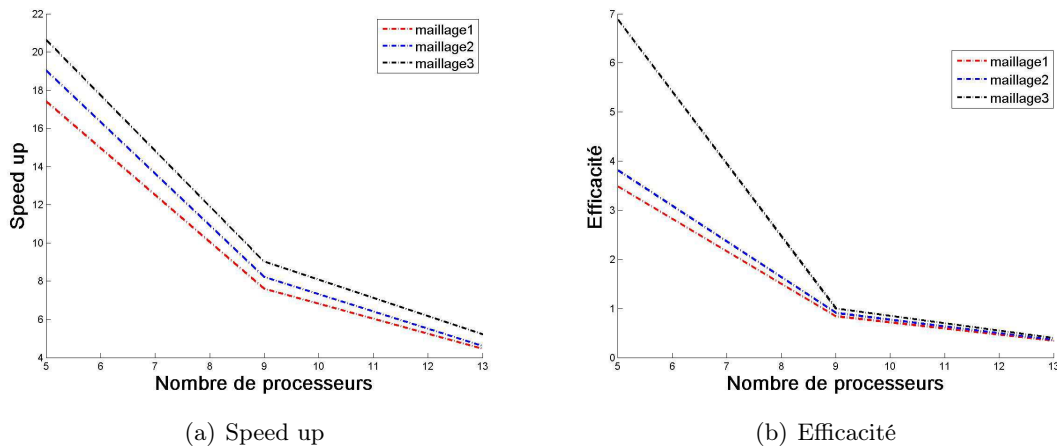


FIGURE 3.11 – 3 sphères tronquées : Speed up et efficacité

De ces résultats, nous retenons les informations suivantes :

- la majeure partie de notre algorithme est consacrée au clustering spectral sur les sous-domaines ;
- le temps consommé pour cette partie correspond au temps du processeur contenant le plus de points : il y a un point de synchronisation à la fin de chaque partie, avant l'étape de regroupement ;
- avec cet exemple, l'interface a le maximum de points ;
- le speed-up est plus grand que le ratio entre le nombre total de points et le nombre maximal de points pour un sous-domaine. Par exemple, avec $N = 4361$ points et 5 processeurs, le ratio est de 2.75 et le speed-up est de 14.12. Ceci peut être expliqué par la non-linéarité du problème lié au calcul des vecteurs propres de la matrice affinité Gaussienne ;
- le clustering spectral par sous-domaines est plus rapide que de considérer l'ensemble entier

S des données. Le calcul des paramètres σ , σ^* et l'étape de regroupement ne pénalisent pas notre stratégie ; le temps consacré à ces parties est négligeable (inférieur à 2% du temps total).

Discussion

La boucle implémentée pour tester les différentes valeurs du nombre de clusters k dans l'algorithme de clustering spectral jusqu'à satisfaire (3.3) devient de moins en moins coûteuse quand le nombre de processus augmente. Passer en revue les valeurs de k revient à concaténer des vecteurs propres, appliquer la méthode k -means, réordonnancer et calculer le ratio η_k . Le calcul des vecteurs propres devient moins coûteux à mesure que l'on réduit la taille des matrices affinités. Aussi, diviser l'ensemble des données S revient implicitement à réduire la matrice d'affinité gaussienne à des sous-blocs diagonaux (après permutation).

Remarque 3.11. *D'autres méthodes d'extraction de valeurs propres et de vecteurs propres associés peuvent être utilisées pour réduire le coût numérique. En effet, une routine classique classique de LAPACK [6] est utilisée pour assurer ce calcul. Les méthodes de Lanczos et d'Arnoldi présentées par [92, 38] seraient certainement efficaces. Une autre technique peut être envisagée : celle de seuiller les très faibles affinités (de l'ordre de la précision machine) pour creuser la matrice affinité et réduire le nombre d'opérations.*

Quand l'ensemble des données est divisé en un grand nombre de sous-domaines, l'interface concentre le maximum des données par sous-domaines et devient le processus le plus coûteux. Utiliser une interface qui connecte toutes les partitions locales peut présenter des limites : plus le domaine est divisée en sous-domaines, plus le volume de l'interface, fonction du nombre de découpages augmente et concentre de points. Il faut donc trouver un compromis entre le découpage et la taille de l'interface. Pour limiter cet inconvénient, on définit un seuil, noté τ , représentant le ratio entre le volume couvert par l'interface et le volume total balayé par les données de S . Ce seuil sera donc fonction du nombre de découpages et des longueurs maximales sur chaque dimension :

$$\tau = \frac{Vol(interface)}{Vol} \quad (3.9)$$

où $Vol(interface)$ est défini par (3.7) et Vol est le volume total couvert par l'ensemble des données. Vol est fonction de l_i défini par (3.6) pour $i = \{1, \dots, p\}$: $Vol = \prod_{i=1}^p l_i$.

Ce seuil permet de limiter les découpages et d'équilibrer le temps de calcul par processus.

3.3 Classification spectrale parallèle avec recouvrement

La stratégie obtenue en séparant l'ensemble des données permettant la connexion entre les partitions issues des différents sous-domaines présente des limites quand le cardinal de l'ensemble interface augmente. Dans la suite, nous proposons de distribuer cet ensemble de points dans les sous-domaines associés. Après une présentation de la stratégie de parallélisation par recouvrement, nous appliquerons les mêmes tests que pour la stratégie avec interface afin de comparer les deux méthodes.

3.3.1 Principe

Afin de pallier l'inconvénient de considérer une interface comme un sous-domaine distinct, l'ensemble des données de l'interface peut être inclus dans les autres sous-domaines. En fait, l'ensemble des données est divisé en q boîtes qui ont une intersection non-vide entre elles. Ainsi, le nombre de

processeurs est réduit ($nbproc = q$) et le calcul du paramètre d'affinité spécifique σ^* pour l'interface est supprimé. Le principal avantage est que la méthode de clustering spectral est utilisée sur tous les sous-domaines avec la même topologie de volume ce qui ne brise pas la distribution isotropique.

Remarque 3.12. *Le seuil τ défini par l'équation (3.9) est préservé pour réduire le temps de l'étape de regroupement et maintenir une cohérence entre le nombre de découpes et l'ensemble des données S . Donc, le volume de l'intersection entre les sous-domaines est majoré par une fraction du volume total balayé par S .*

Le principe de la stratégie avec recouvrement le long des frontières de découpes est représenté sur la figure 3.12 pour le cas $q = 2$.

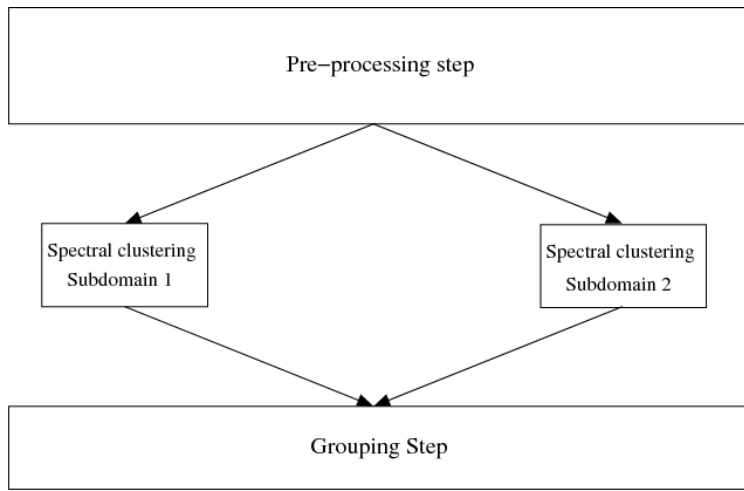


FIGURE 3.12 – Principe de l'alternative clustering spectral parallèle avec recouvrement pour $q = 2$

Etape de prétraitement : partitionner S en q sous-domaines

Incluons toutes les données dans une boîte de de côté l_i pour la i ème dimension, $i = \{1, \dots, p\}$ où :

$$l_i = \max_{1 \leq i_1, i_2 \leq N} |x_{i_1}(i) - x_{i_2}(i)|, \forall i \in \{1, \dots, p\}. \quad (3.10)$$

En prenant la longueur maximale sur chaque dimension, la boîte est divisée en q sous-boîtes où $q = \prod_{i=1}^p q_i$ et q_i correspond au nombre de subdivisions sur la i -ème dimension. Le nombre de processus, noté $nbproc$, est égal à $nbproc = q$.

Décomposition par domaine : sous-domaines

Chaque processus de 1 à $nbproc$ a un sous-ensemble de données S_i , $i = 1 \dots nbproc$ dont les coordonnées sont incluses dans une sous-boîte géométrique. Le paramètre d'affinité gaussienne est défini par l'équation suivante (1.2) à partir de l'ensemble initial S de données. La largeur de la bande γ de recouvrement le long des frontières des sous-domaines est égale à la largeur de l'interface c'est-à-dire $\gamma = 3\sigma$ où σ satisfait l'équation (1.2). Par exemple pour la i^{me} dimension divisée en q sous-domaines, la taille du i^{me} côté d'un pavé est alors égale à $l_i/q + \gamma$.

Spectral Clustering sur les sous-domaines

Le même algorithme de spectral clustering utilisé pour la méthode parallèle avec interface est appliqué sur les sous-domaines. De plus, les détails liés à la parallélisation sont les mêmes.

Etape de regroupement

La partition finale est formée du regroupement des partitions issues des $nbproc$ partitions du clustering spectral appliqué aux $nbprocs$ sous-ensembles de S . Le regroupement est réalisé à l'aide des résultats du spectral clustering sur la zone de recouvrement entre les sous-domaines et de la relation d'équivalence (3.3). Si un point appartient à deux clusters différents, ces deux clusters sont assemblés en un plus grand. En sortie de la méthode parallèle, une partition de l'ensemble S des données et un nombre final de clusters k sont donnés.

De la même façon que pour l'interface, la stratégie avec recouvrement est testée sur le même exemple et représentée sur la figure 3.13. La cible est donc divisée en $q = 4$ sous domaines. Sur la gauche, la partition finale, après l'étape de regroupement, est affichée. Sur la droite, les résultats du clustering spectral sur les 4 sous-domaines sont représentés.

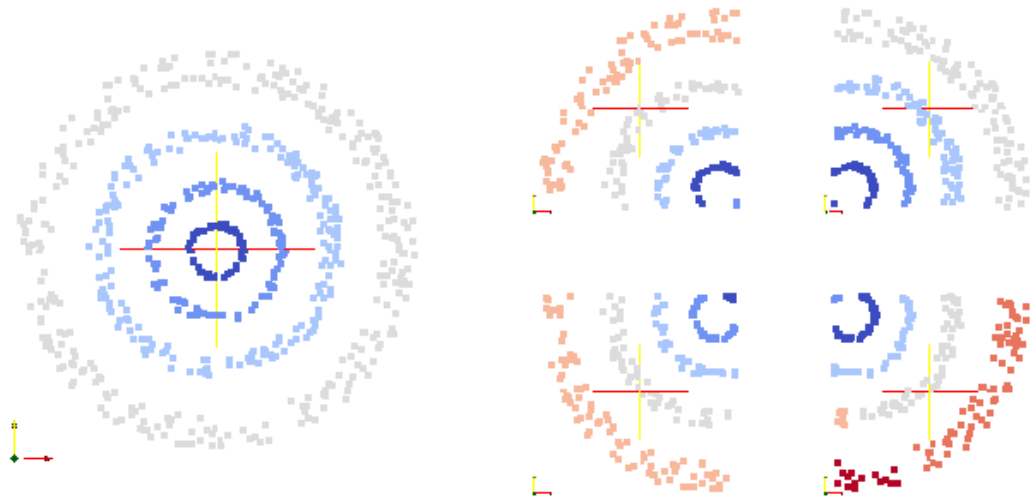


FIGURE 3.13 – Exemple cible : Résultat du spectral clustering avec recouvrement après étape de regroupement et par sous-domaines

3.3.2 Justification de la cohérence de la méthode

Dans cette section, la cohérence de la méthode précédemment introduite est étudiée. On va procéder comme dans la section 3.2.2. Notons par $\hat{S}_k$ l'ensemble des points pour le sous-domaine k , $k \in [1, q]$. L'objectif est de montrer, à nouveau, qu'à partir d'un clustering idéal de l'ensemble des points $\hat{S}_k$ du sous-domaine k , l'étape de regroupement permet de reconstituer les ensembles $S_i = S \cap \Omega_i$. La difficulté est de montrer que le clustering sur les pavés C_k , $k \in [1, q]$ définis par la stratégie par découpe avec recouvrement permet d'obtenir les clusters ω_i^k au sens de l'hypothèse du

théorème 3.3 de sorte à définir le clustering sur les points S au sens du théorème 3.5. La difficulté est encore pour pouvoir appliquer le théorème 3.5 de satisfaire l'hypothèse que si $\omega_i^k \cap \omega_j^l \neq \emptyset$ pour $i \in [1, q]$ alors il existe un point $x \in S_i^k \cap S_j^l$. Toutefois, le clustering est réalisé sur les points de S et donc il existe une infinité de façon de définir des pavés fermés C'_k de manière à satisfaire $\hat{S}_k = S \cap C_k = S \cap C'_k$. L'objectif de cette partie est de définir un ensemble de pavés C'_k "convenables" pour $k \in [1, q]$.

Soient $C_1, \dots, C_q$ les q sous-domaines qui permettent le partitionnement géométrique de l'ensemble des données S et γ représente l'épaisseur du recouvrement pour chaque pavé C_j . On rappelle que $(\omega_i^k)_{k=1..n_i}$ sont les clusters issus du pavé C_i . Soient maintenant les ensembles C'_k pour $k \in [1, q]$ construits de la manière suivante : on commence avec $C'_i = C_i$ pour $i \in [1, q]$ puis on applique la règle suivante représentée sur la figure 3.6 (b),

- si $\omega_i^k \cap \omega_j^l \neq \emptyset$ et $\omega_i^k \cap \omega_j^l \cap S = \emptyset$ pour $i \neq j$ alors on pose $C'_i := C'_i \setminus \Gamma_i^\epsilon$ où $\Gamma_i^\epsilon = \omega_i^k \cap \{x \in \Omega, d(x, \partial\omega_j^l) < \epsilon\}$ pour $\epsilon > 0$ et tel que $\epsilon < d_c$ et $\Gamma_i^\epsilon \cap S = \emptyset$. De même, on pose $C'_j := C'_j \setminus \Gamma_j^\epsilon$ pour un ϵ vérifiant les mêmes hypothèses que précédemment.

Lemme 3.13. *Les C'_k pour $k \in [1, q]$ vérifient les propriétés suivantes :*

1. C'_k est fermé,
2. $\Omega \subset \bigcup_{k \in [1, q]} C'_k$,
3. $C'_k \cap S = \hat{S}_k$.

Démonstration. Même démonstration que pour le lemme 3.9. □

On note par $((\omega')_i^k)_{k=1..n_i}$ les clusters issus du pavé C'_i . Par le lemme 3.13, les C'_k forment un ensemble de pavés "convenables" et donc par le théorème 3.3, permettent de retrouver les ω_i^k . On obtient alors le résultat suivant.

Théorème 3.14. *Supposons que le clustering de S est idéal sur chaque $\hat{S}_i$, $i \in [1, q]$ et que de plus, $\gamma > d_c$. Alors pour tout $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$, il existe $i \in \{1, \dots, n\}$ tel que $\Omega_i \cap S = \bigcup_{S_i^k \in \bar{S}_A} S_i^k$.*

Réciproquement, pour tout $i \in \{1, \dots, n\}$, il existe $\bar{S}_A \in \mathcal{S}/\mathcal{R}^S$ tel que $\bigcup_{S_i^k \in \bar{S}_A} S_i^k = \Omega_i \cap S$.

Démonstration. Il suffit de montrer que si $(\omega')_i^k \cap (\omega')_j^l \neq \emptyset$ pour $i \in [1, q]$ alors il existe un point $x \in S_i^k \cap S_j^l$, le résultat se déduira du théorème 3.5. Or par la définition des C'_k et le lemme 3.13, on a $(\omega')_i^k \cap (\omega')_j^l \cap S = \emptyset$ implique $(\omega')_i^k \cap (\omega')_j^l = \emptyset$. □

Donc si, sur chaque sous-domaine, le clustering est idéal alors l'étape de regroupement de la stratégie par interface permet de reconstituer les clusters $\Omega_i \cap S$ via la relation d'équivalence $\mathcal{R}^S$.

Choix du paramètre γ

Dans cette approche, l'ensemble des points permettant de lier les clusters à support sur plusieurs pavés est inclus dans l'ensemble des données de chaque sous-domaine. D'après le théorème 3.14, l'hypothèse $\gamma > d_c$ doit être satisfaite pour regrouper les clusters. Nous faisons l'hypothèse pour la définition des composantes connexes Ω_k qu'aucune d'entre elles ne peut être subdivisée en deux ouverts recouvrant le même sous-ensemble de points de S et à distance supérieure à γ . La largeur de la bande de recouvrement est donc la même que celle de l'interface dans la stratégie de spectral clustering avec intersection c'est-à-dire égale à 3σ où σ représente la distance de séparation entre les points dans le cas d'une distribution uniforme c'est-à-dire lorsque les points sont tous équidistants les uns des autres. Un choix de $\gamma = 3\sigma$ permet donc a priori de pallier les diverses configurations possibles de distributions des points par cluster.

3.3.3 Expérimentations numériques : exemples géométriques

Les mêmes exemples que ceux de la section 3.2.3 sont testés avec cette nouvelle alternative et réalisés sur Hyperion. Les résultats du spectral clustering sur les 2 sous-domaines sur les exemples des 2 sphères tronquées sont respectivement représentés sur les figures 3.14 (a)-(b) et le résultat final sur les figures 3.14 (c).

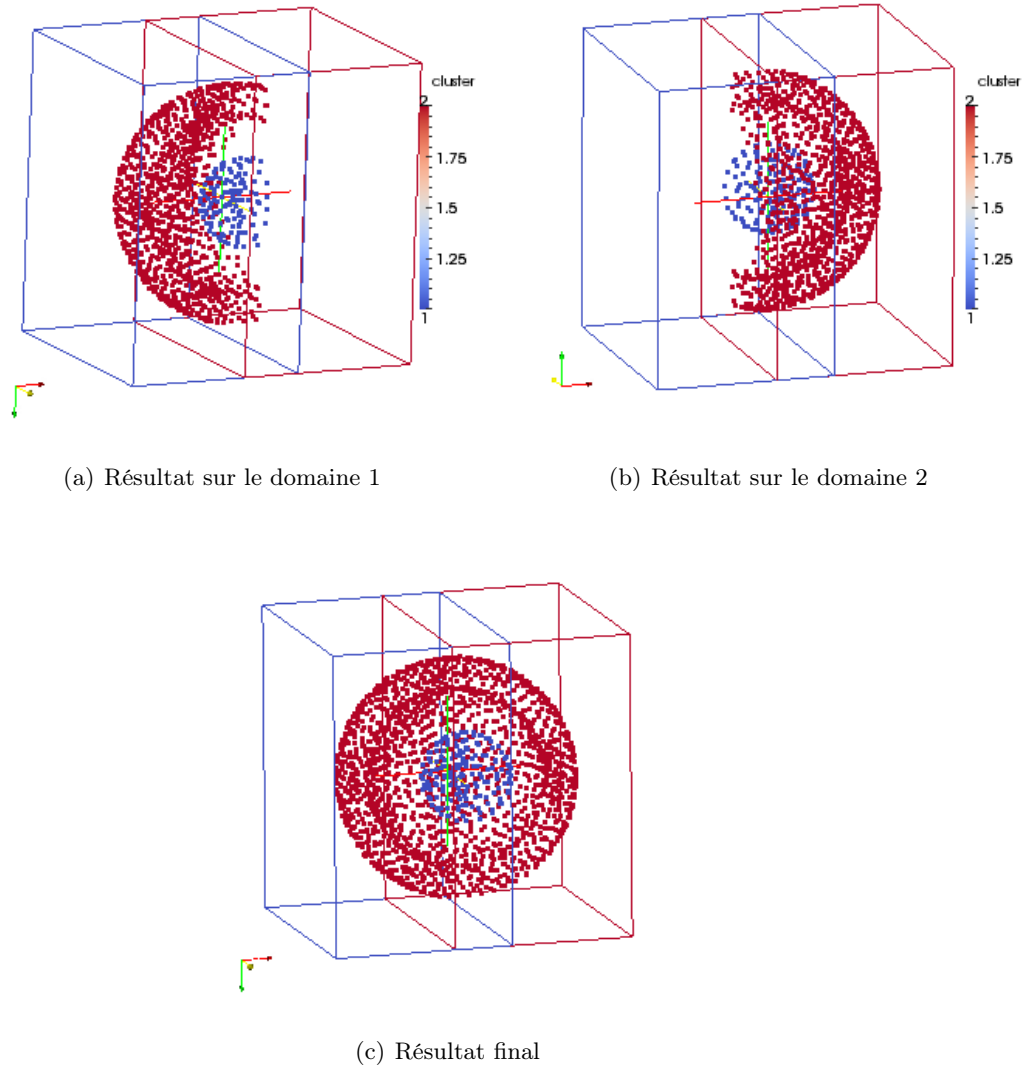
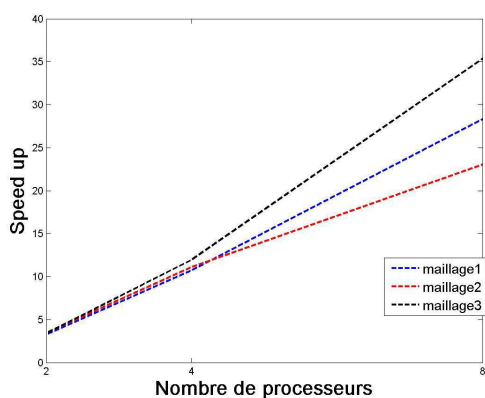


FIGURE 3.14 – Résultats du spectral clustering sur le cas des 2 sphères découpées en 2 clusters

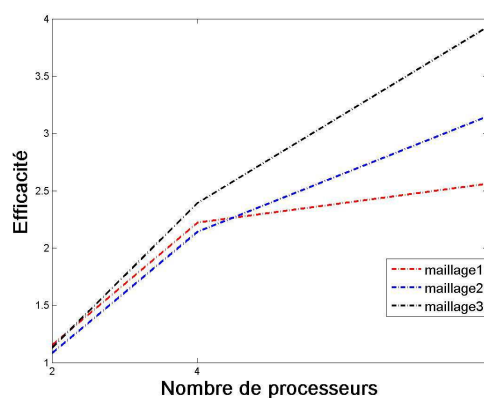
De la même manière, les résultats sont récapitulés, pour des découpages en $q = \{2, 4, 8\}$ sous-domaines, sur les tableaux 3.3 à partir des mesures temporelles sur les étapes respectives du clustering spectral parallèle avec recouvrement. Les graphes de speed-up et d'efficacité sont respectivement représentés sur les figures 3.15 (a) et (b).

N	Nombre de processeurs	Nombre maximal de données par processeur	Temps total (sec)	% du temps total pour le clustering spectral
1905	1	-	21.67	99.58
	2	1265	6.27	98.89
	4	849	1.95	95.90
	8	608	0.94	88.30
3560	1	-	127.27	99.85
	2	2336	39.22	99.59
	4	1548	11.85	98.48
	8	1127	4.49	94.21
10717	1	-	3626.59	99.97
	2	7010	1071.04	99.89
	4	4625	303.29	99.63
	8	3197	102.52	98.73

TABLE 3.3 – 2 sphères tronquées : cas avec recouvrement



(a) Speed up



(b) Efficacité

FIGURE 3.15 – 2 sphères tronquées : Speed up et efficacité

Les résultats du spectral clustering parallèle avec recouvrement sur l'exemple des 3 sphères tronquées sont représentés sur la figure 3.16. On observe que le nombre de clusters déterminés via l'heuristique (3.5) varie d'un sous-domaine à l'autre sur la figure 3.16 (a)-(b). Le sous-domaine 1 définit 2 clusters tandis que le sous-domaine 2 comprend 3 clusters. Pour des découpes en $q = 4, 8, 12$ sous-domaines, le tableau 3.4 récapitulent les mesures temporelles des diverses étapes de la méthode. Les graphes de speed-up et d'efficacité sont tracés sur la figure 3.17.

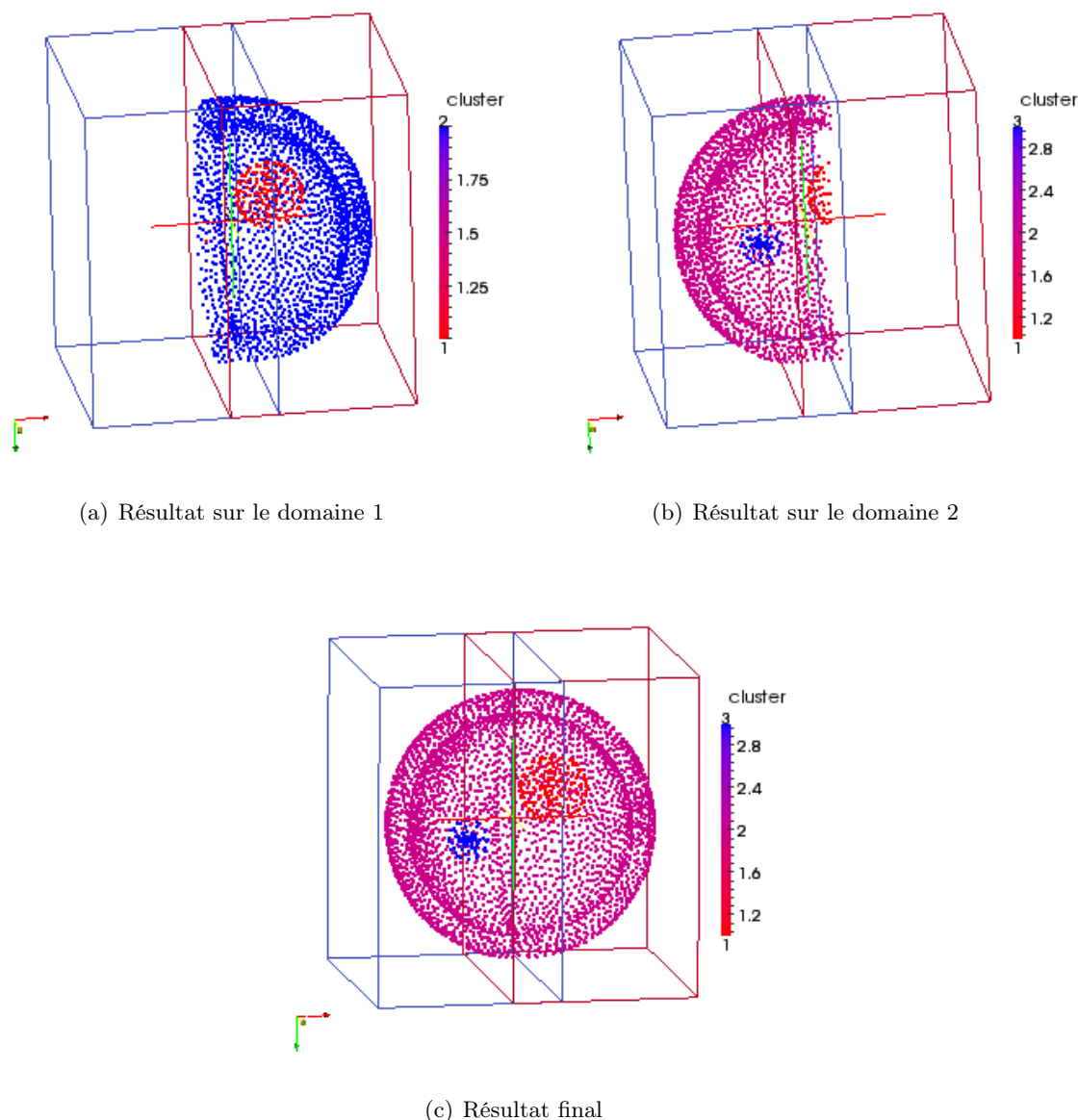


FIGURE 3.16 – Résultats du spectral clustering parallèle sur l'exemple des 3 sphères découpées en 2 sous-domaines

N	Nombre de processeurs	Nombre maximal de données par processeur	Temps total (sec)	% du temps total pour le clustering spectral
4361	1	-	251.12	99.8
	4	1662	14.73	97.2
	8	984	3.37	85.2
	12	1004	3.71	82.5
9700	1	-	2716.34	99.9
	4	3712	157.76	99.4
	8	2265	38.89	97
	12	2283	40.43	96.6
15247	1	-	11348.43	-
	4	5760	578.92	99.6
	8	3531	133.79	98.1
	12	3517	131.83	97.9

TABLE 3.4 – 3 sphères tronquées : cas avec recouvrement

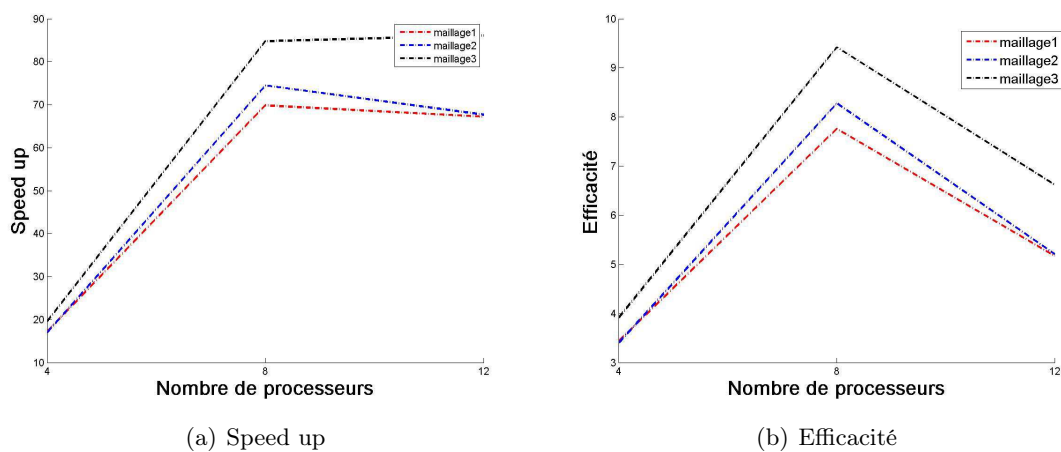


FIGURE 3.17 – 3 sphères tronquées : Speed up et efficacité

A partir de ces résultats, nous observons que cette alternative a le même comportement, les mêmes principales caractéristiques que la stratégie avec interface sur les deux exemples géométriques :

- très bonne performance, plus grande que le ratio entre le nombre total de points et le nombre maximal de points sur un sous-domaines ;
- la majeure partie du temps est consacrée à l'étape de clustering spectral ;
- le temps de l'étape de clustering spectral correspond au temps du processeur avec le maximum de données.

Des remarques spécifiques à cette alternative apparaissent :

- les temps et par conséquent le speed-up et l'efficacité sont nettement meilleurs que ceux de la stratégie avec interface avec un nombre équivalent de processeurs supérieur à 2 : pour l'exemple des 2 sphères, avec $q = 4$ et $N = 3560$ points, le temps total est divisé par 9.3 et pour l'exemple des 3 sphères, avec $q = 12$ et $N = 15247$, le temps total est divisé par 16.5 ;
- le temps est décroissant quand le nombre de subdivisions augmente à condition que le nombre maximal de données sur un processeur décroisse. On observe, par exemple, qu'avec $N = 4361$ points et $q = 12$, le processeur avec le maximum de points a moins de points que son équivalent avec $q = 8$. Ceci explique que le temps maximum pour un processeur soit plus grand pour $q = 12$ que pour $q = 8$. Ceci explique aussi pourquoi la stratégie avec interface plus performante que celle avec recouvrement sur les deux exemples pour un nombre équivalent de sous-domaines égal à 2.

Remarque 3.15. *Cette dernière remarque ouvre une réflexion sur comment diviser le domaine : un découpage qui équilibrerait le nombre de points au sein des processeurs donnerait de meilleurs résultats qu'un découpage géométrique automatique. Cependant, il faut veiller à ne pas trop localiser l'analyse et à garder une cohérence avec les résultats théoriques développés au chapitre précédent. Rappelons que ces données projetées dans l'espace spectral sont la discrétisation de fonctions propres à support sur une seule composante connexe.*

3.4 Segmentation d'images

Des exemples de segmentation d'images en niveaux de gris sont considérés. La parallélisation donne lieu à approfondir les raisonnements développés au chapitre 1. En effet, cette stratégie de spectral clustering avec recouvrement nous permet d'aborder des images de grandes tailles et étudier une définition de l'affinité alliant à la fois les informations géométriques et les informations de luminance. Ce genre d'exemples est bien formaté pour la stratégie en parallèle grâce à sa distribution uniforme en coordonnées géométriques. Dans la suite, toutes les expérimentations numériques sont testées sur Hyperion.

3.4.1 Boîte affinité 3D rectangulaire

L'approche par boîte 3D rectangulaire consiste à définir l'affinité entre deux pixel i et j par :

$$A_{ij} = \exp\left(-\frac{\|x_i - x_j\|_2}{\sigma^2/2}\right),$$

où les données $x_i \in \mathbb{R}^3$ regroupant les coordonnées géométriques et la luminance. Un problème se pose lorsque l'on assimile deux données de natures différentes. En effet, les niveaux de gris sont des entiers appartenant à l'intervalle $[0, 255]$ alors que les coordonnées géométriques dépendent de la dimension $p_1 \times p_2$ de l'image considérée. L'approche 3D doit pondérer l'influence des coordonnées

géométriques par rapport à la luminance. Un procédé classique consiste à normaliser les coordonnées géométriques sur chaque dimension : les pixels (i, j) et $(i + 1, j)$ sont séparés de $\frac{1}{p_1}$ et les pixels (i, j) et $(i, j+1)$ séparés de $\frac{1}{p_2}$.

Remarque 3.16. Dans la théorie développée au chapitre 2, cette normalisation des données revient à considérer une fonction composée du noyau de la chaleur et d'une autre fonction de normalisation. Le produit de convolution avec un produit de fonctions composées revient à considérer un produit tensoriel de l'opérateur de la chaleur appliqué à une fonction propre et d'une fonction à définir. La propriété de clustering sur les fonctions propres de l'opérateur de la chaleur est conservée mais est pondérée par une fonction. Cette normalisation des données ouvre une nouvelle étude à réaliser sur les résultats théoriques précédents et l'influence sur le paramètre.

On considère l'image entière (non seuillée) du bouquet de fleurs donc les exemples de segmentation d'images du chapitre 1 sont extraits. Pour ce genre d'exemple, étant donnée l'uniformité de la distribution des données comme le montre la figure 3.19, on peut se permettre de définir une zone de recouvrement relativement faible. Il s'agit d'une image de 186×230 pixels, à savoir $N = 42780$ points. Le clustering spectral parallèle avec recouvrement est appliqué sur $q = 20$ processeurs. Notons qu'avec $q < 20$, cet exemple ne fonctionne pas en raison d'un manque de mémoire virtuelle.

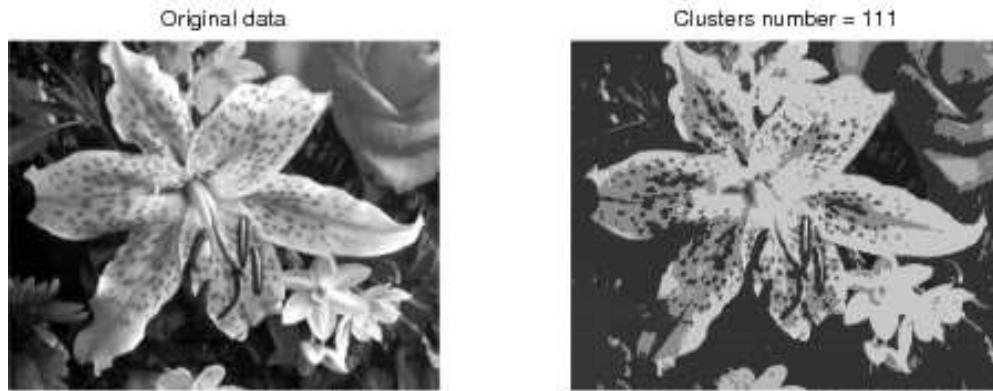


FIGURE 3.18 – Exemple de segmentation d'image testée sur Hyperion

Sur la figure 3.18, les données initiales sont représentées sur la gauche et le clustering final sur la droite. Pour comparer la partition issue du spectral clustering parallèle, à partir du cardinal de chaque cluster et de l'assignation des points, tous les points associés au même cluster se voient affectés un niveau de gris égal à la moyenne des niveaux de gris des pixels correspondant à l'image originale pour chacun de ces points. Ce procédé permet de comparer visuellement les clusters trouvés en fonction de l'image originale. En effet, si les points assignés au même cluster correspondent à différentes couleurs et différentes formes alors ce cluster aura une couleur et une forme différentes de l'image originale. L'algorithme de spectral clustering pour cette image de bouquet de fleurs a déterminé 111 clusters. Comparées à l'image originale, les formes des fleurs sont bien décrites et les détails du lys central sont distingués. Le temps total est de 398.78 secondes pour $N = 42780$ points, ce qui confirme la performance de cette stratégie parallèle avec recouvrement. Etant données les dimensions et la complexité de l'image, l'impact de la normalisation des coordonnées géométriques ne pénalise pas le caractère de compression de données que l'on peut déduire du résultat du spectral clustering. En effet, considérer un caractère géométrique aux données implique la définition de nombreux clusters aux mêmes niveaux de gris mais situés géométriquement à des endroits différents

de l'image. Ici, avec 111 clusters, l'image du bouquet est restitué tout en évitant 256 clusters de niveaux de gris. Pour expliquer le processus du spectral clustering parallèle avec recouvrement sur une image, la découpe est représentée sur la figure 3.19 où l'on représente la découpe géométrique en 4 sous-domaines sur la figure (a) puis la découpe en 5 sous-domaines suivant le niveau de luminance comme le montre la figure (b). Ainsi le résultat du spectral clustering est représenté sur la figure 3.20 pour chaque découpe de niveau de luminance pour le quart de l'image colorée en dégradés de bleu sur la figure 3.19 (a). La couleur marron indique qu'il n'y a pas de donnée pour ce niveau de luminance en ces coordonnées géométriques. Ainsi, on observe que pour chaque découpe de niveau de gris, les clusters regroupent les points d'une ou plusieurs fleurs géométriquement proche et de même niveau de gris. La superposition des clusters sur les différentes découpes est représentée sur la figure 3.20 (f).

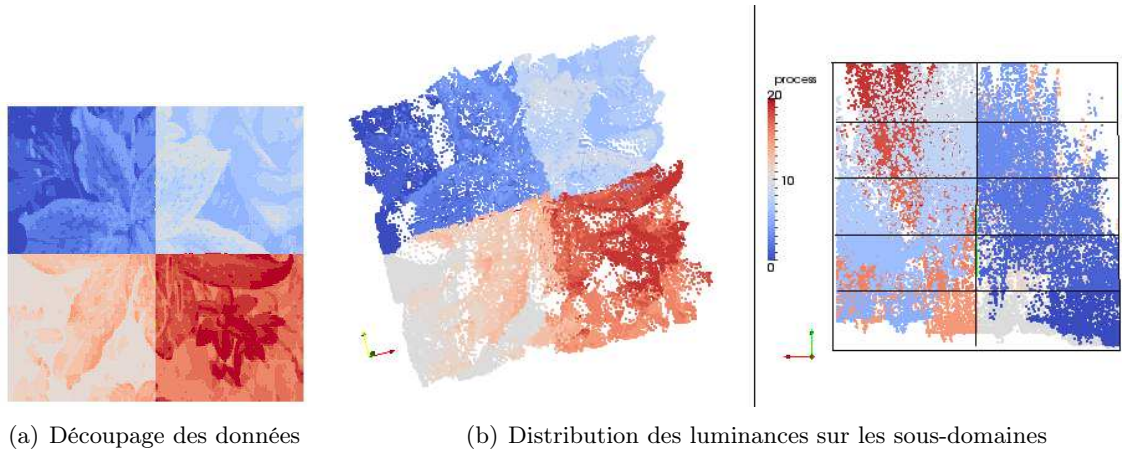
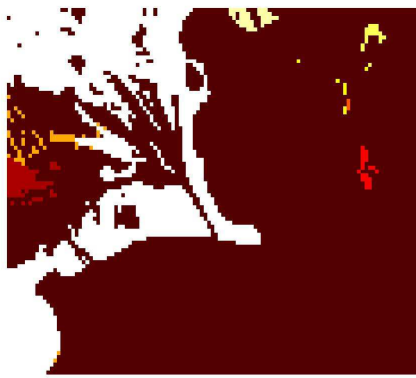


FIGURE 3.19 – Ensemble des données 3D de l'image bouquet

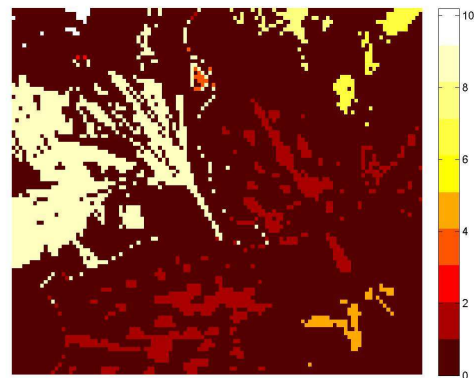
D'autres exemples sont testés et représentés sur la figure 3.21 : sur la gauche figure l'image originale et sur la droite le résultat du clustering spectral parallèle. Ces nouvelles images sont de dimensions plus grandes, présentent des difficultés différentes :

- l'image 200×320 du visage d'un clown comprend $N = 64000$ points. Après l'étape de regroupement, 203 clusters ont été définis en 3110.38 secondes. Cette image présente une distribution des points plus contrastés notamment au niveau des cheveux (cf figure 3.21) d'où un temps total important.
- une photo 537×358 pixels (soit $N = 192246$ points) représentant un croissant de lune à la surface complexe. En effet, le relief et le niveau de gris ne sont pas uniformes. 1860 clusters ont été déterminés en 6842.5 secondes.

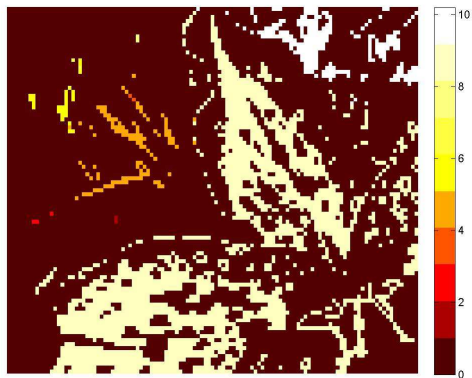
Ces deux exemples indiquent que lorsque la complexité de l'image augmente, l'information géométrique influe et crée des clusters géométriquement distincts mais avec des niveaux de gris identiques. Sur l'image du clown, seuls 3 niveaux de gris en moyenne sont représentatifs de 2 ou 3 clusters. L'information reste néanmoins compressible étant donné que nous pouvons reconstituer l'image avec 203 clusters. Le croissant de lune par contre représente un cas limite où la complexité de la surface alliant relief (donc informations géométriques différentes) et surfaces aux niveaux de gris différents permet uniquement de reconnaître la forme du croissant de lune. En augmentant le nombre de clusters limites par sous-domaine à trouver à 100 ($nblimit = 100$), 1860 clusters sont déterminés.



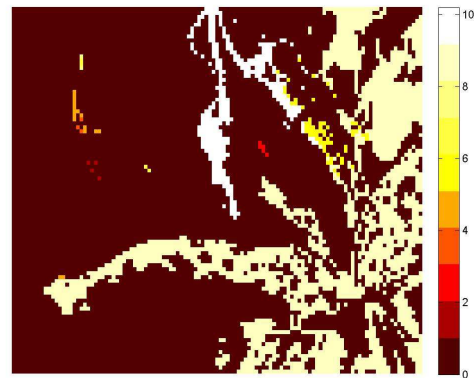
(a) Résultat sous-domaine 1



(b) Résultat sous-domaine 2



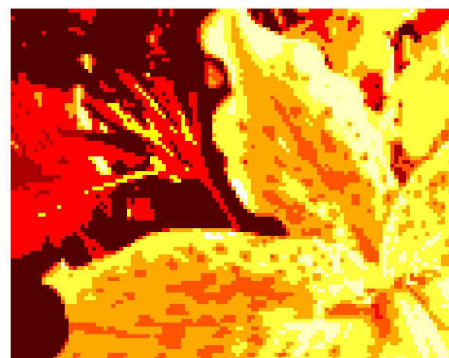
(c) Résultat sous-domaine 3



(d) Résultat sous-domaine 4



(e) Résultat sous-domaine 5



(f) Résultat final

FIGURE 3.20 – Résultats du spectral clustering sur le bouquet

Parmi ces clusters, de nombreux clusters sont géométriquement distants mais ont le même niveau de gris. De plus, l'image n'est pas correctement restituée pour cette valeur de *nblimit*.



(a) Clown : 64000points



(b) Croissant de lune : $N = 192246$ points

FIGURE 3.21 – Autres exemples de segmentation d'images testées sur Hyperion

Enfin, un autre test est effectué sur une planche de bande dessinée fabriquée uniquement sur ordinateur, à l'aide d'un crayon graphique avec une tablette graphique et de logiciels de retouche d'image. Cette image retouchée crée des dégradés continus de niveau de gris rendant difficile la segmentation d'image. Un exemple de mahua réalisé par Zhang Lin est représenté à gauche sur la figure 3.22 auquel on teste notre méthode. Cette image est de taille 476×489 soit $N = 232764$ points et est divisée en 20 sous-domaines et le nombre limite de clusters à tester par vignette est égal à 50, 100 et 200 (*nblimit* = 200). 2475 clusters ont été déterminés pour *nblimit* = 200 pour un temps total de 8803.27 secondes. Le résultat du spectral clustering pour les diverses valeurs de *nblimit* est représenté à droite sur la figure 3.22. Ainsi une évolution des clusters peut être observée au fur et à mesure que le nombre limite de clusters à trouver par sous-domaine augmente. Le nombre de clusters

Images nblimit	Arthus1 100	Arthus1 200	Arthus1 500	Arthus2 100	Arthus2 200	Arthus2 500
Nb $\Lambda \leq 10$	255	706	3549	445	1161	4751
Nb $10 < \Lambda \leq 10^2$	550	1527	3716	455	1187	2760
Nb $10^2 < \Lambda \leq 10^3$	263	236	94	122	138	50
Nb $10^3 < \Lambda \leq 10^4$	4	6	6	4	3	7
Nb $\Lambda > 10^4$	1	1	0	1	3	0

TABLE 3.5 – Nombres de clusters suivant leur cardinal

suivant leur cardinal peut être évalué. Pour $nblimit = 50$, 88 clusters avaient entre 10 et 100 points et 2 clusters comprenaient plus de 10^5 points, à savoir 136050 et 11671 points. Pour $nblimit = 200$, 1290 clusters ont entre 10 et 100 points et 2 clusters comprennent plus de 10^4 points, c'est-à-dire 22835 et 11176 points. En augmentant le nombre de clusters limite, le nombre de clusters final a augmenté et la restitution de l'image s'est affinée. Une étude doit être menée afin de déterminer une limite entre la compression des éléments essentiels d'une image et la restitution complète d'une image. De plus, l'influence de l'information géométrique doit être étudiée en particulier d'un point de vue théorique.

3.4.2 Boîte affinité 5D rectangulaire

La même démarche que pour les images en niveaux de gris est adoptée pour les images couleurs. L'approche boîte 5D rectangulaire consiste donc à définir l'affinité entre deux pixels i et j par l'équation 1.1 où les données $x_i \in \mathbb{R}^5$ regroupant les coordonnées géométriques et les niveaux de couleurs en Rouge, Vert et Bleu.

Deux photos de Yann Arthus Bertrand balayant différentes couleurs sont testées et représentées sur les figures 3.24 et 3.23 :

- La première, de taille 407×316 pixels, représentée sur la figure 3.23 gauche, est un paysage de champs colorés près de Sarraud dans le Vaucluse. L'image contient donc des découpes géométriques mais aussi des arbres et des ombres qui joignent les différents champs.
- La seconde, de taille 480×323 pixels, représente le Geyser du grand Prismatic, dans le parc national de Yellowstone dans le Wyoming aux Etats-Unis. Cette photo présente sur la figure 3.24 gauche est complexe car de nombreux dégradés de couleurs sont présents et les formes de ce bassin thermal s'apparentent à celles d'un soleil donc de géométrie complexe.

A nouveau, pour visualiser l'évolution des résultats du spectral clustering en fonction de $nblimit$, le résultat du spectral clustering est affiché sur la droite des figures 3.23 et 3.24 pour des valeurs de $nblimit$ successivement égales à 100, 200 et 500. On constate qu'un peu moins de 7500 clusters sont nécessaires pour discerner des zones aux couleurs voire à la géométrie complexes. Le tableau 3.5 indique le nombre de clusters Λ suivant leur cardinal et suivant la valeur du $nblimit$ pour les 2 photos. A nouveau, on observe que plus $nblimit$ augmente, plus les clusters comportant un très ensemble de données se décompose en clusters de cardinal plus petit. Ainsi l'image reste potentiellement compressible puisqu'elle réduit 256^3 couleurs en 7500 couleurs localisées.

Concernant la performance de la stratégie parallèle avec recouvrement, l'augmentation du $nblimit$ influe sur la durée totale de la méthode. Pour $nblimit = 100$, le temps total pour les figures 3.23 et 3.24 est respectivement égal à 2264.7s et 1405.53s. Pour $nblimit = 500$, ces temps sont égaux



(a) nblimit=50



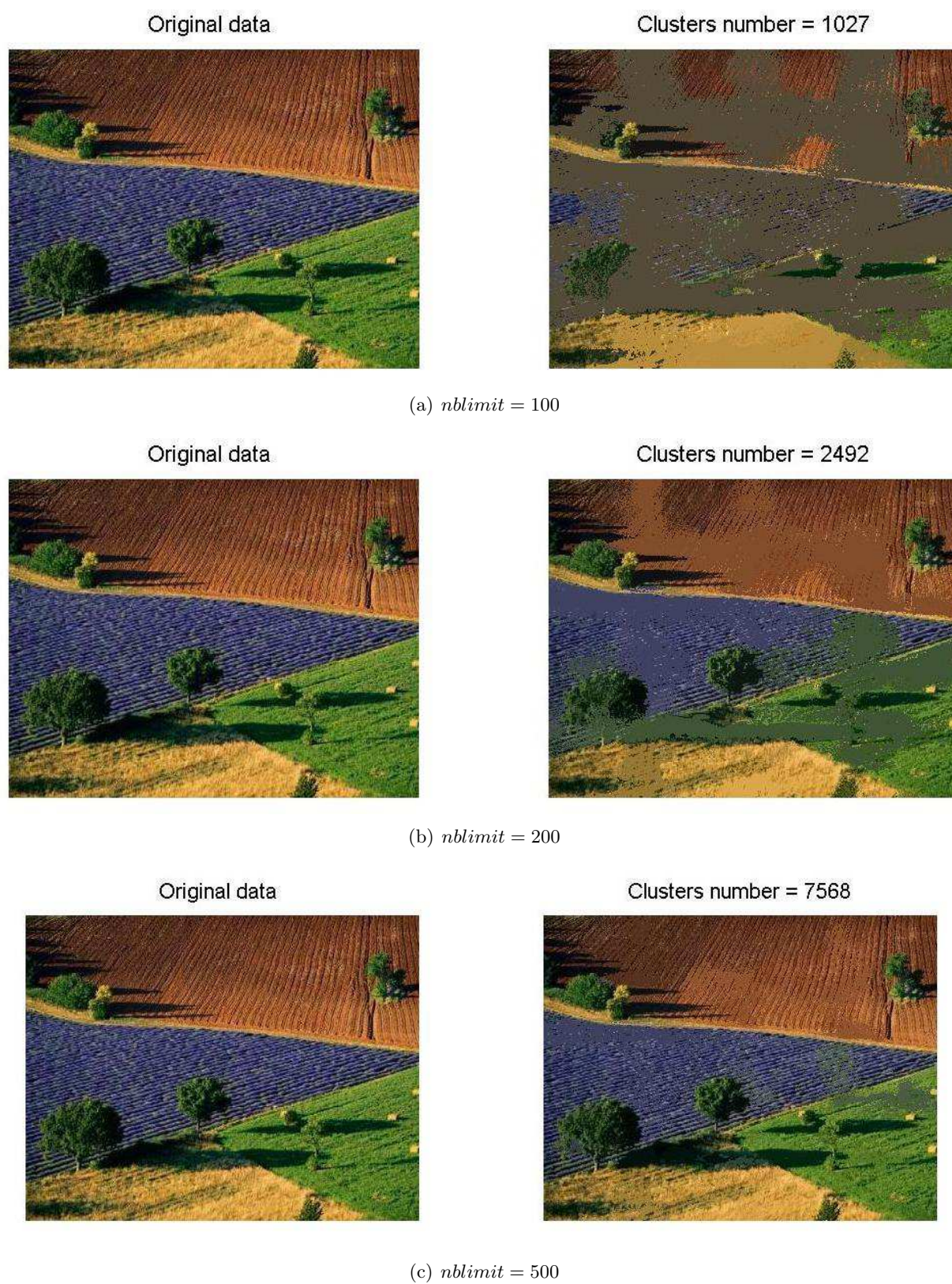
(b) nblimit=100



(c) nblimit=200

FIGURE 3.22 – Test sur un exemple de mahua

respectivement à 6460.1s et 5075.6s. La boucle implémentée pour tester les diverses valeurs de k n'est donc pas d'un coût linéaire en fonction de la valeur du *nblimit*. Ces premiers résultats sur des images couleurs en considérant de données $5D$ sont encourageants mais une étude doit être menée pour tester les limites et dégager des propriétés du spectral clustering dans le cadre de la segmentation d'images.

FIGURE 3.23 – Champs colorés dans le Vaucluse, $N = 128612$ points

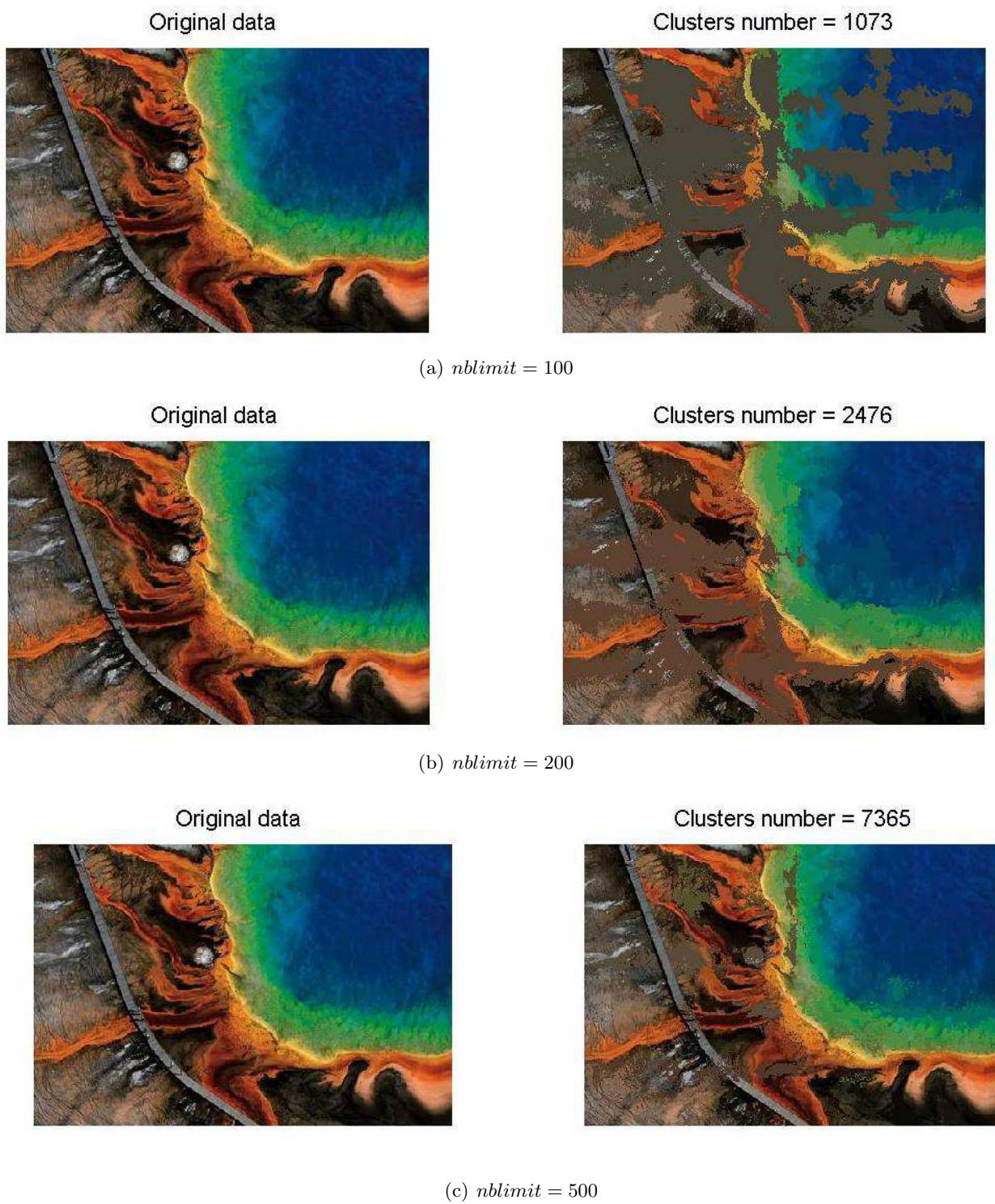


FIGURE 3.24 – Geyser du grand Prismatic, $N = 155040$ points

Chapitre 4

Extraction de connaissances appliquée à la biologie et l'imagerie médicale

L'objectif de ce chapitre est donc de confronter la méthode de spectral clustering à un cadre réel où les données sont bruitées, multidimensionnelles et nombreuses.

Dans la suite, nous abordons deux domaines principalement touchés par l'arrivée des nouvelles technologies et qui requièrent des méthodes non supervisées pour extraire des connaissances. Les domaines concernés sont celui de l'analyse de données génomiques et celui de la Tomographie par Emission de Positons (TEP). Ces études sont effectuées dans le cadre de deux collaborations :

1. *Etude de profils temporels d'expressions de gène issus de biopuce* : collaboration avec le laboratoire Symbiose et Pathologie des Plantes de l'Ecole Nationale Supérieure Agronomique de Toulouse (ENSAT).
2. *Segmentation d'images dynamiques TEP* : collaboration (en cours) avec l'Unité *Imagerie et Cerveau* (UMRS INSERM) de l'université François Rabelais de Tours.

La version entièrement non supervisée de la méthode de spectral clustering est utilisée comprenant les heuristiques (1.2) pour le choix du paramètre et (3.5) pour déterminer le nombre de clusters k . L'information issue des clusters des données réelles est étudiée. Afin de juger de la pertinence et de la validité des résultats issus du spectral clustering, ces résultats seront comparés à ceux issus de méthodes reconnues dans les domaines respectifs à savoir la méthode de cartes de Kohonen [98] pour les données temporelles génomiques et la méthode k -means [112] pour la segmentation d'images TEP. Ces domaines permettront de mettre en application une partie du matériel théorique et numérique développés dans les chapitres précédents et vérifier leur compatibilité pour des données réelles.

Partie 1 :

Classification de profils temporels d'expression de gènes

Après avoir introduit le contexte à travers le cadre de l'étude et le principe des biopuces, nous adapterons dans un premier temps une méthode basée sur les cartes de Kohonen aux données d'expressions temporelles de gènes qui servira par la suite d'outil référence de comparaison. Dans un second temps, la méthode du spectral clustering sera appliquée sur ces mêmes données. Enfin, les résultats seront comparés dans une dernière partie.

4.1 Présentation du problème biologique

Dans la suite, le cadre de l'étude sera introduit ainsi que le principe des biopuces ou puces ADN utilisées pour obtenir l'ensemble des données génomiques.

Cadre de l'étude

Les légumineuses fixent l'azote atmosphérique en établissant une symbiose avec une bactérie racinaire (Rhizobium). Cette aptitude explique la forte teneur en protéines et permet la synthèse de protéines sans avoir recours à une fertilisation azotée chimique. La culture des légumineuses, à la différence des céréales ou des oléo-protéagineux comme le colza et le tournesol, constitue le meilleur moyen de produire des protéines tout en respectant l'environnement.

La légumineuse modèle considérée est le Médicago Truncatula [102]. Cette espèce est la proche d'un point de vue phylogénétique (i.e étude de la formation et de l'évolution des organismes vivants en vue d'établir leur parenté) de la plupart des légumineuses cultivées en Europe. Elle fait partie du groupe des galégoïdes contenant les tribus des Trifoliées (luzernes, trèfles...), des Viciées (pois, féveroles, lentilles..) et Cicérées (pois chiche...). Le Médicago Truncatula est une luzerne proche de la luzerne cultivée (*M.sativa*). La taille du génome est 450 mégabases (10^6 bases) par cellule, équivalente à celle du riz. Il s'agit d'une espèce diploïde (appariement des chromosomes contenus dans la cellule biologique) et autogame (autoféconde). Le principal inconvénient de ces espèces est sa sensibilité à un grand nombre de pathogènes surtout racinaires.

Une bactérie connue pour attaquer plus de 200 espèces incluant beaucoup de cultures et de récoltes agricoles est le *Ralstonia Solanacearum*. Sa forme en seringue permet une infection par la racine des plantes, véritable autoroute pour les pathogènes. Des expérimentations basées sur l'inoculation de la bactérie dans le *Medicago Truncatula* ont dévoilé l'existence de plantes résistantes (mutant HRP). L'observation des gènes concernés par les maladie foliaire et racinaire et l'analyse des gènes issus des plantes résistantes est nécessaire. On a recours aux puces à ADN.

Principes des expériences effectuées

La fabrication de puce à ADN [94] consiste à immobiliser de façon ordonnée sur un support solide des molécules d'ADN, appelés sondes, spécifiques des gènes que l'on souhaite étudier comme

le montre la figure 4.1. Une fois la puce fabriquée, l'expression des gènes d'un échantillon biologique peut être analysée. Les cibles sont obtenues par transcription inverse des ARN messagers extrait de l'échantillon et par marquage simultané. Ce marquage est réalisé à l'aide de fluorophores spécifiques (couleur rouge avec le cyanine5 et couleur verte avec le cyanine3) dans le cas de puce sur lame de verre. Ces solutions sont déposées sur la puce, de telle sorte que les cibles déposées en excès s'hybrident avec la séquence complémentaire sur la puce. L'utilisation des fluorophores différents permet de déposer simultanément sur la puce deux échantillons différents, ce qui constitue un avantage indéniable pour des approches différentielles où l'on compare le niveau d'expression des gènes entre un tissu sain et tumoral, ou entre un organisme sauvage et mutant.

L'expression des gènes est un procédé consistant à convertir une séquence ADN en une protéine (cette dernière participant à la construction des cellules, organes et assurant leur fonction). En d'autres termes, elle permet la synthèse des protéines chargées d'exécuter le programme cellulaire. Elle est principalement régulée au niveau de la copie de gènes en ARNm (transcription). La régulation est assurée par la fixation des protéines inhibitrices (répresseurs) ou activatrices (activateurs) sur les régions de l'ADN où débute la transcription des gènes (régions régulatrices).

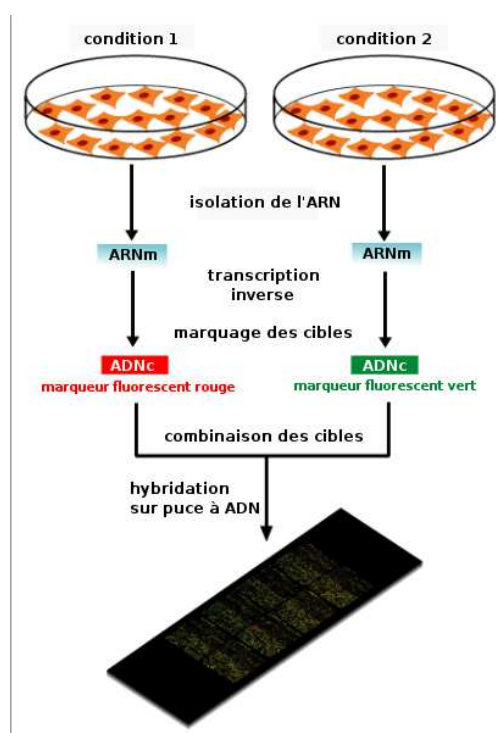


FIGURE 4.1 – Principe des puces à ADN

Dans le cadre de nos expériences, une puce à lame de verre avec des marquages fluorophores est utilisée. Deux cellules sont étudiées : l'espèce 'A17' du *Médicago Truncatula* et son espèce mutante résistante à la bactérie *Ralstonia Solanacearum* notée 'F83'. Dans les deux cas, la bactérie est inoculée à ces deux espèces et les niveaux d'expression, à divers instants, d'une partie du génome sont observés (cf figure 4.2).

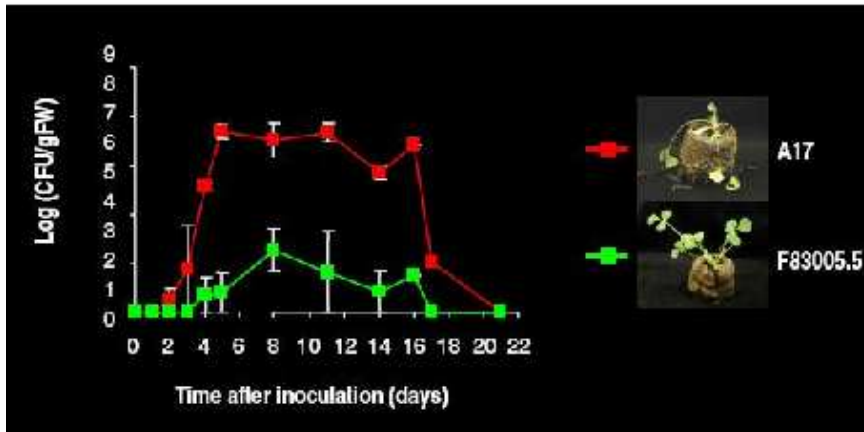


FIGURE 4.2 – Root pathosystems in *Medicago Truncatula*

Nature des données

Les puces à ADN permettent de mesurer le niveau d'expression relatif de chaque gène dans un échantillon cellulaire comparé à un contrôle de référence. Ici le contrôle de référence est la cellule A17. La puce utilisée est une puce dédiée c'est-à-dire une puce où on étudie une partie précise et connue du génome. Les gènes considérés ont des fonctions particulières dans l'exécution du programme cellulaire et sont répertoriés sur le tableau 4.1. Les chiffres obtenus mesurent le niveau absolu d'expression des gènes, exprimés en $\log_{10}$ (et en unités arbitraires). Les données ont été collectées dans huit expérimentations produit cartésien de l'ensemble des lignées par celui des temps puis normalisées (et exprimés en $\log_{10}$). On obtient donc une matrice de taille 2561×7 de mesures répétées et dépendantes du temps.

Une méthode connue pour ses applications dans le domaine de la bioinformatique [98, 99] est la méthode à noyau cartes auto-organisatrices (ou *Self Organizing Maps*). Dans la suite, la méthode SOM sera présentée et adaptée à des profils temporels d'expression de gène. Puis nous testerons le spectral clustering dans le cadre de cette problématique.

4.2 Méthode Self Organizing Maps

Après une présentation de la méthode Self Organizing Maps ou carte de Kohonen, les éléments de la méthode seront adaptés à la classification de profils génomiques temporels. Ces données seront ensuite testées sur les espèces A17, F83 et sur l'expression gène à gène entre les espèces A17 et F83.

4.2.1 Présentation de la méthode Self Organizing Maps

Le "Self Organizing Maps" (SOM) est un des modèles les plus connus dans les réseaux de neurones [68, 69]. Il est basé sur une méthode de scaling multidimensionnel qui projette les données dans un espace de plus petite dimension. Un SOM est formé de neurones placés sur une grille d'une ou deux dimensions. Les neurones de la carte sont connectés aux neurones adjacents par une relation de voisinage imposée par la structure de la carte. En général, la topologie de la carte est rectangulaire ou hexagonale. Le nombre de neurones détermine la "granularité" de la carte ce qui affecte l'exactitude, la précision et la capacité générale du SOM. On fixe les relations topologiques

Catégorie	Nombre	Pourcentage
Paroi cellulaire	93	5
Cytosquelette	15	1
Transport membranaire	51	3
sécrétion	15	1
Métabolisme primaire	436	22
Métabolisme secondaire et hormones	177	9
Métabolisme de l'ADN	27	1
ARN et Transcription	287	15
Synthèse protéique	117	6
Transduction des signaux	104	5
Cycle cellulaire	4	0
Divers	116	6
Défense	186	9
Stress abiotique	32	2
Fonction inconnue	210	11
Homologie inconnue	97	5
TOTAL	1967	100

TABLE 4.1 – Catégories génomiques présentes sur la puce dédiée

et le nombre de neurones afin d'obtenir le nombre de clusters désiré avec une taille de voisinage contrôlant le lissage et ainsi le profil des gènes.

Définitions et notations

Soit l'ensemble de données $\mathcal{X} = \{x_1, \dots, x_N\} \in \mathbb{R}^p$ où p est le nombre d'instantants. La représentation matricielle de cet ensemble est la matrice des données $X : X \in \mathcal{M}_{N,p}, X = [x_i^j]$ où x_i^j est la valeur de l'expression à l'instant j de la donnée x_i .

Soit une grille 2D composée de K neurones. On note par $r^i \in \mathbb{R}^2$ les coordonnées sur la grille du neurone i pour $i \in \{1, \dots, K\}$. Chaque neurone k , $k \in \{1, \dots, K\}$, est représenté par un vecteur poids (ou vecteur représentant) de dimension p , noté $m_i = [m_i^1, \dots, m_i^p]^T$ où p est la dimension des vecteurs des données.

On définit une mesure de similarité s entre les vecteurs poids et les vecteurs de données pour assigner les vecteurs données au neurone dont le vecteur poids associé est le plus semblable. Cette mesure vérifie les propriétés suivantes :

$$\begin{cases} \forall (x_i, m_j), & s(x_i, m_j) = s(m_j, x_i), \\ \forall (x_i, m_j), & 0 \leq s(x_i, m_j) \leq 1, \\ \forall (x_i, m_j), & s(x_i, m_j) = 1 \Leftrightarrow i = j. \end{cases}$$

On appelle BMU (*Best-Matching Unit*) le neurone c le plus semblable au vecteur des données x_j :

$$s(x_j, m_c) = \min_k -s(x_j, m_k)$$

La mise à jour d'un vecteur poids consiste à intégrer la différence d'amplitude entre tous les vecteurs de données qui lui sont associés à partir du BMU et lui-même. Le vecteur poids est le représentant de

ses vecteurs données donc il doit être ajusté à chaque itération jusqu'à convergence de l'algorithme. Le neurone $c = \arg \min_k -s(x_j, m_k)$ et les neurones voisins à ϵ du neurone c sont mis à jour à l'itération t à l'aide du noyau de voisinage suivant :

$$h_{ci}(t) = \alpha(t)h(s(r_c, r_i), \epsilon).$$

Le noyau de voisinage est composé de deux parties :

- *fonction de voisinage* $h(s(.,.), \epsilon)$: seuls les voisins les plus proches sont mis à jours à l'itération t . En général $h(s(.,.), \epsilon)$ est une fonction indicatrice, une gaussienne ou la composée d'une gaussienne et d'une indicatrice et le seuil de voisinage est défini par ϵ . Par exemple, pour une indicatrice : $h(s(.,.), \epsilon) = \mathbb{I}_{s(.,.) \leq \epsilon}$
- *fonction d'apprentissage* α : sert de pondération ; il s'agit d'une fonction strictement décroissante en fonction des itérations t , le rayon initial verifie : $\alpha_0 = \alpha(0) \in [0, 1]$.

Le critère de convergence est défini par la différence, entre deux itérations successives, de la mesure de similarité des vecteurs de données avec leur vecteur représentant. Autrement dit,

$$\exists \delta > 0, \sum_{j=1}^N \sum_{i=1}^K \mathbb{I}_{s(x_j, m_i(t)) \leq \epsilon} s(m_i(t), x_j) - \sum_{j=1}^N \sum_{i=1}^K \mathbb{I}_{s(x_j, m_i(t-1)) \leq \epsilon} s(m_i(t-1), x_j) \leq \delta. \quad (4.1)$$

En cas de divergence, on fixe un nombre d'itération maximum noté t_{max} . On notera L l'itération vérifiant le critère de convergence l'équation (4.1).

L'algorithme de cette méthode est défini comme suit :

Algorithm 4 Algorithme SOM général

1. $t=1$,
2. arrêt=faux.
3. Pour tout $i = 1, \dots, K$ initialisation des vecteurs poids $m_i = [m_i^1, \dots, m_i^p]^T$.
4. Répéter
 - (a) Pour tout $j = 1, \dots, N$, identifier le vecteur poids m_c le plus semblable à x_j , vérifiant :

$$s(x_j, m_c) = \min_k -s(x_j, m_k).$$

- (b) Pour tout $j = 1, \dots, N$, pour tout $i = 1, \dots, K$, mise à jour du vecteur poids m_i :

$$m_i(t+1) = m_i(t) + \alpha(t)h(s(r^i, r^c), \epsilon)[x_j(t) - m_i(t)].$$

- (c) Faire $t=t+1$.

- (d) Mise à jour (arrêt).

5. Jusqu'à arrêt=vrai $\Leftrightarrow t = L$ et $t \leq t_{max}$.
-

Problèmes liés à cette méthode

L'algorithme 4 impose la définition de certains paramètres :

1. *nombre de neurones dans la grille* : comme il s'agit d'une méthode de classification non supervisée, le nombre de clusters est supposé inconnu,
2. *initialisation des vecteurs poids* : une bonne initialisation permet à l'algorithme de converger plus vite vers la meilleure solution. Trois types d'initialisation sont, en général, utilisés :
 - *initialisation aléatoire* : les vecteurs poids initiaux sont des valeurs prises aléatoirement entre le minimum et le maximum des valeurs des vecteurs dans l'espace des données,
 - *initialisation aléatoire sur les gènes* : les vecteurs poids sont des gènes pris aléatoirement dans l'ensemble des données,
 - *initialisation issue de l'ACP* : les vecteurs initiaux sont des vecteurs propres associés aux plus grandes valeurs propres de l'ensemble des données.
3. *Choix de la mesure de similarité* : choix important suivant le problème que l'on traite. En général, la norme euclidienne $\|\cdot\|_2$ est utilisée.
4. *Noyau de voisinage* : à modifier suivant les données à traiter.

4.2.2 Adaptation de la méthode Self Organizing Maps pour données temporelles d'expression de gènes

Dans la suite, chaque problème lié aux cartes de Kohonen recensé dans la section précédente est abordé et adapté à la classification de profils temporels d'expression de gènes.

Dimension de la grille

Nous disposons de données répétées et dépendante du temps : $\mathcal{X} = \{x_1, \dots, x_N\} \in \mathbb{R}^p$. Le but est de définir le nombre de profils différents et de regrouper les gènes correspondants à chaque profil. Tous les profils possibles peuvent être recensés en utilisant un système ternaire c'est-à-dire un système de numération de la base 3 :

- soit l'expression du gène x_i est constant entre deux instants (invariance),
- soit l'expression du gène x_i augmente entre deux instants (répression),
- soit l'expression du gène x_i diminue entre deux instants (inhibition).

Le nombre d'unités K de la grille de neurones est de 3^{p-1} et les dimensions de la grille sont $[3 \ 3^{p-2}]$. La structure de la grille peut être rectangulaire ou hexagonale.

Initialisation par balayage de profils

Le nombre de neurones correspond au nombre de combinaisons possibles de profils. Ainsi les vecteurs poids initiaux associés aux neurones doivent balayer toutes les combinaisons de profils. De plus, cette initialisation doit prendre en compte des informations sur les données à savoir moyenne et écart-type de manière à converger plus vite et donc diminuer le temps de calcul.

Pour chaque vecteur poids, la valeur à l'instant t_0 est égale à la moyenne des données à t_0 . Ensuite, une amplitude adaptée au jeu de données, à savoir l'écart-type entre deux instants consécutifs est définie. Pour finir, la conversion en base 3 des chiffres $\{1, 2, \dots, 3^{p-1}\}$ est utilisée afin de balayer tous les cas. Les divers profils sont ainsi balayés et adaptés à notre jeu de données. L'avantage est donc un gain en terme de temps de calcul et de coût numérique : le calcul de la matrice de variance-covariance et l'extraction des p vecteurs propres sont évités.

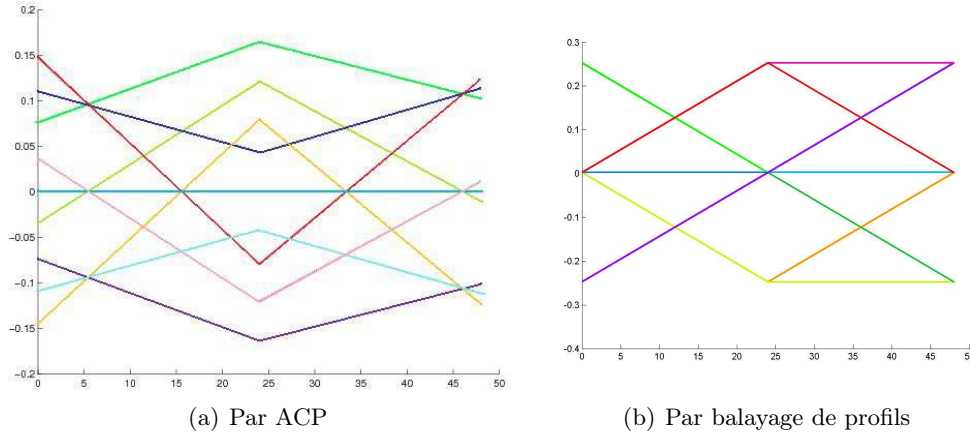


FIGURE 4.3 – Initialisations de la méthode SOM (t en abscisse, niveau d'expression en ordonnée)

Mesure de similarité

Concernant notre problématique, l'utilisation d'une norme (l_1, l_2, l_∞) ne permet pas de distinguer certains types de profils : le fait de sommer (ou de prendre le maximum pour l_∞) les amplitudes entre le vecteur de données et le vecteur poids sur tous les instants peut aboutir à une confusion des profils comme le montre la figure 4.4 :

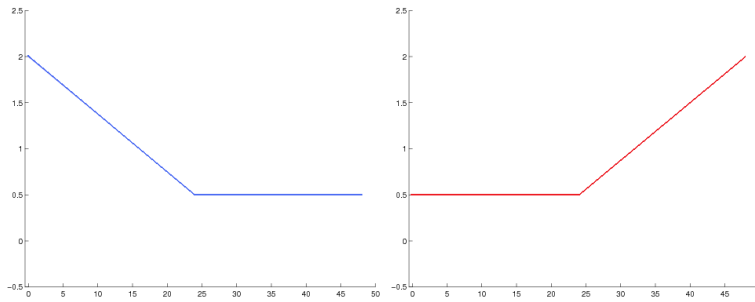


FIGURE 4.4 – Exemple de profils temporels différents mais de même norme (t en abscisse, niveau d'expression en ordonnée)

Dans le cadre de profil temporel, la corrélation entre deux vecteurs x et y est plus adéquate comme mesure de similarité :

$$\text{corr}(x, y) = \frac{x^T y}{\|x\| \|y\|} = \cos(\theta),$$

où θ est l'angle formé par les vecteurs x et y . Plus $\cos(\theta)$ est proche de 1, plus les vecteurs x et y sont colinéaires et décrivent le même profil.

A chaque vecteur de données $x_j \in \mathbb{R}^p, j = 1, \dots, N$, on associe $p - 1$ vecteurs de $\mathbb{R}^2$ définis par : $X_i^j = [t_{i+1} - t_i, x_j^{i+1} - x_j^i]^T, \forall i \in \{1..p - 1\}$.

De même, on associe à chaque vecteur poids $m_k \in \mathbb{R}^p, k \in \{1, \dots, 3^{p-1}\}$, $p - 1$ vecteurs de $\mathbb{R}^2$ définis par : $m_k^i = [t_{i+1} - t_i, m_k^{i+1} - m_k^i]^T, \forall i \in \{1..p - 1\}$.

La corrélation entre un vecteur de données $x_j \in \mathbb{R}^p$ et un vecteur poids $m_i \in \mathbb{R}^p$ est définie par

la relation suivante :

$$corr(x_j, m_i) = \sum_{i=1}^{p-1} \frac{(X_i^j | M_i^k)}{\|X_i^j\| \|M_i^k\|}.$$

Donc le BMU, noté c , est défini par :

$$s(x_j, m_c) = \max_k corr(x_j, m_k) = \min_k -corr(x_j, m_k). \quad (4.2)$$

L'avantage de cette mesure de similarité est l'incorporation de l'information temporelle dans la méthode de classification.

Mise à jour des vecteurs poids

Dans le cadre de notre problématique, on dispose d'un ensemble de données conséquent par rapport au nombre de profils ($N \gg p$). Par conséquent, plusieurs centaines voire milliers de gènes seront associés à un neurone ($N \gg K$). Cela implique un réajustement de la fonction d'apprentissage α lors de la mise à jour des vecteurs poids. Le poids associé à la différence d'amplitude entre le vecteur de données et le vecteur poids le plus semblable doit être adapté aux données. Le rayon initial est, par conséquent, défini en fonction du nombre de données N , du nombre de neurones dans la grille K et de l'écart-type des données noté σ :

$$\alpha(0) = \frac{\sigma K}{N}.$$

Dans le cadre de l'étude, $K = 3^{p-1}$ où p est le nombre d'instantanés considérés.

L'algorithme adapté aux profils temporels d'expression de gène est décrit ci-dessous.

Algorithm 5 Algorithme SOM adapté à des données temporelles

1. $t=1$,
2. arrêt=faux,
3. Pour tout $i = 1, \dots, K$ initialisation des vecteurs poids par balayage des profils.
4. Répéter
 - (a) Pour tout $j = 1, \dots, N$, identifier le vecteur poids m_c le plus semblable à x_j , vérifiant :

$$s(x_j, m_c) = \min_k -corr(x_j, m_k).$$

- (b) Pour tout $j = 1, \dots, N$, pour tout $i = 1, \dots, K$, mise à jour du vecteur poids m_i :

$$m_i(t+1) = m_i(t) + \alpha(t)h(s(r^i, r^c), \epsilon)[x_j(t) - m_i(t)],$$

$$\text{avec } \alpha(0) = \frac{\sigma K}{N}.$$

- (c) faire $t=t+1$
 - (d) mise à jour (arrêt)

5. Jusqu'à arrêt=vrai $\Leftrightarrow t = L$ et $t \leq t_{max}$
-

Remarque 4.1. Ce cas d'étude peut s'avérer complexe si le nombre d'instantanés augmente. En effet, l'initialisation et le nombre de noeuds du maillage augmentent en puissance de 3 à mesure qu'on ajoute des instantanés.

4.2.3 Simulations

On considère un ensemble composé de 2561 gènes. On dispose de données représentant l'expression (i.e le log-ratio d'expression) à trois instants différents ($t_0 = 0$, $t_1 = 24$, $t_2 = 48$) de cet ensemble de gènes pour deux espèces différentes : l'espèce A17 et F83. Donc $n = 2561$ et $p = 3$.

Soit $x_i \in \mathbb{R}^3$ (resp. $y^i \in \mathbb{R}^3$) l'expression de l'espèce A17 (resp. F83) du gène i . De plus, on définit, pour un gène i , $i \in \{1, \dots, 2561\}$ l'expression de la différence entre les deux espèces par $(z_i) = (x_i - y_i) \in \mathbb{R}^3$. On utilise donc une grille rectangulaire composée de 3^2 neurones. La fonction voisinage $h(s(\cdot, \cdot), \epsilon)$ est la composée d'une gaussienne et d'une indicatrice. On restreint la mise à jour des vecteurs poids aux neurones adjacents au neurone $c = \arg \min_k -s(x_j, m^k)$ i.e $\epsilon = 1$:

$$\forall j \in \{1, \dots, K\}, r^j \in \mathbb{R}^3, h(s(r^j, r^c), 1) = e^{-s(r^j, r^c)} \mathbb{I}_{\{s(r^j, r^c) \leq 1\}}$$

La fonction d'apprentissage α est linéaire :

$$\alpha(t) = \alpha_0 \left(\frac{1-t}{t_{max}} \right) \text{ avec } \alpha_0 = \frac{3^{p-1} \sigma}{N}$$

et σ l'écart-type des données et t_{max} le nombre maximal d'itérations.

La méthode SOM adaptée est appliquée à chacune de ces espèces puis sur la différence d'expression entre les deux espèces. Pour chaque clusters, le vecteur poids après convergence de l'algorithme et les centroïdes des gènes associés aux clusters sont représentés sur les figures 4.5, 4.6 et 4.7. Le nombre de gènes appartenant à chaque cluster est aussi indiqué sur les figures 4.5, 4.6 et 4.7.

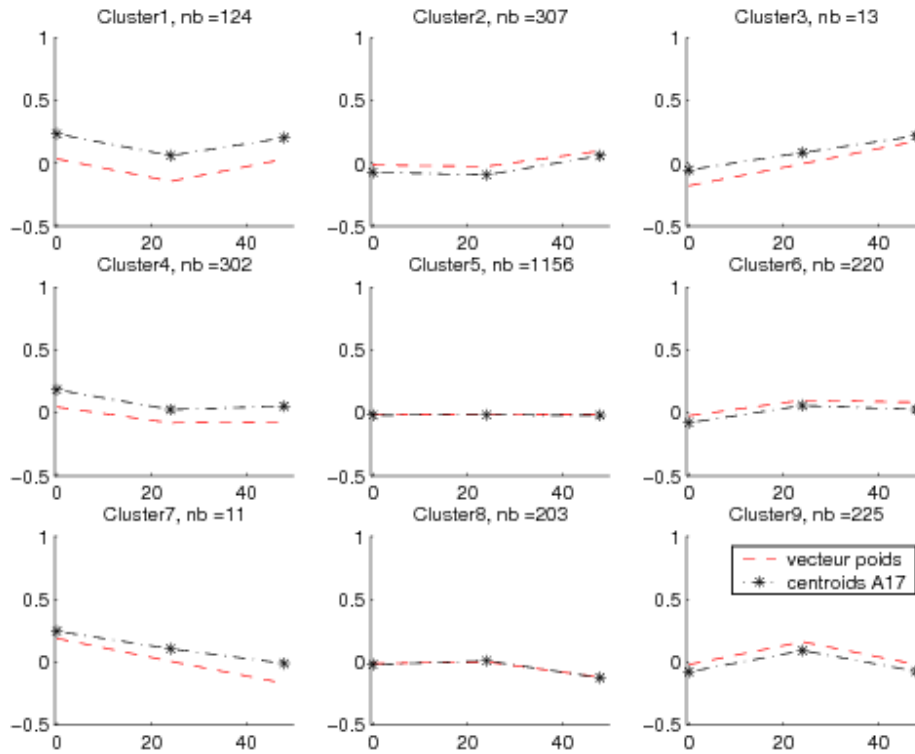


FIGURE 4.5 – Résultats pour l'espèce A17 (t en abscisse, niveau d'expression en ordonnée)

Concernant l'espèce A17, les différents profils possibles sont retrouvés après convergence de l'algorithme et sont affichés sur la figure 4.5. Le cluster 5 contient près de la moitié des gènes. En raison du faible nombre de gènes appartenant aux clusters 1, 3 et 7, le centroïd des gènes et le vecteur poids final sont distants l'un de l'autre dans ces clusters. Pour une étude plus fine du cluster 5, il faudrait à nouveau appliquer la méthode SOM aux 1156 gènes afin d'adapter l'étude à cet ensemble donc à un écart-type plus petit.

Effectuons la même simulation pour l'espèce F83 dont les résultats sont affichés sur la figure 4.6. Pour l'espace F83, le cluster 5 contient, à nouveau, un peu plus de la moitié des gènes et que, dans

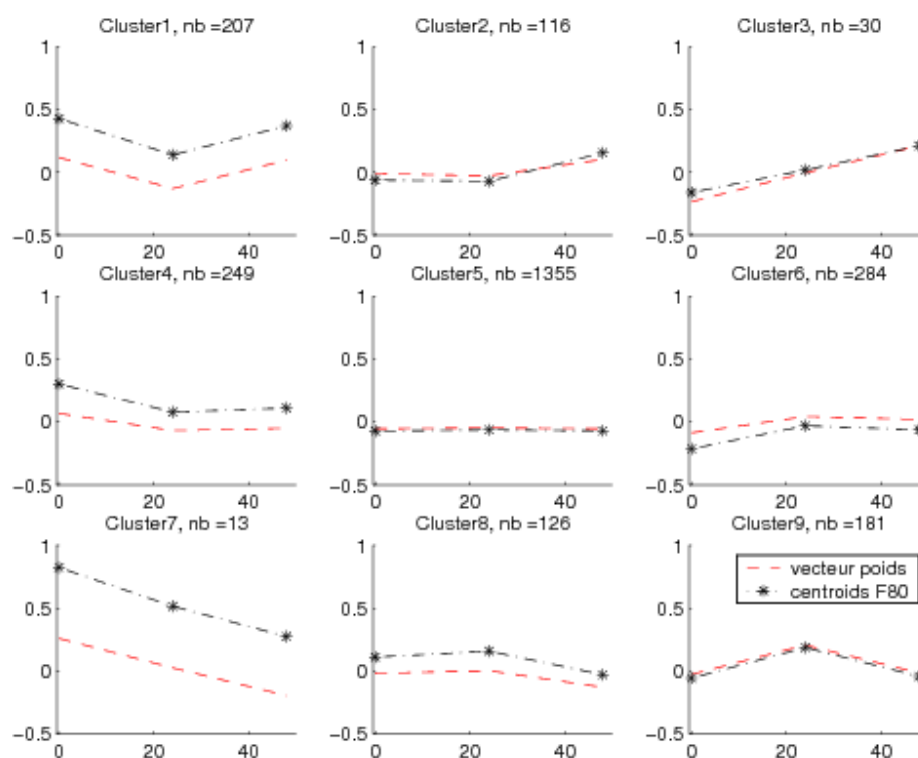


FIGURE 4.6 – Résultats pour l'espèce F83 (t en abscisse, niveau d'expression en ordonnée)

les clusters 1, 7 et 8, le centroïd des gènes et le vecteur poids final sont distants l'un de l'autre. L'écart-type de ce jeu de données est supérieur à celui de l'espèce précédente : $\sigma = 0.25$ contre 0.2 pour A17. Cette étude permet de distinguer des gènes qui ont un comportement identique face à l'injection de la bactérie le cluster 5 ne contient plus que 13% des gènes et les centroïdes des gènes et les vecteurs poids après convergence de l'algorithme sont quasiment tous confondus. Ces résultats ont tous été validés par les experts. Ces résultats sont en totale cohérence avec les observations des spécialistes en génomique. Ainsi les résultats de cette méthode serviront de référence pour juger des résultats du spectral clustering.

Remarque 4.2. *L'appréciation des résultats dépend de la problématique biologique fixée au départ. Cet algorithme classe les principaux profils de gènes. Pour une étude plus fine, notamment au niveau des amplitudes au sein d'un même cluster (d'un même profil temporel de gène), il faudra approfondir l'étude des nombreux paramètres définis pour l'adaptation de la méthode SOM : étape*

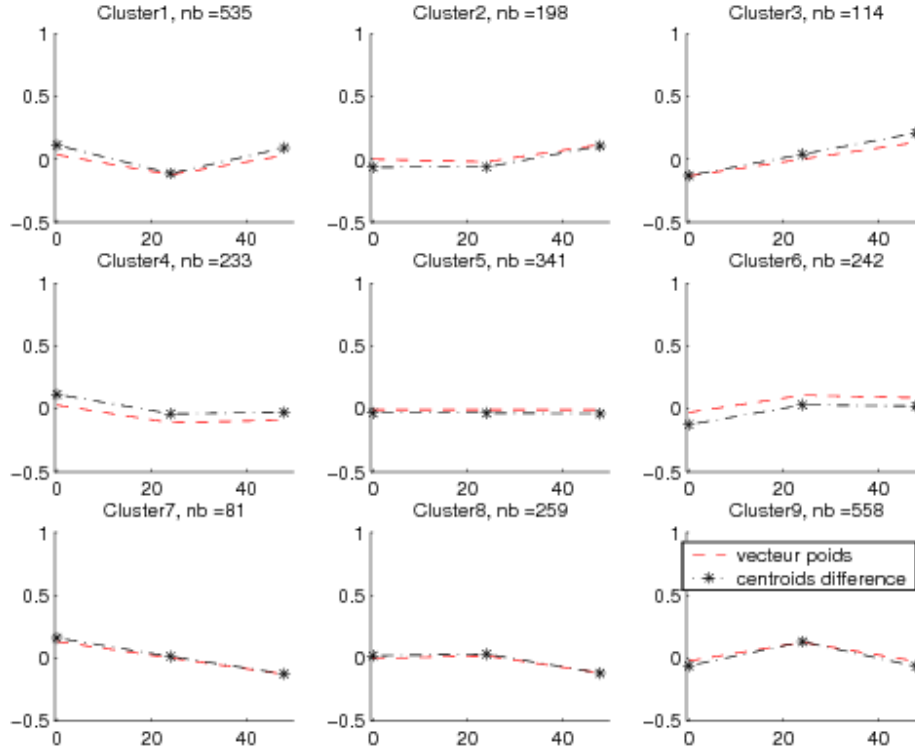


FIGURE 4.7 – Différence entre les espèces A17 et F83 (t en abscisse, niveau d'expression en ordonnée)

d'initialisation et fonctions pour la mise à jour des vecteurs poids, dimension de la grille...

4.3 Spectral Clustering sur les données génomiques

Cette partie est consacrée à l'application de la méthode clustering spectral sur les données génomiques présentées dans la section précédente. Les diverses espèces A17 et F83 seront étudiées séparément puis la différence entre les deux espèces dans l'expression de chaque gène seront étudiée et comparée par rapport à la méthode à noyaux SOM adaptée aux profils temporels.

Dans toute la suite de l'étude, l'heuristique (3.5) du choix du nombre de clusters, développée dans le chapitre 3, est utilisée pour extraire une partition de ces données. Le paramètre σ de l'affinité gaussienne vérifie l'équation (1.2) développée au chapitre 1, à savoir :

$$\sigma = \frac{D_{\max}}{2N^{\frac{1}{p}}},$$

avec $D_{\max} = \max_{1 \leq i, j \leq N} \|x_i - x_j\|$. Le spectral clustering est testé pour un minimum de la fonction ratio η le plus proche du nombre de profils possibles. Comme résultat du spectral clustering, il est représenté le vecteur moyen des points associés au même cluster. Un profil temporel moyen de chaque cluster est ainsi représenté.

Etude de l'espèce A17

Tout d'abord, le spectral clustering est appliqué sur les données initiales. Pour l'espèce A17, la figure 4.8 regroupe donc :

- (a) : la représentation dans $\mathbb{R}^3$ des données de l'espèce A17 ;
- (b) : le ratio $\eta = \frac{2}{k(k-1)} \sum_{j=i+1}^k r_{ij}$ en fonction du nombre de classes k ;
- (c) : le résultat du spectral clustering pour $k = 4$. Chaque couleur représente un même profil temporel.

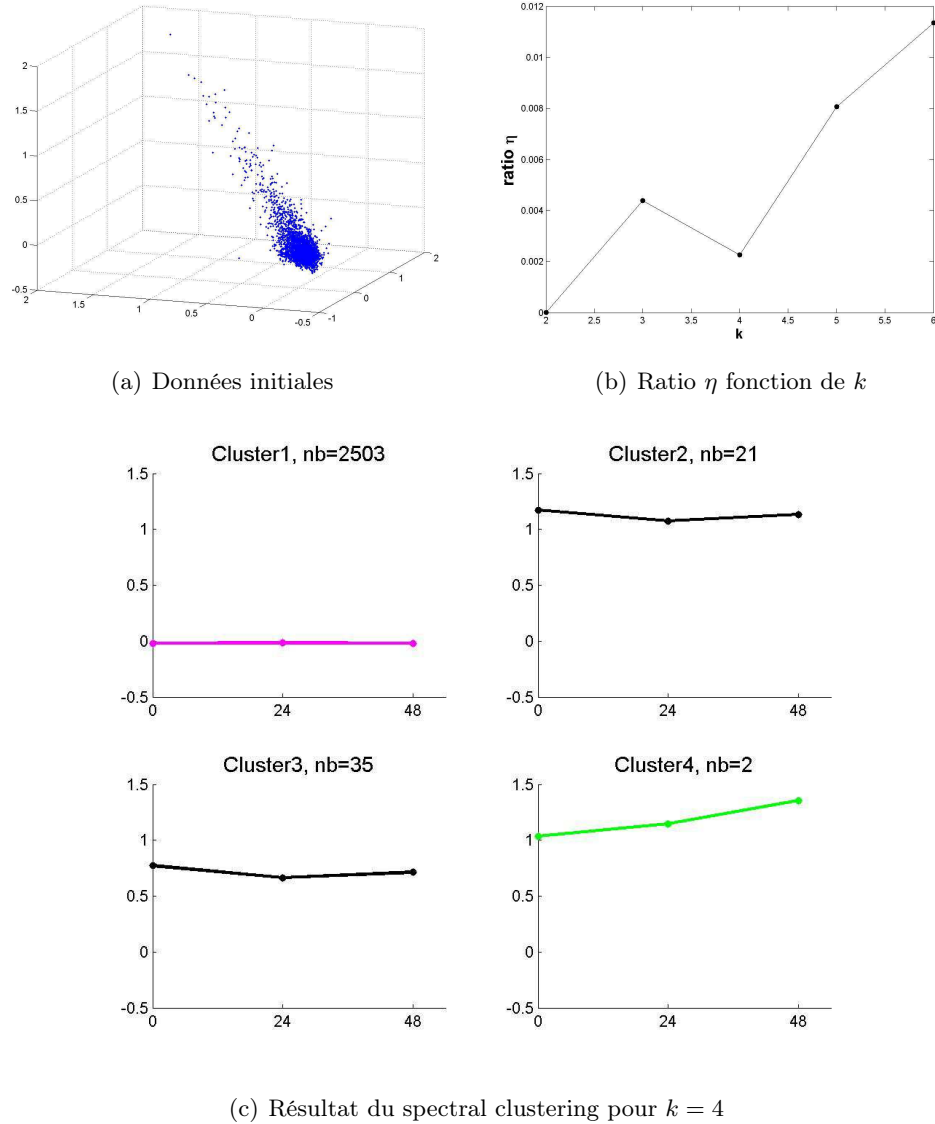


FIGURE 4.8 – Etude avec les données des profils temporels de l'espèce A17

D'après la figure 4.8 (b), on observe que, pour le deuxième minimum du ratio η ($k = 4$), le cluster qui rassemble la quasi totalité des points a un profil invariant et de moyenne nulle. Les autres clusters, représentant environ 2% des gènes, présentent des profils aux amplitudes fortes de répression ou bien

d'inhibition suivi de répression. La distribution des points dans $\mathbb{R}^3$ étant compacte, seuls les points extérieurs à ce nuage de points c'est-à-dire les gènes aux amplitudes et sujets à de fortes réactions de répression et/ou d'inhibition sont considérés comme clusters. De plus, étant donné que l'objectif est de distinguer le profil sans faire intervenir l'amplitude des valeurs des log-ratios d'expressions de gènes, une normalisation des données peut être envisagée. Ainsi la distribution des points est plongée dans une sphère unité où les coordonnées représentent la pondération, l'évolution, par rapport au temps du niveau d'expression du gène de l'espèce envisagée. La figure 4.9 regroupe donc le nuage de points, le ratio η fonction de k et le résultat du spectral clustering pour le deuxième minimum local de η c'est-à-dire pour $k = 11$ clusters.

Pour juger et comparer les résultats de la méthode de spectral clustering, nous comparons les profils des gènes associés à un même cluster à ceux de la méthode SOM. Pour l'espèce A17, le cluster regroupant le plus de gènes correspond à un profil en moyenne constante ce qui est conforme aux résultats avec les SOM et les observations des manipulateurs. Cependant il y a une différence de 418 gènes avec la méthode SOM : les gènes supplémentaires inclus dans ce cluster et les gènes associés au cluster 1 présentent un profil invariant soit sur $[t_0, t_1]$ ou bien sur $[t_1, t_2]$ et de faibles amplitudes sur le complémentaire. Les autres clusters représentent des profils marqués d'inhibition et/ou de répression.

Etude de l'espèce F83

L'étude avec les données brutes de l'espèce F83 aboutit à considérer seulement 4 clusters dont 3 regroupent 5% des gènes. Les données normalisées sont donc, à nouveau, considérées. Et on procède de façon similaire pour déterminer le nombre de clusters puis effectuer l'étude de la qualité du clustering. On observe que le minimum est atteint pour la valeur $k = 19$ clusters. Pour l'espèce F83, les différences entre la méthode SOM et la méthode de Spectral Clustering sont moins notables. Le cluster associé au profil constant/invariant possède quasiment les mêmes gènes. De plus, de nouveaux profils apparaissent comme ceux des clusters 3, 10, 12, 13 et 16. Plusieurs clusters représentent le même profil mais avec des amplitudes différentes. Pour cette espèce, la méthode de spectral clustering partitionne les données par profil et amplitude. En effet, pour l'espèce mutante, les écarts-type ont augmenté respectivement de 35%, 14.5% et 20% suivant les instants $\{t_0, t_1, t_2\}$ par rapport à l'espèce A17 et les moyennes ont augmenté respectivement de 50%, 20% et 43%. On distingue donc des réactions d'inhibition et/ou de répression marquées dans l'amplitude. Ces derniers forment, après normalisation des données, des clusters compacts, déterminés ensuite par le spectral clustering. Donc par rapport aux résultats avec l'espèce A17, les résultats du spectral clustering sont plus satisfaisants.

Etude de la différence entre les espèces A17 et F83

On considère maintenant la différence normalisée, gène à gène, des log-ratios d'expressions entre les espèces A17 et F83. En comparaison avec les résultats de la méthode SOM, le profil invariant contient approximativement le même nombre de gènes et 8 profils différents sont recensés. Le nombre de gènes associé à un même profil est du même ordre de grandeur entre les deux méthodes. A nouveau, des clusters ayant le même profils mais d'amplitudes différentes sont dissociés par cette méthode.

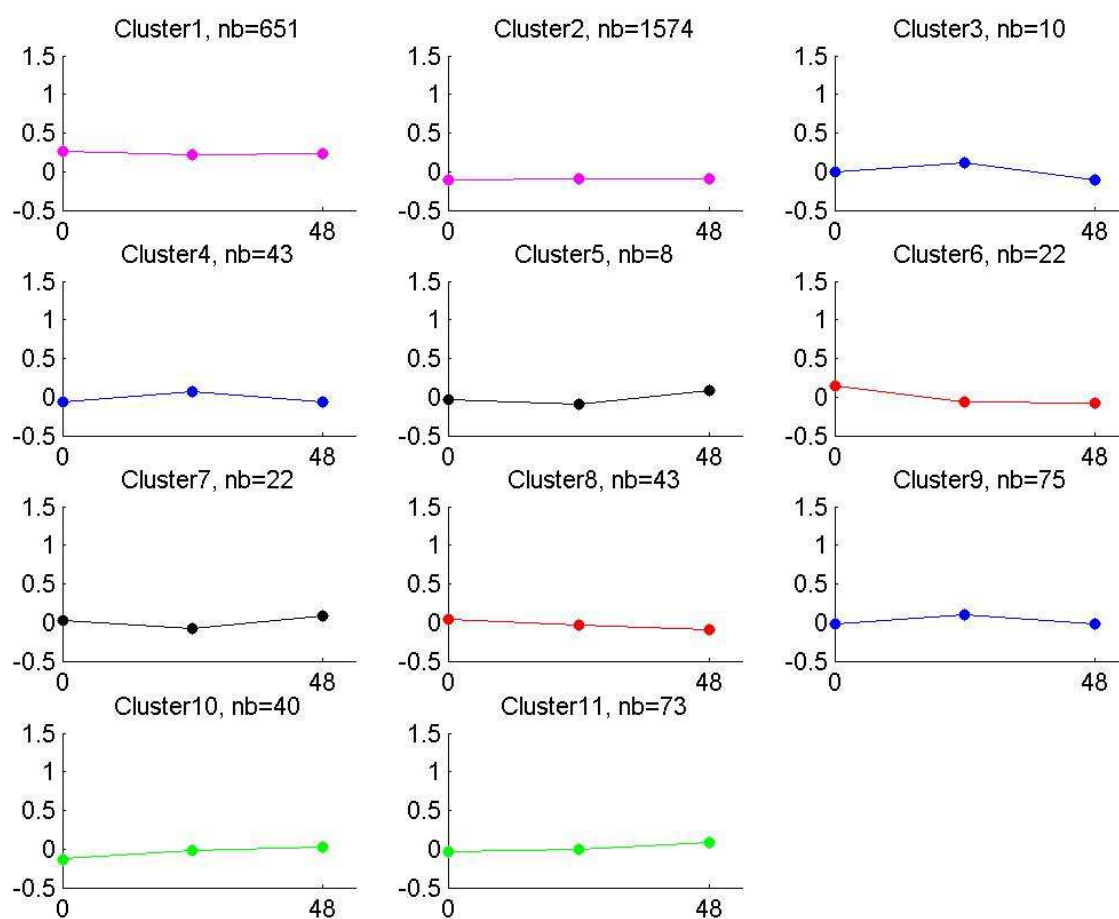
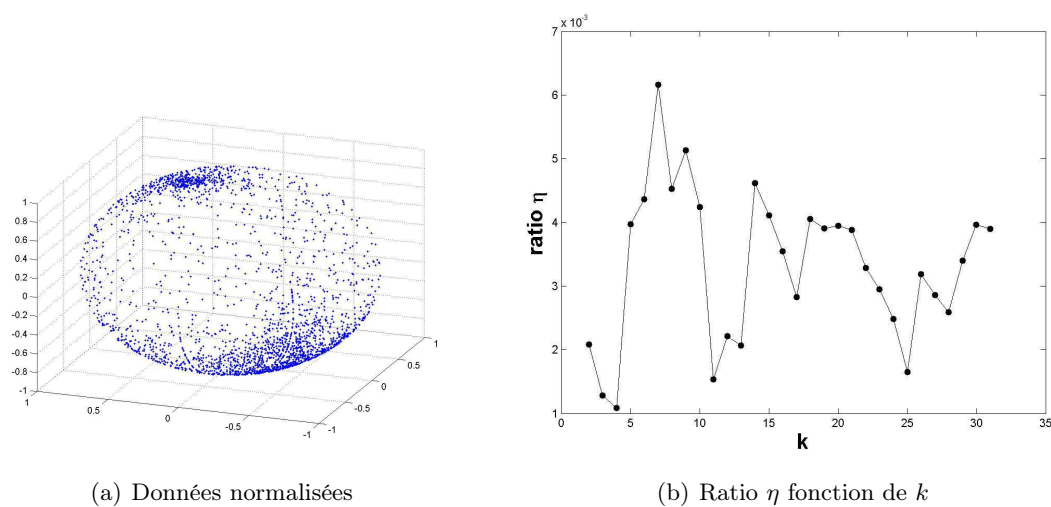
(c) Résultat du spectral clustering $k = 11$

FIGURE 4.9 – Etude avec les données normalisées des profils temporels de l'espèce A17

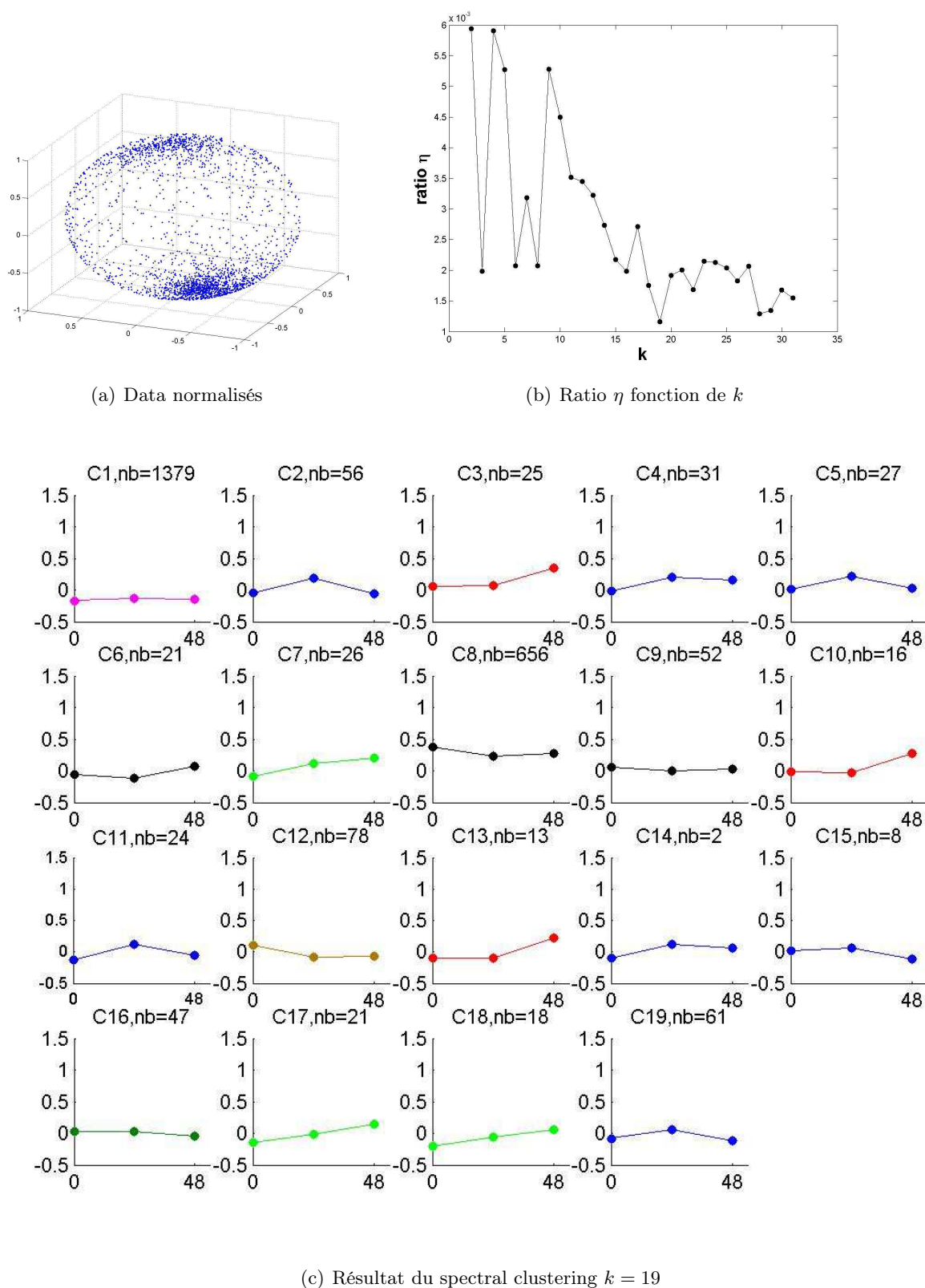


FIGURE 4.10 – Etude avec les données normalisées des profils temporels de l'espèce F83

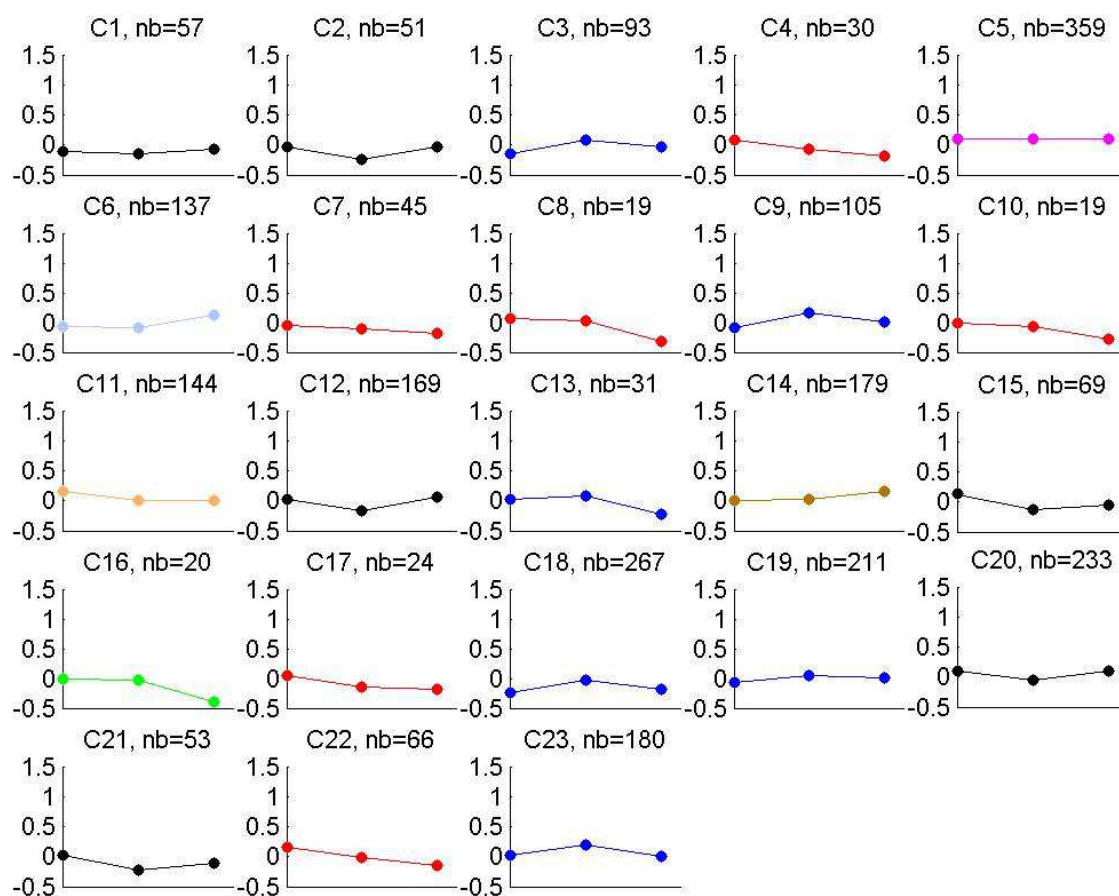
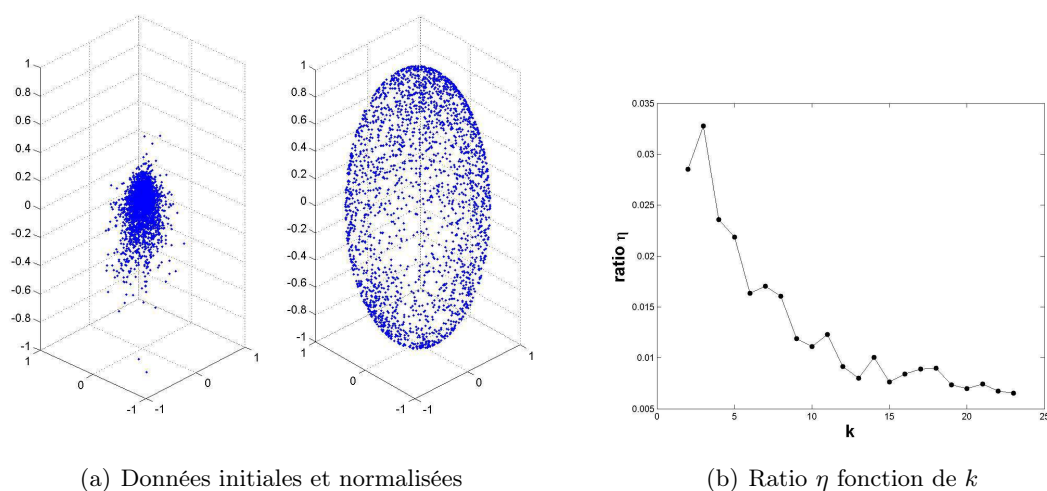
(c) Résultat du spectral clustering $k = 23$

FIGURE 4.11 – Distribution et normalisation des données de la différence entre les profils temporels des espèces A17 et F83

4.4 Discussion

La méthode SOM adaptée fournit des résultats validés par les biologistes. En effet, les informations apportées dans l'adaptation de la méthode améliorent considérablement la distinction des profils et par conséquent les résultats répondent aux attentes des experts. Le spectral clustering donne des résultats exploitables mais moins satisfaisants : les principaux profils sont restitués et certains profils (profils 2,4,6,8 de la figure 4.12) sont englobés dans d'autres (profils 1,3,5,9 de la figure 4.12). Par ailleurs, ces deux méthodes offrent deux différentes approches. La première répond à une problématique précise à l'aide de nombreux paramètres (nombre de neurones, initialisation et mise à jour des vecteurs poids, mesure de similarité) spécialement définis alors que la seconde, non supervisée, appliquée sur les données normalisées ouvrent deux études : les clusters par profils temporels et les clusters par amplitudes d'expression de gènes. A nouveau, il reste à correctement définir la problématique biologique. Dans cette étude, la méthode distingue différents profils temporels mais aussi des clusters au sein d'un même profil temporel aux amplitudes différentes. Suivant la problématique à résoudre, d'autres techniques de normalisation peuvent être envisagées pour pré-traiter les données et spécifier le clustering sur le profil et/ou l'amplitude des expressions de gènes. D'autres problèmes peuvent aussi être étudiés comme la prise en compte du bruit de fond et autres erreurs d'opérande issus de l'expérimentation via les puces ADN. Ainsi des techniques de clustering avec intervalles d'erreur peuvent être envisagées à la fois dans la méthode SOM et pour la méthode de spectral Clustering.

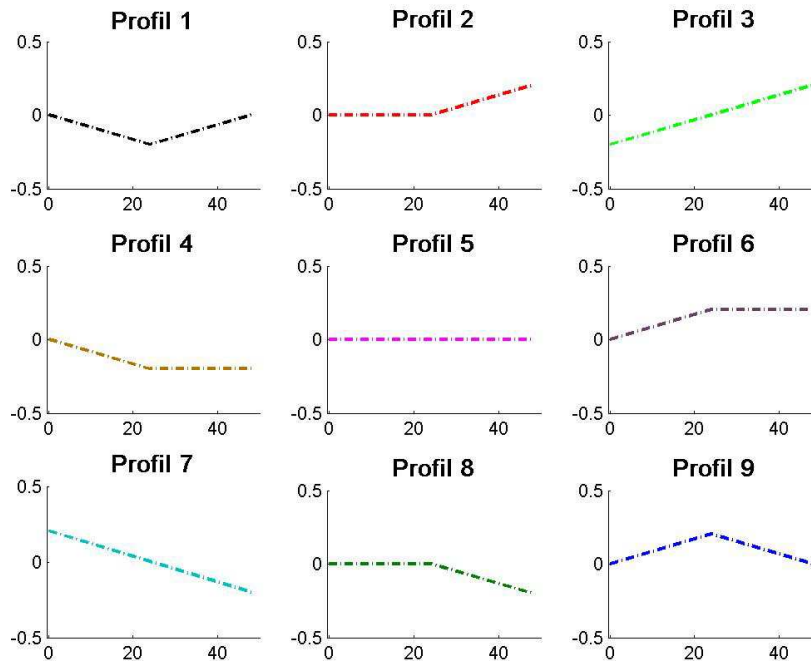


FIGURE 4.12 – Balayage des différents profils

Partie 2 :

Imagerie médicale : segmentation d'images de Tomographie par Emission de Positons

Dans cette partie, nous proposons une méthode de segmentation d'images TEP basée sur la cinétique du radiotraceur dans l'organisme et sur le spectral clustering. Nous commençons par décrire le principe de l'imagerie TEP ainsi que la problématique. Puis, après un état de l'art, nous présentons le spectral clustering pour une segmentation automatique d'images dynamiques. Nous étudions sa validation par des expérimentations numériques, des critères et une comparaison avec la méthode référence du k -means.

4.5 L'imagerie Tomographique par Emission de Positons

Nous présentons ici succinctement le principe physique de l'imagerie Tomographique par Emission de Positons (TEP), ainsi que la quantification de cibles moléculaires par modélisation compartimentale [115].

Principe de la TEP

La Tomographie par Emission de Positons (ou TEP) est une technique d'imagerie médicale fonctionnelle quantitative permettant de visualiser les activités du métabolisme [64]. Elle permet de mesurer in vivo et de manière a-traumatique pour le patient la distribution volumique au sein des organes d'un paramètre physiologique comme par exemple le métabolisme cellulaire ou la densité de récepteurs d'un système de transmission neuronal.

La TEP repose sur le principe général de la scintigraphie qui consiste à injecter un traceur dont on connaît le comportement et les propriétés biologiques pour obtenir une image du fonctionnement d'un organe. Ce traceur est marqué par un atome radioactif (en général le fluor 18, ^{18}F) qui émet des positons dont l'annihilation produit elle-même deux photons comme le montre la figure 4.13 (a). La détection de la trajectoire de ces photons par le collimateur de la caméra TEP est réalisée par un ensemble de détecteurs répartis autour du patient (cf figure 4.13(b)), par effet photoélectrique [101]. La reconstruction tomographique permet ensuite de localiser le lieu de leur émission et donc la concentration du traceur en chaque point de l'organe. C'est cette information quantitative que l'on représente sous la forme d'une image dont les niveaux de gris reflètent la concentration du traceur. Depuis le milieu des années 1990, le radiomarqueur ou traceur le plus connu et le plus utilisé est le glucose marqué au fluor 18. L'imagerie TEP du métabolisme du glucose permet la détection et la localisation des tumeurs et le suivi des patients après la mise en oeuvre des traitements en oncologie. La faisabilité de cette technique repose sur le fait que les cellules cancéreuses ont un métabolisme du glucose plus important que les cellules des tissus, et vont donc apparaître dans l'image comme des régions présentant une forte radioactivité. L'imagerie TEP est une source d'information constamment enrichie par la production de nouveaux radiotraceurs spécifiques de processus d'intérêt. Dans l'étude de la maladie d'Alzheimer par exemple, de nouveaux radiotraceurs permettent de visualiser la quantité de plaques beta-amyloïdes et de dégénération neuro-fibrillaire, qui sont les caractéristiques physiologiques de cette pathologie.

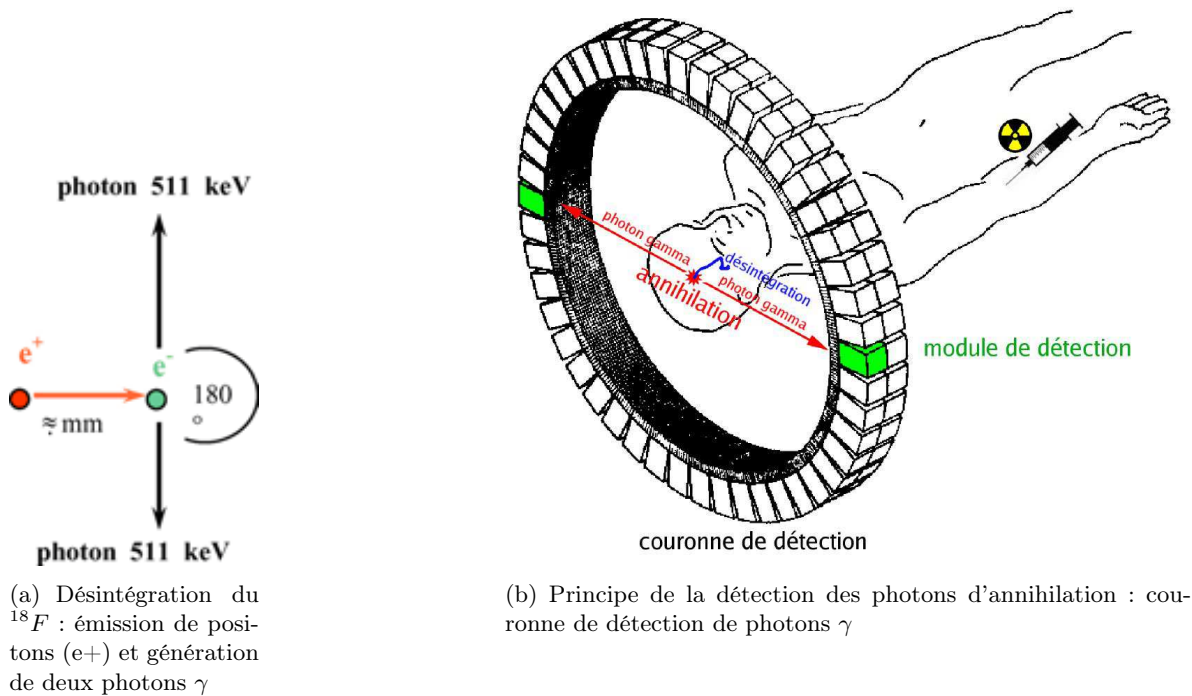


FIGURE 4.13 – Principe de la tomographie par émission de positons

Modélisation compartimentale : présence de cibles moléculaires et cinétiques des traceurs

La quantification en TEP nécessite l'injection du radiotracer dans le sang, et l'enregistrement dynamique de sa distribution. La cinétique de la distribution du radiotracer dépend de son interaction avec l'organisme :

- le passage des barrières tissulaires et cellulaires,
- son affinité pour les cibles spécifiques et non-spécifiques,
- le métabolisme.

Les courbes temps-activité (TAC) calculées à partir de volumes d'intérêt (VOIs) définis sur l'image sont alors utilisées pour résoudre les équations d'un modèle compartimental pour quantifier la cible.

Pour obtenir une quantification de la captation cérébrale, l'acquisition TEP démarre dès l'injection du radiotracer. Une fois les phénomènes dégradant le signal corrigés (temps mort des détecteurs, décroissance physique du radiotracer, phénomènes de diffusion, coïncidences aléatoires...), les données TEP représentent la concentration du traceur (en Bq/ml) pour un instant donné. Pour interpréter ces données TEP, on procède alors à une modélisation pour laquelle on suppose que la substance du traceur se répartit dans des compartiments physiologiquement séparés. Par exemple, dans un modèle à quatre compartiments, on considère un compartiment pour le plasma (sang artériel), un compartiment dit "libre", un compartiment de fixation non-spécifique et enfin un compartiment de fixation spécifique. Ce dernier compartiment est celui dont on cherche à connaître la concentration car elle est directement liée à la quantité de la cible que l'on cherche à évaluer (e.g. la quantité de plaques amyloïdes dans la maladie d'Alzheimer).

Pour estimer simultanément les différents paramètres du modèle compartimental et obtenir la concentration de la cible en un volume donné du cerveau, on s'appuie sur la cinétique de l'acti-

tivité dans le temps mesurée dans ce volume [112]. On parle de Courbe Temps-Activité (TAC), dont on peut voir un exemple dans la figure 4.14 pour des régions d'intérêt anatomiques [76].

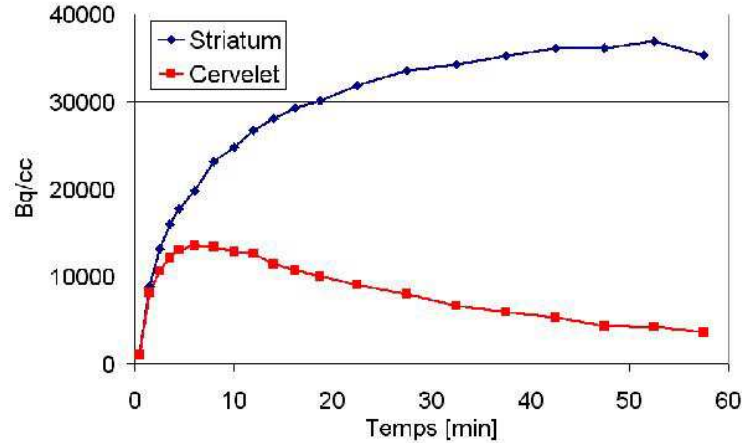


FIGURE 4.14 – Courbes temps activité de deux régions cérébrales

Dans une TAC, la hauteur et la localisation temporelle du pic, la cinétique de croissance et décroissance et l'éventuelle stabilisation sont autant d'indices transcrivant la vitesse de perfusion du volume considéré, la quantité de radiotracer présente ainsi que la quantité de cibles. Il s'agit d'étudier comment fluctue l'activité pour estimer la concentration de la molécule cible du radiomarqueur dans les différentes régions du cerveau.

Nature des données

Une fois acquises, les données TEP sont corrigées et reconstruites en images. Elles forment une séquence d'images $3D + t$ qui traduit l'évolution de la radioactivité dans le temps du volume correspondant au champ de vue de l'appareil. La résolution spatiale d'un imageur TEP est voisine de 5 mm. Les pixels sont généralement cubiques avec une taille de l'ordre de 2 mm par dimension. Un pixel de la matrice 3D constituant une image TEP correspond donc à un volume de 8 mm^3 dans le cerveau étudié. La valeur de l'image en ce pixel est proportionnelle à la radioactivité présente dans le volume.

4.6 Classification en TEP dynamique

Malgré la résolution spatiale limitée de la TEP, la précision avec laquelle les VOIs sont définis est essentielle pour obtenir des TACs réalistes et pour extraire de l'information pertinente du modèle. Intuitivement, un VOI correspond à une région ou une structure anatomique. Cependant, le contraste dans les images TEP est d'origine purement moléculaire et l'information anatomique doit être obtenue à l'aide d'une autre source. Cela peut être soit une image anatomique obtenue par une autre technique d'imagerie (CT, IRM...) qui est recalée avec l'image TEP, soit de l'information a priori par les connaissances en anatomie de l'opérateur. Chacune de ces méthodes présente des désavantages. Le recalage avec une image anatomique nécessite un second système d'imagerie, et comporte de l'imprécision liée à la méthode de recalage utilisée et aux artefacts de mouvements, difficilement compensables. Et surtout, les images anatomiques ne sont pas nécessairement pertinentes

pour délimiter des régions fonctionnellement différentes au sein d'un organe. C'est notamment le cas dans le foie, les reins ou encore le cerveau. D'un autre côté, la définition manuelle de VOIs est opérateur-dépendante, subjective et non facilement reproductible, ce qui peut entraîner des erreurs dans les paramètres du modèle. Enfin cette méthode manuelle est difficilement applicable à un grand volume de données à traiter. Les méthodes automatiques pour la classification en TEP dynamique ont donc fait l'objet d'une attention particulière. Plusieurs études ont proposé d'utiliser des méthodes non-supervisées pour délimiter les VOIs dans les images TEP, sans a priori anatomique, pour dégager des régions dont les TACs sont homogènes. On peut classer les méthodes de classification proposées dans la littérature traitant de la classification des images TEP en trois domaines suivants.

Le premier domaine regroupe les méthodes de classification utilisant la méthode k -means. La méthode k -means est introduite par Wong et al [112] pour traiter des images TEP dynamiques à partir des distances entre les points. Le k -means est par ailleurs utilisé dans les méthodes de classification basée sur une réduction de l'espace avec l'analyse en composantes principales (ACP) ou l'analyse factorielle. Kimura et al. [67] proposent d'appliquer k -means sur la projection des TACs sur les premiers vecteurs propres de l'ACP. Cependant, l'ACP est appliquée sur des données simulées sans bruit. Frouin et al. [46] appliquent le k -means sur les facteurs obtenus par une analyse factorielle. La méthode k -means à noyaux a aussi été étudiée par Acton et al. [1] dans le cadre des c -means flous pour prendre en compte la nature intrinsèque floue de la Tomographie par Emission Mono-Photonique (TEMP) et plus généralement dans le cadre de la segmentation de structures cérébrales par Farinha et al. [39].

Le deuxième domaine englobe les méthodes basées sur des modèles probabilistes. Ashburner et al. [8] utilisent l'hypothèse que les données TEP satisfont une distribution gaussienne et proposent un algorithme maximisant la vraisemblance d'un pixel appartenant à un cluster à partir de la TAC du pixel. Une segmentation par Espérance-Maximisation avec incorporation des chaînes Markov a été proposée par Chen et al. [24]. Brankov et al. [21] ont défini une nouvelle métrique de distance pour les TACs des pixels et utilisé cette métrique dans une méthode de classification par Espérance-Maximisation. Krestyannikov et al. [70] séparent les données TEP dynamiques par une méthode paramétrique basée sur les moindres carrés dans l'espace du sinogramme, évitant ainsi les artefacts que peut apporter la reconstruction de l'image.

Enfin, le dernier domaine traite des méthodes hiérarchiques. Une méthode de classification hiérarchique des voxels a été proposée par Guo et al [50]. Maroy et al. [75] ont proposé une méthode de classification basée sur plusieurs pré-traitements de l'image pour identifier les régions homogènes connexes, suivis d'une méthode hiérarchique ascendante pour regrouper les régions dont les TACs moyennes sont identiques.

Les méthodes proposées dans la littérature souffrent d'une ou plusieurs des limitations suivantes : impossibilité de séparer des domaines non convexes, sensibilité à l'initialisation de centres, choix de la métrique, paramètres d'échelle ajustés non automatiquement, connaissances a priori nécessaires sur la distribution des données, complexité calculatoire prohibitive pour de grands jeux de données, hypothèse implicite de clusters de tailles homogènes, sélection a priori du nombre de classes à trouver.

4.7 Segmentation non supervisée basée sur la dynamique des voxels

Pré-filtrage des données

Pour diminuer l'influence du bruit contenu dans les images TEP sur le résultat de la classification, nous effectuons un pré-filtrage de l'image 4D par un noyau gaussien :

$$G(x, y) = \frac{1}{\sqrt{2\pi}\theta^2} e^{-\frac{x^2+y^2}{2\theta^2}} \quad (4.3)$$

où θ est l'écart-type de la distribution gaussienne. Nous avons utilisé un filtre de taille 7×7 , avec θ calculé pour correspondre à la moitié de la taille de la fenêtre du filtre. Il est donc fixé à $\theta = 7/(4 * \text{sqr}(2 * \log(2)))$.

Prise en compte de l'effet de volume partiel

Les effets de volume partiel (EVP) sont caractérisés par la perte de signal dans les tissus ou organes de taille comparable à la résolution spatiale. Les EVP se manifestent par la contamination croisée d'intensité entre des structures adjacentes ayant des concentrations de radioactivité différentes. Dans ce dernier phénomène, parfois dénommé "spill-over", l'activité haute d'une région donnée peut se répandre et contaminer une région proche et d'activité moindre, ce qui conduit à des mesures surestimées ou sous-estimées des activités réellement présentes dans ces régions. Les techniques actuelles de correction des effets de volume partiel font systématiquement appel à une connaissance a priori de l'anatomie [93]. Nous proposons de ne pas corriger directement les voxels de l'EVP par une méthode de la littérature, mais plutôt de normaliser les vecteurs x_i avant le calcul de la matrice d'affinité. La conséquence est que la classification spectrale que nous proposons ne tient pas compte de l'amplitude de l'intégrale des TACs des voxels. Elle se base uniquement sur la forme des TACs. Cela permet de réduire l'influence de l'EVP qui contamine fortement les TACs lorsque deux régions voisines présentent un écart d'activité important. Sans cette normalisation, les voxels entre les deux régions ont leur TACs qui varient en amplitude et cela perturbe fortement le calcul des distances entre les x_i . Les résultats obtenus sans cette normalisation préalable confirment l'importance de cette étape (voir section résultat de ce chapitre).

Nous proposons donc d'appliquer la méthode de spectral clustering sur les données normalisées, puis d'effectuer un post-filtrage du résultat de la classification pour prendre en compte les différences importantes en amplitude d'intégrale de TACs entre deux régions dont le profil temporel des TACs serait proche à une homothétie près.

Choix du paramètre

Le paramètre est choisi tel qu'il vérifie l'heuristique (1.2) développée dans le chapitre 1. Rappelons que ce paramètre est global et donne une information de densité sur la distribution des points. Dans ce cadre, il fournit un profil caractéristique de la distribution de l'ensemble des profils de TACs.

Estimation du nombre de clusters

De nombreux critères existent pour évaluer le nombre de clusters. Dans le cas de méthodes comme k -means, l'erreur moyenne quadratique (MSE) [75] permet de calculer, suivant une norme pondérée dont les poids sont définis par l'utilisateur, l'écart moyen entre les points d'un même

cluster et son centre associé. Cependant, ce critère étant décroissant quand le nombre de clusters augmente, un seuil défini a priori est requis pour définir un nombre optimal de clusters. D'autres critères comme le critère d'information d'Akaike [3] ou le critère de Schwarz [89] sont basés sur des modèles probabilistes et nécessitent d'émettre des hypothèses quant aux lois de probabilité (en général gaussienne) que les TAC doivent suivre. Dans cette étude, nous utilisons le critère basé sur les normes de Frobenius (3.5) développé dans le chapitre 3 pour définir le nombre de clusters k .

4.8 Validation par simulation de courbes temps-activité

Pour estimer la concentration de la molécule cible du radiomarqueur dans les différentes régions de cerveau, nous nous intéressons à la segmentation de l'image TEP 3D par classification des TACs associées aux pixels. L'objectif étant de regrouper dans une même classe les pixels de l'image correspondant à des comportements pharmaco-cinétiques proches.

4.8.1 Simulation TACs

Les TACs associées aux différentes régions d'intérêt sont générées à partir de l'équation cinétique suivante :

$$Ca_t = A_0 ((A_1 t - A_2 - A_3)) e^{-\lambda_1 t} + A_2 e^{-\lambda_2 t} + A_3 e^{-\lambda_3 t}, \quad (4.4)$$

avec $A_0 \in [1e4, 3e5]$, $A_1 \in [0, 0.8]$, $A_2 \in [0, 1 - A_1]$, $A_3 = 1 - A_1 - A_2$ représentant des coefficients d'amplitudes et où $\lambda_1 \in [30, 45]$, $\lambda_2 \in [\lambda_1, 180]$, $\lambda_3 \in [\lambda_1, 180]$ sont des valeurs temporelles en minutes. Cette équation (4.4) est définie à partir de l'estimation de l'activité fonctionnelle dans le cerveau à travers un modèle du traceur cinétique flurodeoxyglucose (FDG) [75, 40]. Il s'agit d'un modèle différentiel composé de 3 fonctions : la concentration du traceur dynamique FDG dans le plasma/sang puis dans le tissu et enfin la concentration du FDG-6-phosphate dans le tissu. A partir de (4.4), les TAC initiales définies pour chaque compartiment du cerveau pour les 9 régions (cf figure 4.15) à localiser sont représentées sur la figure 4.16.

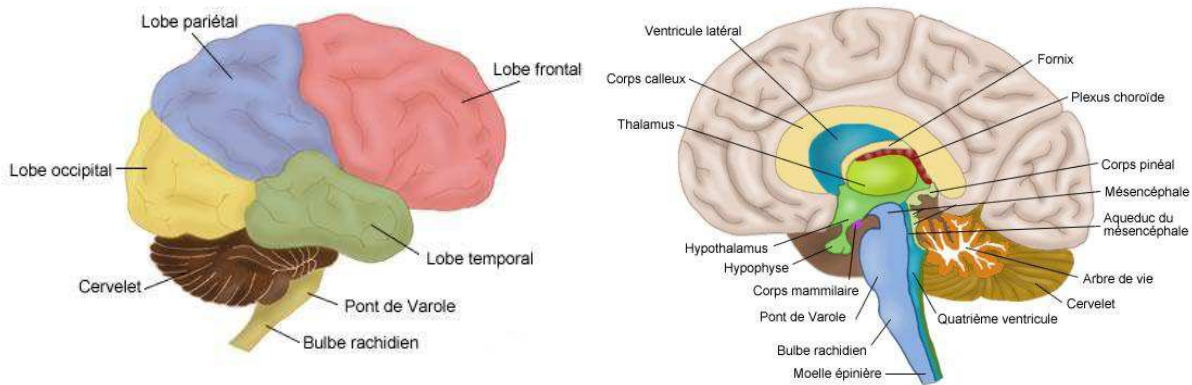


FIGURE 4.15 – Différentes régions du cerveau

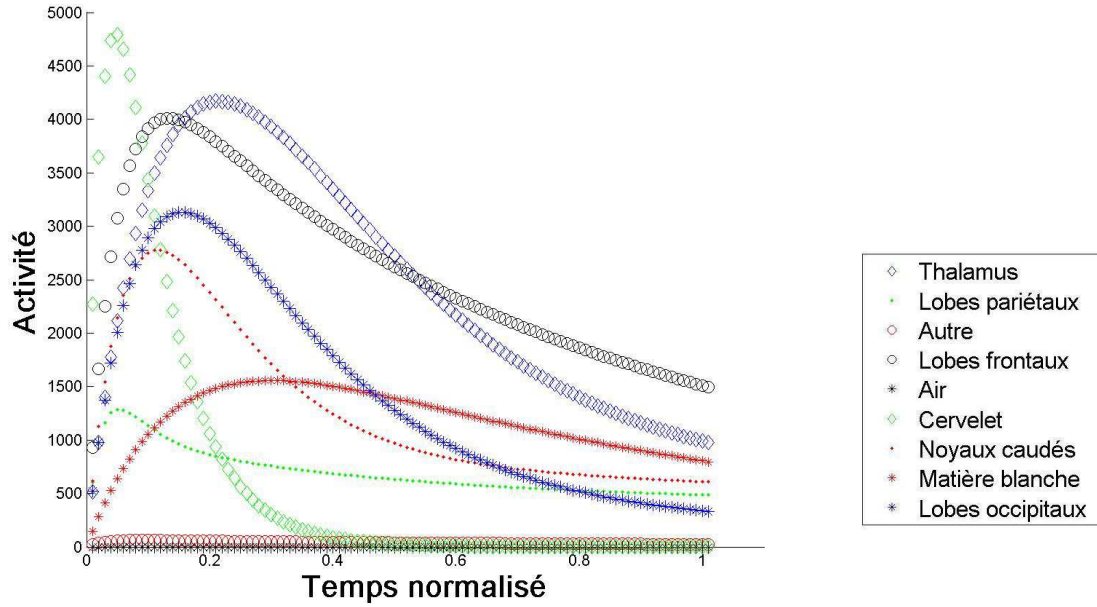


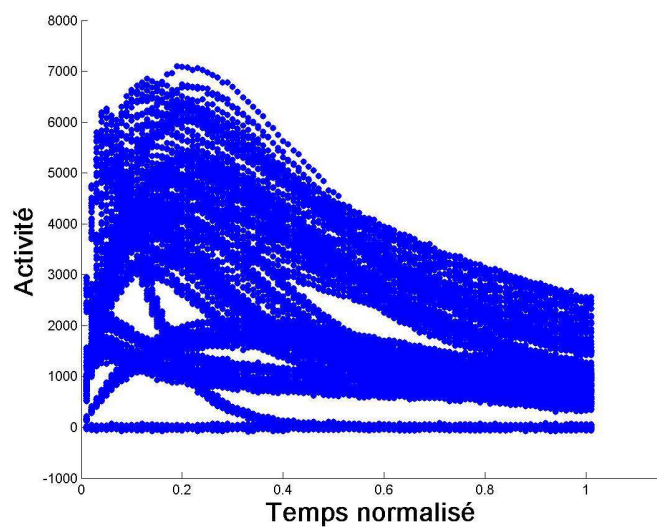
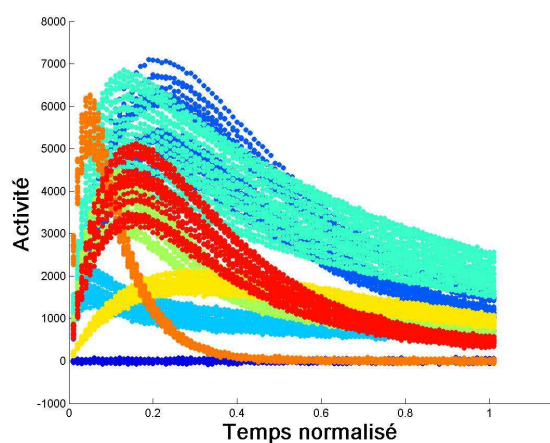
FIGURE 4.16 – Exemple de TACs affectées à différentes régions cérébrales

4.8.2 Spectral clustering sur TACs

Etant donné ce nouveau contexte où les données sont multidimensionnelles, le spectral clustering doit être étudié pour valider les heuristiques développées dans les chapitres précédents. Par conséquent, un contexte de cadre réel est défini à partir de l'équation (4.4). Nous perturbons donc en amplitude ces courbes théoriques et nous y ajoutons du bruit blanc via des lois normales. Autrement dit, soit $S = \{x_i\}_{i=1..N} \subset \mathbb{R}^p$ où $x_i \in \mathbb{R}^p$ est une courbe TAC définie en p instants vérifiant (4.4) telle que :

$$\forall j \in [1, p], x_i(j) = (1 + \varepsilon)x_i(j) + \beta \quad (4.5)$$

avec $\varepsilon \propto \mathcal{N}(0, 1)$ et $\beta \propto \mathcal{N}(\bar{m}, \mu)$ où $\bar{m}$ et μ sont fixés en fonction de la moyenne et de l'écart-type de la distribution des points x_i . A partir de ces équations, l'ensemble des profils x_i de S pour $i \in [1, n]$ vérifiant (4.5) est représenté sur la figure 4.17. L'algorithme 1 de spectral clustering défini dans le chapitre 1 est appliqué sur l'ensemble de données S de la figure 4.17 avec pour paramètre σ d'affinité gaussienne l'heuristique (1.2). La figure 4.18 représente sur la gauche le résultat du spectral clustering et, sur la droite, l'espace de projection spectrale. Rappelons que chaque couleur représente un cluster à savoir un profil de TAC différent. Le partitionnement des données de S issu du spectral clustering correspond exactement au nombre de courbes TAC x_i associées au même profil TAC généré. De plus, les points dans l'espace de projection spectrale associés au même cluster sont distincts de ceux d'un cluster différent. Ce qui signifie que le paramètre de l'affinité gaussienne et les données originales permettent une distinction entre les profils de TACs. Cependant, il faut vérifier la robustesse du choix du paramètre σ et l'heuristique pour déterminer le nombre de clusters.

FIGURE 4.17 – Ensemble des profils de S 

(a) Résultat du spectral clustering

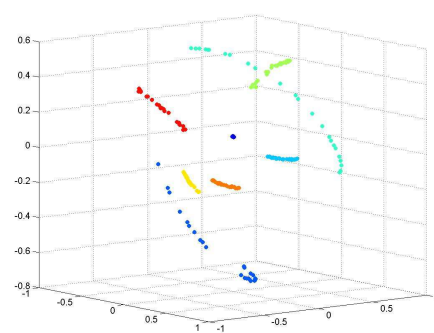
(b) Espace de projection spectrale dans $\mathbb{R}^3$

FIGURE 4.18 – Application du spectral clustering sur des TACS bruitées et perturbées en amplitude

Choix du paramètre σ

Dans ce nouveau contexte où les données sont multidimensionnelles avec $p > 10$, les deux mesures de qualité introduites dans le chapitre 1 sont utilisées pour valider le choix du paramètre : la première basée sur la matrice de confusion estime le pourcentage de points mal classés et la seconde, basée sur les ratios de normes de Frobenius, évalue le rapport r_{ij} entre l'affinité entre les clusters et l'affinité intra-cluster. Ces deux mesures sont calculées en fonction des valeurs de σ pour un nombre de classes k égal à 8, le nombre de profils différents générés par l'équation (4.5) et sont représentées dans la figure 4.19. Dans les deux cas, la valeur de sigma définie par (1.2) est indiquée par une ligne noire verticale en pointillée. Pour cette valeur, le critère est nul, validant la pertinence de l'heuristique. Par conséquent, dans la suite, le choix du paramètre de l'affinité gaussienne est définie par l'heuristique (1.2) du le chapitre 1.

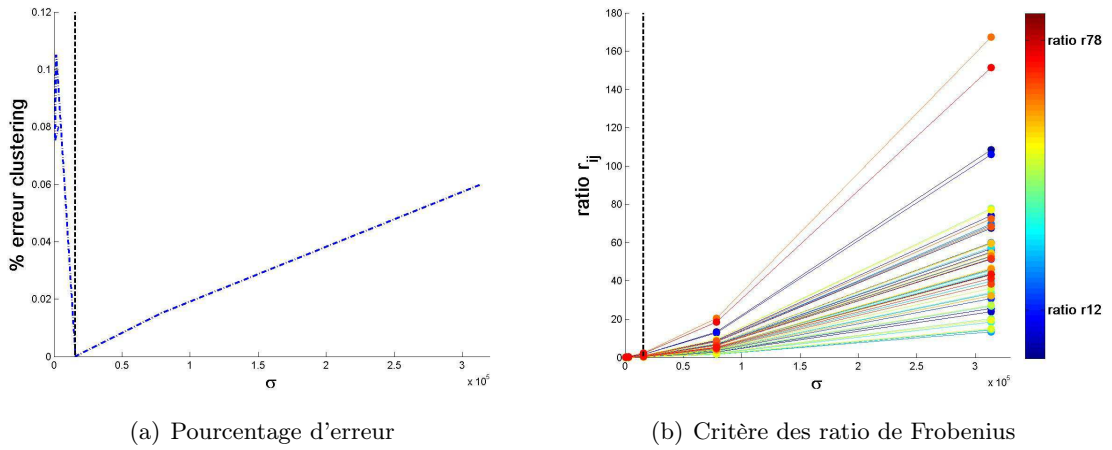


FIGURE 4.19 – Mesures de qualité en fonction des valeurs de σ

Estimation du nombre de clusters

D'après l'heuristique (3.5) développée pour le choix du nombre de clusters k dans le chapitre 3, le ratio η_k fonction de la mesure de qualité basée sur les ratios de normes de Frobenius est tracé en fonction du nombre de clusters $k \in [2, 9]$ et est représenté sur la figure 4.20. Rappelons que le nombre de clusters k final est la valeur de k pour laquelle le ratio η_k est minimal. D'après la courbe, ce minimum est obtenu pour $k = 8$ correspondant au nombre exact de profils définis par l'équation (4.5). Donc ce critère du nombre de clusters (3.5) est conservé.

4.9 Validation par simulation d'images TEP réalistes

Dans la suite, nous présentons le processus pour simuler des images TEP réalistes puis la méthode de spectral clustering sera étudiée avec les estimations automatiques des paramètres d'affinité gaussienne σ et du nombre de clusters k définis précédemment puis comparée à la méthode du k -means, méthode la plus utilisée.

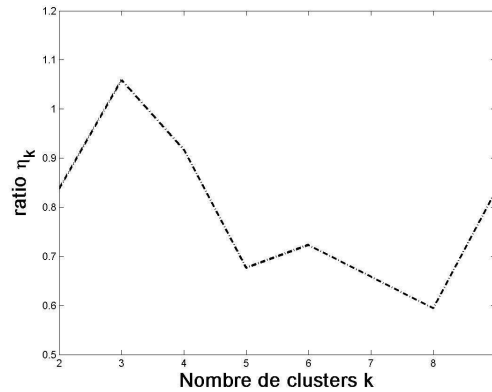


FIGURE 4.20 – Estimation du nombre de clusters

4.9.1 Simulation des images TEP réalistes

La simulation numérique est couramment utilisée dans le domaine du traitement d'image. Elle constitue une aide précieuse pour le développement et l'évaluation de méthodes car elle permet de disposer d'une vérité terrain à laquelle on va comparer les résultats obtenus.

Dans l'imagerie médicale, ces simulations sont généralement effectuées selon la méthode Monte-Carlo, fondée sur les théories des probabilités, particulièrement adaptées à la physique médicale à cause de la nature stochastique des processus d'émission, de transport et de détection. Pour reproduire aussi fidèlement que possible les différents phénomènes physiques impliqués dans l'acquisition réelle d'images TEP, nous nous sommes appuyés sur la plateforme de simulation GATE (Geant4 Application for Tomographic Emission), qui est l'outil de référence actuel. Cette plate-forme de simulation Monte-Carlo générique est dédiée à la simulation d'applications en SPECT et TEP.

Cette plateforme est construite à partir de la bibliothèque GEANT4 [2] développée par le CERN pour la simulation du transport et de l'interaction de particules dans la matière. La plateforme GATE permet de modéliser des processus physiques dépendants du temps, de reproduire la physique d'un imageur TEP (cf figure 4.21) et d'y introduire virtuellement un fantôme numérique qui représente le patient subissant l'examen [5].

Un fantôme numérique du corps humain est une description géométrique virtuelle de la morphologie de l'organisme. Il existe principalement deux types de représentations. D'une part la représentation des structures de l'organisme comme un ensemble d'objets de forme géométrique de complexité variable (e.g. des ellipses, des B-splines généralisées NURBS...). D'autre part, les fantômes numériques provenant de la segmentation d'examens anatomiques réels (TDM ou IRM) avec une labellisation de chacun des voxels de l'image. Pour ce travail autour du cerveau, nous avons choisi d'utiliser le fantôme de tête proposé par Zubal [117]. Il correspond à la description d'une IRM de tête, où chaque voxel est associé au label de la région du cerveau d'appartenance. A chaque label, nous associons une TAC virtuelle, représentée sur la figure 4.16, et donc une valeur de la concentration du radiotraceur exprimé en Becquerel par millilitre (Bq/ml), ainsi qu'une valeur d'atténuation du tissu lui correspondant. A partir des ces trois types d'information (le fantôme, l'activité du radiotraceur, et la valeur d'atténuation), le simulateur GATE (figure 4.22) permet de reproduire les différents phénomènes d'une acquisition TEP réelle comme les processus statistiques de désintégrations radioactives, d'interaction entre particules et de détection de photons.

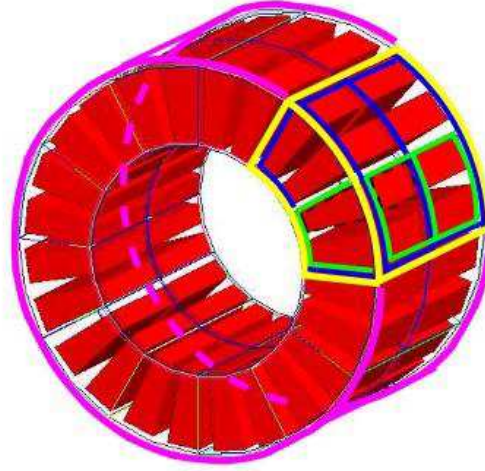


FIGURE 4.21 – Simulation d'un imageur TEP dans GATE

4.9.2 Critères d'évaluation quantitative

Pour valider notre méthode, nous nous appuyons sur deux critères complémentaires : l'index d'inclusion et l'index de similarité [75].

Index d'inclusion

La précision de la classification est évaluée pour chaque organe par l'index d'inclusion, défini pour une classe donnée comme :

$$I_{inc} = \frac{\text{card}(\mathcal{A}_{\text{spectral}} \cap \mathcal{A}_{\text{gtruth}})}{\text{card}(\mathcal{A}_{\text{gtruth}})}, \quad (4.6)$$

où card est la fonction de cardinalité, $\mathcal{A}_{\text{spectral}}$ est la classe résultant de l'application de notre méthode et $\mathcal{A}_{\text{gtruth}}$ est la vérité terrain que l'on cherche à obtenir.

Ce critère n'est pas symétrique, et donne un score maximum lorsque l'ensemble constituant la vérité terrain est inclus dans l'ensemble formé par la classe résultant de la méthode proposée, même si la cardinalité de cette dernière est largement supérieure à celle de la vérité terrain. Nous considérons toutefois ce critère, car il permet de réduire l'écart implicite dû aux fortes différences de résolution des deux modalités. L'imagerie TEP simulée et l'image IRM du fantôme numérique sont respectivement formées de voxels dont le côté mesure 4mm et 1.1mm. La finesse de la définition des régions est par conséquent bien meilleure en IRM, et un index classique ne prend pas en compte cette différence. L'index d'inclusion donne une bonne indication sur la part de la région vérité terrain qui est incluse dans la région résultante de notre méthode. Il doit cependant être accompagné d'un critère qui prend aussi en compte l'écart de taille entre les deux classes que l'on compare. Il s'agit du critère que nous présentons dans le paragraphe suivant.

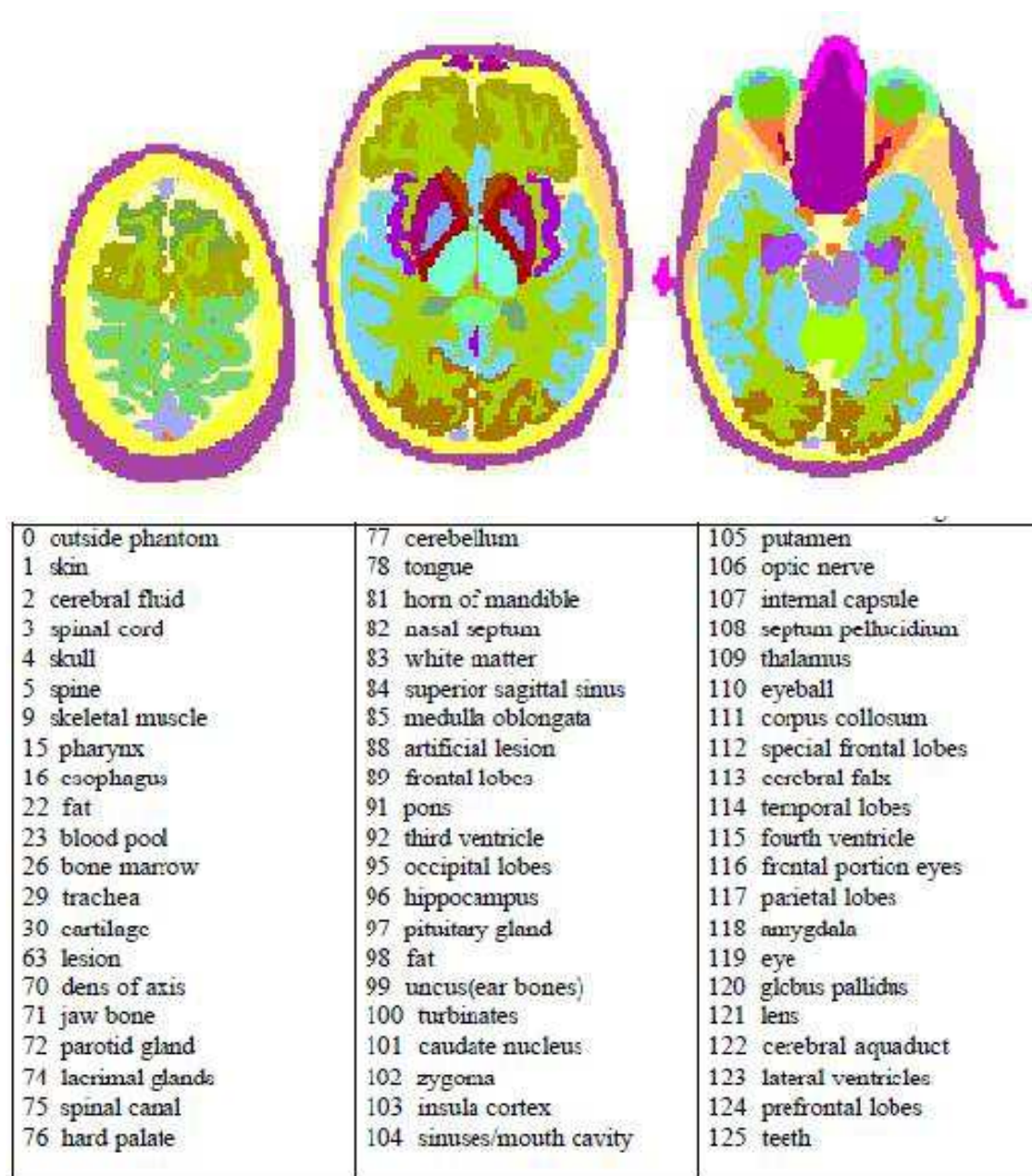


FIGURE 4.22 – Fantôme de cerveau Zubal utilisé pour la simulation

Index de similarité

Pour prendre en compte de façon symétrique les tailles des régions comparées, nous calculons l'index de similarité suivant :

$$I_{sim} = \frac{\text{card}(\mathcal{A}_{\text{spectral}} \cap \mathcal{A}_{\text{gtruth}})}{\text{card}(\mathcal{A}_{\text{spectral}} \cup \mathcal{A}_{\text{gtruth}})}. \quad (4.7)$$

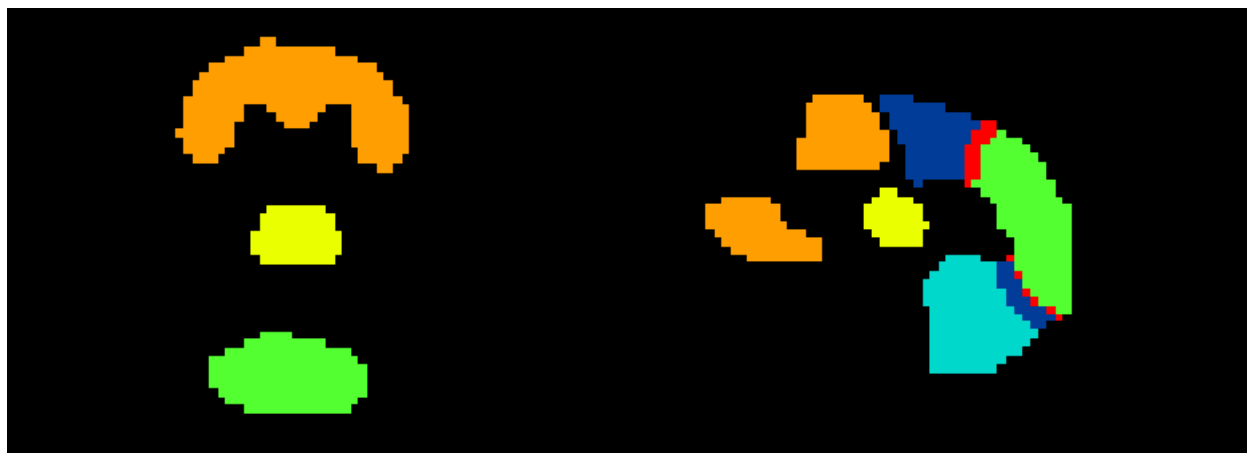
Ce critère est à la fois sensible à la différence de la taille des régions comparées et à leur position relatives. Il sera par contre implicitement affecté par la différence de résolution des deux modalités et donc moins élevé que l'index d'inclusion.

4.9.3 Résultats du spectral clustering

Pour cette étude, nous ne considérons pas directement les données 3D temporelles mais deux tranches 2D temporelles correspondant à une coupe sagittale et une coupe transverse du cerveau. Ainsi l'information de volume, à savoir les données réparties dans tout l'espace, est conservée. Rappelons que la méthode de spectral clustering est appliquée sur les données temporelles filtrées par l'équation (4.3) puis normalisées. Enfin un post-filtrage du résultat de la classification est effectué pour prendre en compte les différences importantes en amplitude d'intégrale de TACs entre deux régions dont le profil temporel des TACs serait proche à une homothétie près. Les TACs définies par l'équation (4.4) sont associées aux 5 régions cérébrales suivantes : le lobe frontal, le lobe occipital, le thalamus, le lobe pariétal, le cervelet. Les clusters sont ensuite localisés dans le cerveau suivant leurs coordonnées géométriques. Certains clusters décrivent alors des régions particulières du cerveau.

Le résultat du spectral clustering est reconstitué suivant les coupes du cerveau et représenté sur la figure 4.23(a). Avec l'heuristique (3.5) pour déterminer le nombre de clusters, $k = 6$ clusters ont été déterminés après le post-filtrage. Pour juger des contours des régions du cerveau, la solution, décrite en couleurs, est superposée au résultat du spectral clustering, décrite en niveaux de gris, sur la figure 4.23 (b). Le résultat du clustering est aussi superposé à l'image IRM pour situer les clusters aux zones cérébrales. On observe que les 5 régions sont correctement délimitées. Les régions de la vérité terrain sont plus petites en raison de la différence de résolution spatiale entre la vérité de l'image IRM et la simulation d'images TEP. Un sixième cluster caractérise le reste de l'image et les zones de "transition" entre le lobe pariétal, le lobe occipital et le cervelet. Ce cluster est imputable à l'effet de volume partiel [93]. Cette zone de transition présente une courbe TAC qui est un mélange des TACs caractéristiques des lobes pariétal et occipital et du cervelet. Le résultat du spectral clustering sur les TACs normalisées est affiché sur la figure 4.24. Les profils des 5 régions sont compacts et décrivent un profil caractéristique. Le dernier cluster en bleu regroupe des TACs principalement affectées à l'effet de volume partiel ainsi que le fond de l'image, d'activité constante.

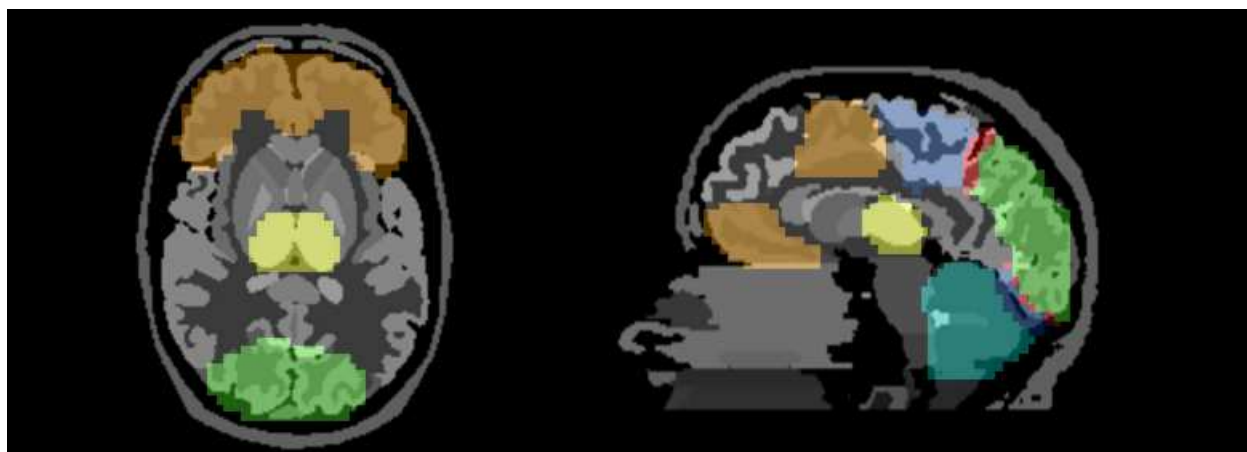
Les critères basés sur l'index de similarité et d'inclusion sont calculés sur la partition issue du spectral clustering et indiqués sur le tableau 4.2. Ces indices indiquent que la forme des clusters et celle des régions cérébrales coïncident fortement entre 68% pour le lobe pariétal et 98% pour le thalamus. Le spectral clustering présente donc des résultats satisfaisants pour la segmentation d'images dynamiques TEP.



(a) Résultat du spectral clustering



(b) Superposition de la solution et du résultat du spectral clustering



(c) Fusion du résultat du spectral clustering avec l'IRM

FIGURE 4.23 – Résultat du spectral clustering pour $k = 6$

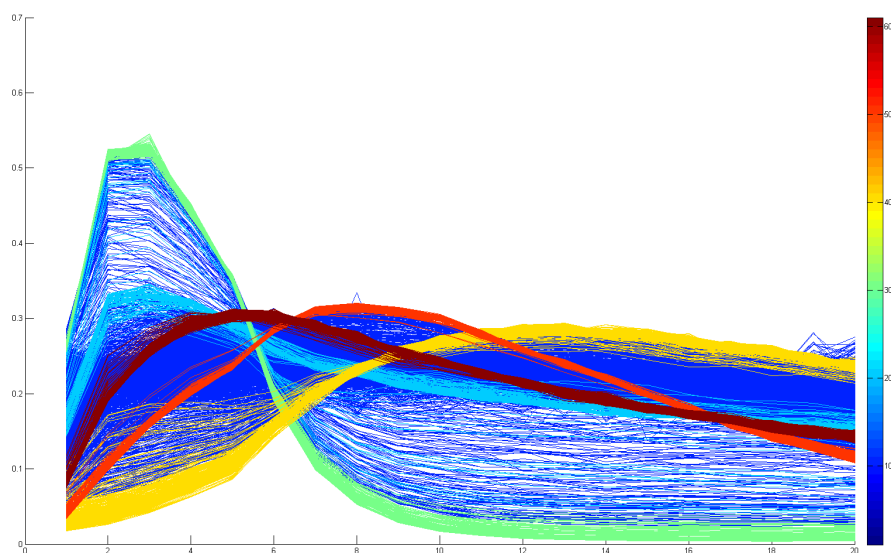


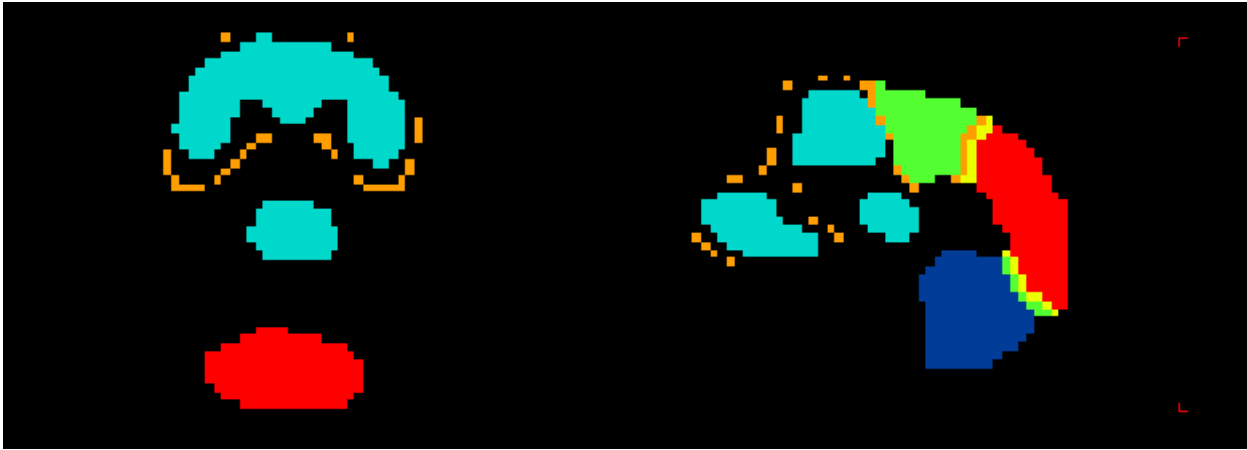
FIGURE 4.24 – Résultat du spectral clustering sur les TACs normalisées

Régions cérébrales	Index de similarité en pourcentage	Index d'inclusion en pourcentage
Cervelet	73.11	87.53
Lobe frontal	69.83	94.76
Lobe occipital	67.50	84.79
Thalamus	81.85	98.41
Lobe pariétal	44.48	68.07

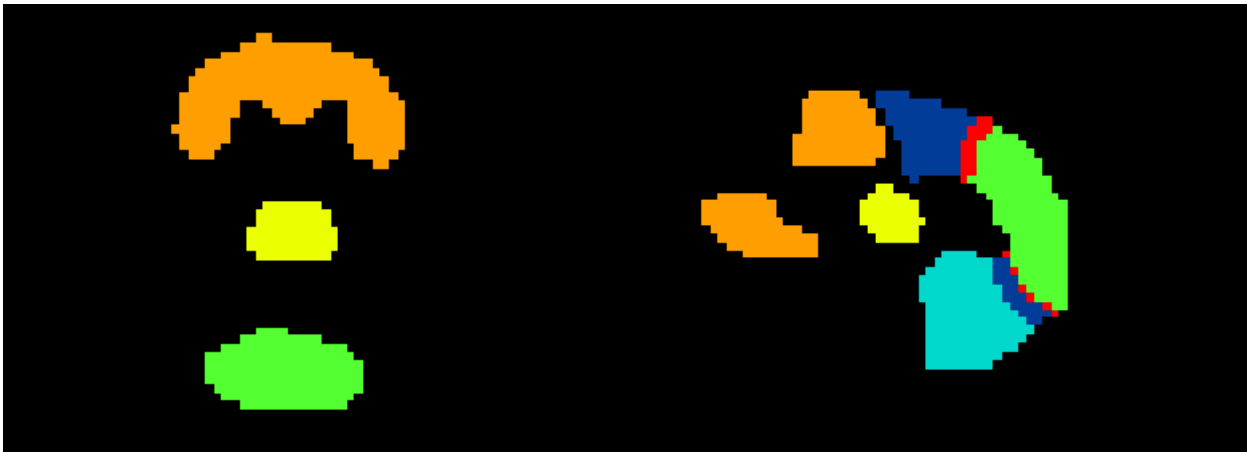
TABLE 4.2 – Critères d'évaluation quantitative

4.9.4 Comparaison avec la méthode k -means

Pour valider les résultats du spectral clustering, la méthode k -means, méthode la plus utilisée en segmentation d'images dynamiques TEP [75], est appliquée sur ce même exemple. Les mêmes pré-traitement et post-traitement sont effectués sur les données. La méthode k -means est donc appliquée sur des données TACs filtrées et normalisées. Le nombre de clusters est fixé à 6 et le résultat de la méthode est représenté sur la figure 4.25(a).



(a) Résultat du k -means



(b) Résultat du spectral clustering

FIGURE 4.25 – Comparaison avec la méthode k -means pour $k = 6$

Comparés aux résultats du spectral clustering sur la figure 4.25 (b), la plupart des formes des régions du spectral clustering sont quasi-identiques à celles du k -means. Par contre, un cluster à support entre le lobe frontal, le lobe pariétal et occipital (cluster orange sur la figure 4.25(a)) est créé par la méthode k -means. De plus, les régions du thalamus et du lobe frontal sont englobées en un seul cluster (cluster bleu sur la figure 4.25(a)). Les TACs associées à ces deux régions sont représentées sur la figure 4.26 et indiquent des profils à supports différents.

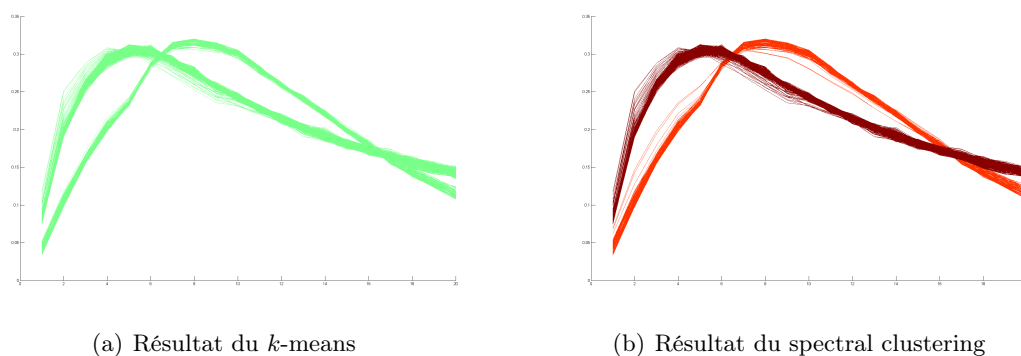


FIGURE 4.26 – Comparaison des TACs du thalamus et du lobe frontal

Sur cet exemple, la méthode de spectral clustering présente des résultats meilleurs que le k -means. Ces derniers résultats associés à ceux issus des critères d'évaluation quantitative prouvent que le spectral clustering est une méthode valable et prometteuse dans le domaine de la segmentation d'images TEP dynamiques. Mais une étude plus approfondie doit être menée, notamment sur le problème des effets de volumes partiels et son impact dans les résultats de clustering. Considérer l'image entière 3D permettrait de réaliser une étude plus précise sur la segmentation des TACs et donc des régions cérébrales. La stratégie parallèle développée dans le chapitre 3 s'inscrit comme une perspective logique dans ce contexte.

Conclusion et perspectives

Dans ce travail, nous avons étudié des éléments de la méthode de spectral clustering. L'objectif était de comprendre le fonctionnement de cette méthode afin de l'appliquer dans un cadre non supervisé sur de grands flots de données.

Une première partie s'est intéressée au choix des paramètres de la méthode. Etant donné qu'elle repose sur la mesure d'affinité gaussienne entre les points et son paramètre, une étude précise du paramètre de l'affinité gaussienne a été développée. Deux heuristiques basées sur la distribution des points ont donc été définies et une mesure de qualité permettant d'observer l'influence du paramètre de l'affinité gaussienne sur le résultat de la partition finale a été introduite. La particularité des heuristiques proposées est qu'elles demandent très peu de connaissance sur les données à utiliser et ainsi leurs évaluations ne pénalisent pas le coût global du calcul tout en donnant d'excellents résultats, y compris dans le cas de domaines non convexes emboîtés où d'autres approches échouent. La mesure de qualité, quant à elle, ne nécessite pas la connaissance a priori des clusters à retrouver (a contrario de la mesure d'erreur) et elle nous a permis de proposer une automatisation du processus de recherche du nombre de clusters uniquement à partir des résultats de clustering obtenus par la méthode. Des premiers tests sur du traitement d'images seuillées ont été réalisés. Ils ont montré la validité de nos propositions et exploré les limites numériques de cette méthode.

Cette étude du paramètre nous a permis d'améliorer les résultats mais il reste difficile de comprendre comment le spectral clustering fonctionne. En effet, ce paramètre conditionne de manière radicale la séparation des clusters dans l'espace de projection spectrale. Comme les éléments spectraux de la matrice affinité ne fournissent pas explicitement de critère topologique pour un ensemble discret de données, une formulation continue de la méthode a été définie dans une seconde partie de cette thèse. Ainsi, dans ce formalisme, les clusters représentent des sous-ensembles disjoints et l'affinité gaussienne est une discrétisation du noyau de la chaleur défini sur $\mathbb{R}^p$. Cette interprétation a soulevé des problèmes. Tout d'abord, le spectre de l'opérateur de la chaleur en espace libre est essentiel ce qui empêche l'interprétation directe en valeurs propres et vecteurs propres. Toutefois, un lien avec un problème d'équation de la chaleur dans un domaine borné a été défini. Par ailleurs, il est difficile d'établir une comparaison entre des données discrètes, à savoir l'affinité gaussienne entre les points, et les fonctions L^2 solutions propres de l'équation de la chaleur. L'introduction d'une discrétisation explicite via les éléments finis a permis cette comparaison entre données discrètes et continues. Enfin, cette interprétation via l'équation de la chaleur et les éléments finis doit être liée à l'analyse de la méthode de spectral clustering, c'est-à-dire permettre de déterminer l'appartenance d'un point à un cluster uniquement à partir des éléments spectraux étudiés. Pour ce faire, des propriétés de clustering ont été établies au fur et à mesure des approximations successivement réalisées et ont permis alors d'explicitier la séparation des données dans l'espace de projection spectrale. Ces interprétations ont permis de définir un intervalle de valeurs pour le paramètre de l'affinité gaussienne dans lequel les propriétés de clustering du continu sont vérifiées. Cette approche a été validée

numériquement.

Partant du lien précédemment établi entre les clusters, pris comme ensembles de points au sein d'une même composante connexe, et la double heuristique proposée dans la première partie, nous nous sommes aussi intéressés à la mise en oeuvre en grande dimension du spectral clustering. La connaissance de choix automatiques du paramètre affinité gaussien et du choix du nombre de clusters à trouver, via la mesure de qualité en norme de Frobenius validée en première partie, ainsi que la description géométrique du comportement du spectral clustering, nous ont en effet permis de proposer deux schémas de parallélisation de la méthode tout en restant non-supervisée. Le point clé de chacune de ces méthodes parallèles résidait dans la définition d'une zone "d'intersection" entre les sous-domaines (ou processus). Cette dernière proposée soit sous la forme d'un domaine supplémentaire appelé interface (méthode de spectral clustering avec interface), soit par introduction de recouvrement entre les sous-domaines adjacents (méthode de spectral clustering avec recouvrement). Le principe de la parallélisation étant alors d'effectuer des résolutions par spectral clustering sur chacun des sous-domaines ainsi définis, le résultat final était alors obtenu par assemblage des résultats obtenus à l'aide d'une relation d'équivalence spécifique. Il a enfin été démontré que, sans une hypothèse sur la taille de la zone "d'intersection", les deux schémas proposés donnaient exactement les résultats de clustering attendus.

D'un point de vue numérique, les expérimentations menées sur des exemples géométriques, pour les deux méthodes donnent des résultats très performants comparés à la résolution complète. Cependant, la stratégie par interface présente une limite lorsque le nombre de découpe augmente : l'interface concentre alors le plus de points causant une diminution de l'efficacité. La seconde méthode présente des résultats beaucoup plus performants que la première. Cette dernière, de plus, est alors testée sur des cas de segmentation d'images où l'on traite des données alliant à la fois géométrie et niveau de gris ou couleurs. Des comparaisons visuelles avec l'image originale sont menées et la stratégie semble prometteuse et peut, suivant les cas, compresser les informations inhérentes à l'image en quelques clusters.

Enfin, la dernière partie de cette thèse était consacrée à la validation de l'intérêt et de la pertinence du spectral clustering dans un contexte applicatif très spécifique. En effet, deux collaborations dans le monde de la génomique et de l'imagerie médicale ont permis de confronter le spectral clustering à un contexte réel où les données sont bruitées, nombreuses et où nous disposons de peu d'information sur les clusters à obtenir. La première application concerne la classification de profils temporels d'expression de gènes. La méthode de spectral clustering a été comparée à une méthode de cartes de Kohonen adaptée à ce type de données temporelles. Les résultats du spectral clustering sont exploitables mais moins satisfaisants qu'avec l'autre méthode. Cependant, elle a révélé un clustering d'une nature différente ouvrant une classification suivant le profil et l'amplitude. La seconde application traite de la segmentation d'images dynamiques de Tomographie par Emission de Positons (TEP). Appliquer le spectral clustering sur des données temporelles permet de segmenter des régions cérébrales comportant la même courbe de temps-activité. Les résultats ont été comparés avec succès à ceux de la méthode k -means et via des critères d'évaluation quantitative et s'avèrent pertinents et prometteurs. La méthode de spectral clustering a par conséquent prouvé avoir sa place dans l'imagerie TEP dynamique.

Perspectives

Comme on a pu le voir dans ce qui précède, le spectral clustering est une méthode assez complexe dans son interprétation mais son étude nous a toutefois permis de déboucher sur des nouveaux

schémas parallèles et de nouveaux domaines d'applications relevant proprement de la classification non supervisée. La complexité de cette discipline étant justement dans le caractère "aveugle" des méthodes à développer/utiliser, cela ouvre donc plusieurs perspectives à ce travail.

Tout d'abord, d'un point de vue théorique, nous avons pu montrer que l'algorithme du spectral clustering sans l'étape de normalisation revenait à rechercher, au sein de la matrice affinité, des vecteurs propres "proches" (en un sens que nous avons quantifié en fonction de la distribution des données et du choix du paramètre gaussien) d'une discrétisation conforme de fonctions propres d'un problème de la chaleur avec conditions de Dirichlet au bord. Ces dernières furent reliées à des propriétés dites "de clustering", c'est-à-dire permettant de déterminer l'appartenance à une composante connexe de tout point donné. Toutefois, le résultat nécessite une famille de vecteurs propres afin de pouvoir classer, sans détermination possible, tous les points d'un ensemble donné. De plus, cette famille n'est a priori ni finie ni limitée aux premiers vecteurs propres (ceux associés aux plus grandes valeurs propres). Les exemples numériques ont néanmoins montré que, dans le cas de l'algorithme de spectral clustering avec étape de normalisation, il était possible de se limiter aux k plus grands vecteurs propres pour avoir un critère déterminant exactement les clusters de l'ensemble à traiter. L'étape de normalisation semble donc déterminante pour simplifier la recherche des "bons" éléments spectraux à utiliser. Comme évoqué dans le chapitre 2, il semblerait qu'elle permette d'apparenter la méthode à la recherche de fonctions propres du problème de la chaleur avec conditions de Neumann. Ceci serait alors en accord avec le caractère constant par morceaux observé sur les vecteurs propres de la matrice affinité normalisée. Une première perspective serait donc de démontrer effectivement s'il s'agit bien d'une approximation des solutions propres du problème avec conditions de Neumann et surtout de quantifier l'ordre d'erreur induit par les approximations successives afin, le cas échéant, de re-qualifier l'heuristique proposée sur le paramètre affinité gaussien.

Par ailleurs, dans le cadre du traitement d'images, des boîtes affinités 3D pour les niveaux de gris ou 5D pour les couleurs ont été définies pour allier les informations géométrique et de couleur d'une image. Ce concept utilise une normalisation des données afin d'équilibrer le poids des informations de natures différentes. Une étude sur l'influence de cette pondération sur les résultats théoriques pourrait être menée.

Enfin cette étude théorique ouvre la voie à l'étude d'autres méthodes à noyaux, à des méthodes basées sur les noyaux de Mercer notamment, le k -means à noyaux et les estimateurs à noyaux.

Du point de vue de la mise en oeuvre et des performances des stratégies de parallélisation proposées, plusieurs points pourraient être étudiés.

En premier lieu, grâce à un bon choix de paramètre, la méthode du spectral clustering présente les meilleurs résultats sur nos exemples. Cependant d'autres étapes de l'algorithme peuvent être améliorées. En effet, le travail de cette thèse se concentrait sur l'étude de la matrice d'affinité gaussienne, de son paramètre et de ses éléments spectraux afin d'expliquer la séparation des données dans l'espace de projection spectrale et conditionner cette séparation. L'étape de k -means qui constitue l'étape de clustering dans l'espace de projection spectrale n'a pas été à proprement étudiée. En effet, le choix de la métrique, l'initialisation des centres et des estimateurs de la qualité de la partition peuvent être adaptés aux données de l'hypersphère unité de dimension k . Des gains importants, notamment en temps de calculs, pourraient être attendus d'une amélioration de cette étape.

Dans un second temps, d'un point de vue numérique, de nombreuses améliorations peuvent être apportées à la stratégie parallèle du spectral clustering. Afin d'optimiser les performances de la parallélisation, une étude sur la découpe des données visant à équilibrer les données par processus serait nécessaire. De plus, des techniques pour creuser la matrice affinité gaussienne via des tech-

niques de seuillages pourraient être envisagées toujours dans le but de réduire les coûts. De même, le coût numérique lié au calcul des vecteurs propres et des valeurs propres peut être réduit à l'aide de méthodes de Lanczos et d'Arnoldi par exemple. L'aspect "réduit" ou même relativement "cubique" des sous-domaines de calcul n'est pas pris en compte et pourrait orienter plus précisément le choix des algorithmes à employer. Par ailleurs, un lien avec la recherche d'une interprétation du spectral clustering comme approximation d'un problème spectral de chaleur avec conditions de Neumann, l'introduction a priori des informations sur les vecteurs propres recherchés permettraient de simplifier les calculs (limitation de l'espace de recherche des vecteurs propres ou utilisation de méthodes de projection adaptées,...).

Enfin, en ce qui concerne les applications à la biologie et à l'imagerie médicale introduites, deux pistes d'améliorations se détachent.

Tout d'abord, la confrontation avec le monde génomique a montré qu'un pré-traitement des données via une normalisation permet d'améliorer considérablement les résultats. Le mode expérimental propre aux biopuces nous pousse à considérer les erreurs issues de l'hybridation des sondes. Ainsi envisager une méthode de spectral clustering avec des données par intervalles pourrait définir des clusters avec une marge d'erreur.

Finalement, dans le cadre de l'imagerie TEP, il serait intéressant d'étudier dans quelle mesure l'effet des volumes partiels biaise, dans la théorie, les résultats du spectral clustering. En outre, la stratégie parallèle reste à appliquer dans le cas de voxels. Enfin, dans ces deux domaines, nous avons traité des données temporelles. Il serait intéressant d'incorporer cette information d'évolution dans le fonctionnement de la méthode et dans la définition du paramètre de l'affinité gaussienne. De même, la stratégie parallèle reposant sur une décomposition en sous-domaines, un problème serait de savoir comment inclure l'information en amont sur l'ensemble des données d'un sous-domaine.

Bibliographie

- [1] P.D. Acton, L.S. Pilowsky, H.F. Kung, and P.J. Ell. Automatic segmentation of dynamic neuroreceptor single-photon emission tomography images using fuzzy clustering. *European Journal of Nuclear Medicine and Molecular Imaging*, 26(6) :581–590, 1999.
- [2] S. Agostinelli, J. Allison, K. Amako, J. Apostolakis, H. Araujo, P. Arce, M. Asai, D. Axen, S. Banerjee, G. Barrand, et al. Geant4-a simulation toolkit. *Nuclear Instruments and Methods in Physics Research-Section A Only*, 506(3) :250–303, 2003.
- [3] H. Akaike. A new look at the statistical model identification. *IEEE transactions on automatic control*, 19(6) :716–723, 1974.
- [4] M.B. Al-Daoud and S.A. Roberts. New methods for the initialisation of clusters. *Pattern Recognition Letters*, 17(5) :451–455, 1996.
- [5] J. Allison, K. Amako, J. Apostolakis, H. Araujo, P.A. Dubois, M. Asai, G. Barrand, R. Capra, S. Chauvie, R. Chytrcek, et al. Geant4 developments and applications. *IEEE Transactions on Nuclear Science*, 53(1) :270–278, 2006.
- [6] E. Anderson, Z. Bai, C. Bischof, S. Blackford, J. Demmel, J. Dongarra, J. Du Croz, A. Greenbaum, S. Hammarling, A. McKenney, et al. *LAPACK Users' guide*. Society for Industrial Mathematics, 1999.
- [7] N. Aronszajn. *Theory of reproducing kernels*. Harvard University Cambridge, 1950.
- [8] J. Ashburner, J. Haslam, C. Taylor, VJ Cunningham, and T. Jones. A cluster analysis approach for the characterization of dynamic PET data. Chapter 59, 1996.
- [9] G.H. Ball and D.J. Hall. *ISODATA, a novel method of data analysis and pattern classification*. Storming Media, 1965.
- [10] C. Bardos and O. Laffitte. Une synthese de resultats anciens et recents sur le comportement asymptotique des valeurs propres du Laplacien sur une variete riemannienne. 1998.
- [11] M. Belkin and P. Niyogi. Laplacian eigenmaps and spectral techniques for embedding and clustering. *Advances in Neural Information Processing Systems*, 14(3), 2002.
- [12] M. Belkin and P. Niyogi. Laplacian eigenmaps for dimensionality reduction and data representation. *Neural computation*, 15(6) :1373–1396, 2003.
- [13] M. Belkin and P. Niyogi. Towards a theoretical foundation for Laplacian-based manifold methods. *Journal of Computer and System Sciences*, 74(8) :1289–1308, 2008.
- [14] S. Ben-David, U. Von Luxburg, and D. Pál. A sober look at clustering stability. *Learning Theory*, pages 5–19, 2006.
- [15] A. Ben-Hur, A. Elisseeff, and I. Guyon. A stability based method for discovering structure in clustered data. In *Pacific Symposium on Biocomputing 2002 : Kauai, Hawaii, 3-7 January 2002*, page 6. World Scientific Pub Co Inc, 2001.

- [16] A. Ben-Hur, D. Horn, H.T. Siegelmann, and V. Vapnik. Support vector clustering. *The Journal of Machine Learning Research*, 2 :137, 2002.
- [17] P. Berkhin. A survey of clustering data mining techniques. *Grouping Multidimensional Data*, pages 25–71, 2006.
- [18] R. Bhatia and C. Davis. A bound for the spectral variation of a unitary operator. *Linear and Multilinear Algebra*, 15(1) :71–76, 1984.
- [19] P.S. Bradley and U.M. Fayyad. Refining initial points for k-means clustering. In *Proceedings of the Fifteenth International Conference on Machine Learning*, pages 91–99. Citeseer, 1998.
- [20] M. Brand and K. Huang. A unifying theorem for spectral embedding and clustering. *9th International Conference on Artificial Intelligence and Statistics*, 2002.
- [21] JG Brankov, NP Galatsanos, Y. Yang, and MN Wernick. Segmentation of dynamic PET or fMRI images based on a similarity metric. *IEEE Transactions on Nuclear Science*, 50(5 Part 2) :1410–1414, 2003.
- [22] H. Brezis. *Analyse fonctionnelle et applications*. Masson, Paris, 1983.
- [23] F. Camastra and A. Verri. A novel kernel method for clustering. *Biological and Artificial Intelligence Environments*, pages 245–250.
- [24] J. Chen, S. Gunn, M. Nixon, and R. Gunn. Markov random field models for segmentation of PET images. In *information processing in medical imaging*, pages 468–474. Springer, 2001.
- [25] F.R.K. Chung. Spectral Graph Theory (CBMS Regional Conference Series in Mathematics, No. 92). *American Mathematical Society*, 3 :8, 1997.
- [26] P.G. Ciarlet. *The finite element method for elliptic problems*. North-Holland, 1978.
- [27] N. Cristianini, J. Shawe-Taylor, and J. Kandola. Spectral kernel methods for clustering. In *Proceedings of Neural Information Processing Systems*, volume 14, 2001.
- [28] R. Dautray and J.L. Lions. *Analyse mathématique et calcul numérique pour les sciences et les techniques*. Tomes 2 : L’opérateur de Laplace. 1987.
- [29] I. Dhillon, Y. Guan, and B. Kulis. A fast kernel-based multilevel algorithm for graph clustering. In *Proceedings of the eleventh ACM SIGKDD international conference on Knowledge discovery in data mining*, page 634. ACM, 2005.
- [30] I. Dhillon, Y. Guan, and B. Kulis. A unified view of kernel k-means, spectral clustering and graph cuts. *Technical Report TR-04*, 25, 2005.
- [31] I.S. Dhillon, Y. Guan, and B. Kulis. Kernel k-means : spectral clustering and normalized cuts. In *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 551–556. ACM, 2004.
- [32] C. Ding. A tutorial on spectral clustering. In *Talk presented at ICML*. (Slides available at <http://crd.lbl.gov/~cding/Spectral/>), 2004.
- [33] C. Ding, X. He, H. Zha, M. Gu, and H.D. Simon. A min-max cut algorithm for graph partitioning and data clustering. In *Proceedings of the 2001 IEEE International Conference on Data Mining*, pages 107–114, 2001.
- [34] R. Dubes and A.K. Jain. Clustering methodologies in exploratory data analysis. *Advances in Computers*, 19 :113–228, 1980.
- [35] R.O. Duda and P.E. Hart. *Pattern classification and scene analysis*. Wiley, New York, 1973.
- [36] S. Dudoit and J. Fridlyand. A prediction-based resampling method for estimating the number of clusters in a dataset. *Genome Biology*, 3(7), 2002.

- [37] M. Ester, H.P. Kriegel, J. Sander, and X. Xu. A density-based algorithm for discovering clusters in large spatial databases with noise. In *Proc. KDD*, volume 96, pages 226–231, 1996.
- [38] Wen-Yen Chen et Song Yangqiu et Hongjie Bai et Chih-Jen Lin et Edward Y. Chang. Parallel Spectral Clustering in Distributed Systems. *Preprint of IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2010.
- [39] R.J.P.C. Farinha, U. Ruotsalainen, J. Hirvonen, L. Tuominen, J. Hietala, J.M. Fonseca, and J. Tohka. Segmentation of striatal brain structures from high resolution pet images. *Journal of Biomedical Imaging*, 2009 :6, 2009.
- [40] D. Feng, K.P. Wong, C.M. Wu, and W.C. Siu. A technique for extracting physiological parameters and the required input function simultaneously from PET image measurements : Theory and simulation study. *IEEE Transactions on Information Technology in Biomedicine*, 1(4) :243–254, 1997.
- [41] M. Filippone, F. Camastra, F. Masulli, and S. Rovetta. A survey of kernel and spectral methods for clustering. *Pattern Recognition*, 41(1) :176–190, 2008.
- [42] I. Fischer and J. Poland. Amplifying the block matrix structure for spectral clustering. In *Proceedings of the 14th Annual Machine Learning Conference of Belgium and the Netherlands*, pages 21–28. Citeseer, 2005.
- [43] E. Forgy. Cluster analysis of multivariate data : Efficiency vs. interpretability of classifications. *Biometrics*, 21(3) :768, 1965.
- [44] C. Fowlkes, S. Belongie, F. Chung, and J. Malik. Spectral grouping using the Nystrom method. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 26(2) :214–225, 2004.
- [45] C. Fraley and A.E. Raftery. How many clusters? Which clustering method? Answers via model-based cluster analysis. *The Computer Journal*, 41(8) :578, 1998.
- [46] F. Frouin, P. Merlet, Y. Bouchareb, V. Frouin, J.L. Dubois-Rande, A. De Cesare, A. Herment, A. Syrota, and A. Todd-Pokropek. Validation of Myocardial Perfusion Reserve Measurements Using Regularized Factor Images of H215O Dynamic PET Scans. *Journal of Nuclear Medicine*, 42(12) :1737, 2001.
- [47] P.L. George and H. Borouchaki. Triangulation de Delaunay et maillage. Applications aux éléments finis. *Hermes, Paris*, 1997.
- [48] M. Girolami. Mercer kernel-based clustering in feature space. *IEEE Transactions on Neural Networks*, 13(3) :780–784, 2002.
- [49] G.H. Golub and C.F. Van Loan. *Matrix computation*. The John Hopkins University Press Baltimore, 1996.
- [50] H. Guo, R. Renaut, K. Chen, and E. Reiman. Clustering huge data sets for parametric PET imaging. *Biosystems*, 71(1-2) :81–92, 2003.
- [51] L. Hagen and B. Kahng. New spectral methods for ratio cut partitioning and clustering. *IEEE Transactions on computer-aided design*, 11(9), 1992.
- [52] J. Han and M. Kamber. *Data mining : concepts and techniques*. Morgan Kaufmann, 2006.
- [53] J. Han, M. Kamber, and A.K.H. Tung. Spatial clustering methods in data mining : A survey. *Geographic Data Mining and Knowledge Discovery*. Taylor and Francis, 21, 2001.
- [54] P. Hansen, E. Ngai, B.K. Cheung, and N. Mladenovic. Analysis of global k-means, an incremental heuristic for minimum sum-of-squares clustering. *Journal of classification*, 22(2) :287–310, 2005.

- [55] D. Harel and Y. Koren. Clustering spatial data using random walks. In *Proceedings of the seventh ACM SIGKDD international conference on Knowledge discovery and data mining*, page 286, 2001.
- [56] J.A. Hartigan and MA Wong. A k-means clustering algorithm, Algorithm AS 136. *Applied Statistics*, 28(1), 1979.
- [57] J. He, M. Lan, C.L. Tan, S.Y. Sung, and H.B. Low. Initialization of cluster refinement algorithms : a review and comparative study. In *2004 IEEE International Joint Conference on Neural Networks, 2004. Proceedings*, volume 1, 2004.
- [58] R. Herbrich. *Learning kernel classifiers : theory and algorithms*. The MIT Press, 2002.
- [59] A. Hinneburg and D.A. Keim. An efficient approach to clustering in large multimedia databases with noise. *Knowledge Discovery and Data Mining*, 5865, 1998.
- [60] F. Hirsch and G. Lacombe. *Elements of functional analysis*. Springer, 1999.
- [61] AK Jain and RC Dubes. *Algorithms for cluster analysis*. Prentice Hall, Englewood Cliffs, NJ, 1988.
- [62] M. Kac. Can one hear the shape of a drum? *American Mathematical Monthly*, pages 1–23, 1966.
- [63] R. Kannan and A. Vetta. On clusterings : Good, bad and spectral. *Journal of the ACM*, 51(3) :497–515, 2004.
- [64] V. Kapoor, B.M. McCook, and F.S. Torok. An Introduction to PET-CT Imaging1. *Radio-graphics*, 24(2) :523, 2004.
- [65] L. Kaufman and PJ Rousseeuw. Finding groups in data : an introduction to cluster analysis. *NY John Wiley & Sons*.
- [66] S.S. Khan and A. Ahmad. Cluster center initialization algorithm for K-means clustering. *Pattern Recognition Letters*, 25(11) :1293–1302, 2004.
- [67] Y. Kimura, M. Senda, and NM Alpert. Fast formation of statistically reliable FDG parametric images based on clustering and principal components. *Physics in Medicine and Biology*, 47 :455, 2002.
- [68] T. Kohonen. The self-organizing map. *Neurocomputing*, 21.
- [69] T. Kohonen, J. Hynninen, J. Kangas, and J. Laaksonen. SOM PAK : The self-organizing map program package. *Report A31, Helsinki University of Technology, Laboratory of Computer and Information Science*, 1996.
- [70] E. Krestyannikov, J. Tohka, and U. Ruotsalainen. Segmentation of dynamic emission tomography data in projection space. *Computer Vision Approaches to Medical Image Analysis*, pages 108–119, 2006.
- [71] JM Lambert and WT Williams. Multivariate methods in plant ecology : VI. Comparison of information-analysis and association-analysis. *The Journal of Ecology*, 54(3) :635–664, 1966.
- [72] GN Lance and WT Williams. A general theory of classificatory sorting strategies : 1. Hierarchical systems. *The computer journal*, 9(4) :373, 1967.
- [73] GN Lance and WT Williams. A general theory of classificatory sorting strategies : II. Clustering systems. *The Computer Journal*, 10(3) :271, 1967.
- [74] KW Lau, H. Yin, and S. Hubbard. Kernel self-organising maps for classification. *Neurocomputing*, 69(16-18) :2033–2040, 2006.

- [75] R. Maroy, R. Boisgard, C. Comtat, V. Frouin, P. Cathier, E. Duchesnay, F. Dollé, P.E. Nielsen, R. Trébossen, and B. Tavitian. Segmentation of rodent whole-body dynamic PET images : an unsupervised method based on voxel dynamics. *IEEE Transactions on Medical Imaging*, 27(3) :342–354, 2008.
- [76] JC Mazziotta, CC Pelizzari, GT Chen, FL Bookstein, and D. Valentino. Region of interest issues : The relationship between structure and function in the brain. *Journal of cerebral blood flow and metabolism*, 11(2) :A51–A56, 1991.
- [77] JB Mcqueen. Some methods for classification and analysis of multivariate observations. u : Proceedings of Berkeley Symposium on mathematical statistics and probability. *Berkeley, CA, itd : University of California Press*, 1 :281–297, 1967.
- [78] M. Meilă. Comparing clusterings. In *Proceedings of the Conference on Computational Learning Theory*, 2003.
- [79] M. Meila and J. Shi. A random walks view of spectral segmentation. *Artificial Intelligence and Statistics*, 2001, 2001.
- [80] B. Nadler and M. Galun. Fundamental limitations of spectral clustering. *Advances in Neural Information Processing Systems*, 19 :1017, 2007.
- [81] B. Nadler, S. Lafon, R. Coifman, and I. Kevrekidis. Diffusion Maps-a Probabilistic Interpretation for Spectral Embedding and Clustering Algorithms. *Principal Manifolds for Data Visualization and Dimension Reduction*, pages 238–260.
- [82] B. Nadler, S. Lafon, R.R. Coifman, and I.G. Kevrekidis. Diffusion maps, spectral clustering and eigenfunctions of fokker-planck operators. *Preprint*, 2005.
- [83] B. Nadler, S. Lafon, R.R. Coifman, and I.G. Kevrekidis. Diffusion maps, spectral clustering and reaction coordinates of dynamical systems. *Applied and Computational Harmonic Analysis*, 21(1) :113–127, 2006.
- [84] A. Y. Ng, Michael I. Jordan, and Yair Weiss. On spectral clustering : analysis and an algorithm. *Proceedings of Neural Information Processing Systems*, 2002.
- [85] Freeman W.T. Perona, P. A factorization approach to grouping. *European Conference on Computer Vision*, 1998.
- [86] H. Queffélec. *Topologie*. Masson, 1998.
- [87] Thomas Raviart. *Introduction à l'analyse numérique des équations aux dérivées partielles*. Dunod, 1993.
- [88] E. Schikuta and M. Erhart. The BANG-clustering system : grid-based data analysis. *Advances in Intelligent Data Analysis Reasoning about Data*, pages 513–524, 1997.
- [89] G. Schwarz. Estimating the dimension of a model. *The annals of statistics*, 6(2) :461–464, 1978.
- [90] SZ Selim and M. Ismail. K-means-type algorithms : a generalized convergence theorem and characterization of local optimality. *IEEE Transaction of Pattern Analysis and Machine intelligence*, 6(1) :81–86, 1984.
- [91] J. Shi and J. Malik. Normalized cuts and image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22(8) :888–905, 2000.
- [92] Y. Song, W.Y. Chen, H. Bai, C.J. Lin, and E.Y. Chang. Parallel spectral clustering. In *Proceedings of European Conference on Machine Learning and Pattern Knowledge Discovery*. Springer, 2008.

- [93] M. Soret, S.L. Bacharach, and I. Buvat. Partial-volume effect in PET tumor imaging. *Journal of Nuclear Medicine*, 48(6) :932, 2007.
- [94] P. Soularue and X. Gidrol. Puces à ADN. *Techniques de l'Ingénieur*, 6 :1–10, 2002.
- [95] D. Steinley. K-means clustering : A half-century synthesis. *British Journal of Mathematical and Statistical Psychology*, 59(1) :1–34, 2006.
- [96] G.W. Stewart and J. Sun. *Matrix perturbation theory*. Academic press New York, 1990.
- [97] S. Still and W. Bialek. How many clusters? an information-theoretic perspective. *Neural Computation*, 16(12) :2483–2506, 2004.
- [98] P. Tamayo, D. Slonim, J. Mesirov, Q. Zhu, S. Kitareewan, E. Dmitrovsky, E.S. Lander, and T.R. Golub. Interpreting patterns of gene expression with self-organizing maps : methods and application to hematopoietic differentiation. *Proceedings of the National Academy of Sciences of the United States of America*, 96(6) :2907, 1999.
- [99] P. Törönen, M. Kolehmainen, G. Wong, and E. Castrén. Analysis of gene expression data using self-organizing maps. *FEBS letters*, 451(2) :142–146, 1999.
- [100] J.W. Tukey. Exploratory data analysis. *Massachusetts : Addison-Wesley*.
- [101] T.G. Turkington. Introduction to PET instrumentation. *Journal of Nuclear Medicine Technology*, 29(1) :4, 2001.
- [102] F. Vailleau, E. Sartorel, M.F. Jardinaud, F. Chardon, S. Genin, T. Huguet, L. Gentzbittel, and M. Petitprez. Characterization of the interaction between the bacterial wilt pathogen *Ralstonia solanacearum* and the model legume plant *Medicago truncatula*. *Molecular Plant-Microbe Interactions*, 20(2) :159–167, 2007.
- [103] V.N. Vapnik. *The nature of statistical learning theory*. Springer Verlag, 2000.
- [104] D. Verma and M. Meila. A comparison of spectral clustering algorithms. *University of Washington, Technical Report UW-CSE-03-05-01*, 2003.
- [105] U. Von Luxburg. A tutorial on spectral clustering. *Statistics and Computing*, 17(4) :395–416, 2007.
- [106] U. Von Luxburg, M. Belkin, and O. Bousquet. Consistency of spectral clustering. *Annals of Statistics*, 36(2) :555, 2008.
- [107] U. Von Luxburg, O. Bousquet, and M. Belkin. Limits of spectral clustering. *Advances in neural information processing systems*, 17, 2004.
- [108] C. Wagschal. *Dérivation, intégration*. Hermann, 1999.
- [109] D.L. Wallace. Comment. *Journal of the American Statistical Association*, 78(383) :569–576, 1983.
- [110] W. Wang, J. Yang, and R. Muntz. STING : A statistical information grid approach to spatial data mining. In *Proceedings of the International Conference on Very Large Data Bases*, pages 186–195. Citeseer, 1997.
- [111] J.H. Ward Jr. Hierarchical grouping to optimize an objective function. *Journal of the American statistical association*, 58(301) :236–244, 1963.
- [112] K.P. Wong, D. Feng, SR Meikle, and MJ Fulham. Segmentation of dynamic PET images using cluster analysis. In *2000 IEEE Nuclear Science Symposium Conference Record*, volume 3, 2000.
- [113] R. Xu and D. Wunsch. Survey of clustering algorithms. *IEEE Transactions on neural networks*, 16(3) :645–678, 2005.

- [114] L. Yen, D. Vanvyve, F. Wouters, F. Fouss, M. Verleysen, and M. Saerens. Clustering using a random walk based distance measure. In *Proceedings of the 13th Symposium on Artificial Neural Networks (ESANN 2005)*, pages 317–324, 2005.
- [115] H. Zaidi and W.D. Erwin. Quantitative analysis in nuclear medicine imaging. *Journal of Nuclear Medicine*, 48(8) :1401, 2007.
- [116] L. Zelnik-Manor and P. Perona. Self-tuning spectral clustering. *Advances in Neural Information Processing Systems*, 17(1601-1608) :16, 2004.
- [117] I.G. Zubal, C.R. Harrell, E.O. Smith, Z. Rattner, G. Gindi, and P.B. Hoffer. Computerized three-dimensional segmented human anatomy. *Medical Physics-Lancaster*, 21 :299–299, 1994.

Publications

Conférences internationales avec actes et comité de sélection

- [3] S. Mouysset, J. Noailles, and D. Ruiz. On a strategy for Spectral Clustering with parallel computation. In *High Performance Computing for Computational Science : 9th International Conference, San Francisco, CA USA, 22/06/2008-25/06/2010*. Springer-Verlag, 2010.
- [2] S. Mouysset, J. Noailles, and D. Ruiz. On an Interpretation of Spectral Clustering via Heat Equation and Finite Elements Theory. In *International Conference on Data Mining and Knowledge Engineering (ICDMKE), Londres, 30/06/2010-02/07/2010*. IAENG, 2010 (Best Paper Award).
- [1] S. Mouysset, J. Noailles, and D. Ruiz. Using a Global Parameter for Gaussian Affinity Matrix in Spectral Clustering. In *International Meeting High Performance Computing for Computational Science (VECPAR), Toulouse, 24/06/2008-27/06/2008*. Springer-Verlag, juin 2008.

Conférences internationales avec comité de sélection

- [3] S. Mouysset, J. Noailles, and D. Ruiz. A parallel strategy for spectral clustering. In *2nd Japanese French Laboratory for Informatics Workshop (JFLI), Paris, October 12-13, 2010*.
- [2] S. Mouysset, J. Noailles, and D. Ruiz. On the efficiency of Spectral Clustering : interpretation, parallel computation and results. In *European Conference on Operational Research (EURO'10), Lisbon, Portugal, June 30- July 02, 2010*.
- [1] S. Mouysset, J. Noailles, and D. Ruiz. On the efficiency of Spectral Clustering : interpretation and results. In *Sparse Days Meeting, Toulouse, June 18-19, 2009*.

Rapports de Recherche

- [4] S. Mouysset and A. Senescau. Analyse de données génomiques : hybridation de sondes. Rapport de recherche RT-APO-10-07, IRIT, Université Paul Sabatier, Toulouse, février 2010.
- [3] S. Mouysset, J. Noailles, and D. Ruiz. On a new interpretation of Spectral Clustering via Heat equation and Finite Elements theory. Rapport de recherche RT-APO-09-02, IRIT, Université Paul Sabatier, Toulouse, juin 2009.
- [2] S. Mouysset, J. Noailles, and L. Gentzbittel. Classification de profils d'expression de gènes : méthode Self Organizing Maps. Rapport de recherche RT-APO-09-04, IRIT, Université Paul Sabatier, Toulouse, décembre 2009.
- [1] S. Mouysset and J.Noailles. Classification non supervisée : une approche spectrale et une approche via vecteurs support. Technical Report RT-APO-07-7, IRIT, Septembre 2006.

Contributions à l'étude de la classification spectrale et applications

Résumé : La classification spectrale consiste à créer, à partir des éléments spectraux d'une matrice d'affinité gaussienne, un espace de dimension réduite dans lequel les données sont regroupées en classes. Cette méthode non supervisée est principalement basée sur la mesure d'affinité gaussienne, son paramètre et ses éléments spectraux. Cependant, les questions sur la séparabilité des classes dans l'espace de projection spectral et sur le choix du paramètre restent ouvertes. Dans un premier temps, le rôle du paramètre de l'affinité gaussienne sera étudié à travers des mesures de qualités et deux heuristiques pour le choix de ce paramètre seront proposées puis testées. Ensuite, le fonctionnement même de la méthode est étudié à travers les éléments spectraux de la matrice d'affinité gaussienne. En interprétant cette matrice comme la discrétisation du noyau de la chaleur définie sur l'espace entier et en utilisant les éléments finis, les vecteurs propres de la matrice affinité sont la représentation asymptotique de fonctions dont le support est inclus dans une seule composante connexe. Ces résultats permettent de définir des propriétés de classification et des conditions sur le paramètre gaussien. A partir de ces éléments théoriques, deux stratégies de parallélisation par décomposition en sous-domaines sont formulées et testées sur des exemples géométriques et de traitement d'images. Enfin dans le cadre non supervisé, le classification spectrale est appliquée, d'une part, dans le domaine de la génomique pour déterminer différents profils d'expression de gènes d'une légumineuse et, d'autre part dans le domaine de l'imagerie fonctionnelle TEP, pour segmenter des régions du cerveau présentant les mêmes courbes d'activités temporelles.

Mots-clés : classification non supervisée, classification spectrale, noyau gaussien, équation de la chaleur, éléments finis, parallélisation, imagerie médicale.

Contributions to the study of spectral clustering and applications

Abstract : The Spectral Clustering consists in creating, from the spectral elements of a Gaussian affinity matrix, a low-dimension space in which data are grouped into clusters. This unsupervised method is mainly based on Gaussian affinity measure, its parameter and its spectral elements. However, questions about the separability of clusters in the projection space and the spectral parameter choices remain open. First, the rule of the parameter of Gaussian affinity will be investigated through quality measures and two heuristics for choosing this setting will be proposed and tested. Then, the method is studied through the spectral element of the Gaussian affinity matrix. By interpreting this matrix as the discretization of the heat kernel defined on the whole space and using finite elements, the eigenvectors of the affinity matrix are asymptotic representation of functions whose support is included in one connected component. These results help define the properties of clustering and conditions on the Gaussian parameter. From these theoretical elements, two parallelization strategies by decomposition into sub-domains are formulated and tested on geometrical examples and images. Finally, as unsupervised applications, the spectral clustering is applied, first in the field of genomics to identify different gene expression profiles of a legume and the other in the imaging field functional PET, to segment the brain regions with similar time-activity curves.

Key-words : clustering, spectral clustering, Gaussian kernel, heat equation, finite elements, parallelization, medical imaging.