

UNIVERSITÉ TOULOUSE I CAPITOLE
ECOLE DOCTORALE MATHÉMATIQUES, INFORMATIQUE ET
TÉLÉCOMMUNICATIONS DE TOULOUSE

**MÉTHODES VARIATIONNELLES POUR L'ESTIMATION,
L'INFÉRENCE ET LA DÉCISION
DANS LES MODÈLES GRAPHIQUES**

MÉMOIRE D'HABILITATION PRÉSENTÉ PAR

Nathalie Peyrard

pour l'obtention du diplôme
D'HABILITATION À DIRIGER DES RECHERCHES

présenté et soutenu publiquement le 4 octobre 2013 devant

Le Jury

Président : Jean-Michel Marin, Professeur de l'Université Montpellier II
Rapporteurs : Philippe Leray, Professeur Polytech Nantes
Stéphane Robin, Directeur de Recherche INRA
Christine Thomas-Agnan, Professeur Université Toulouse 1 Capitole
Examineur : Joël Chadœuf, Directeur de Recherche INRA

Remerciements

Je suis reconnaissante à Christine Thomas-Agnan d'avoir accepté d'être tutrice de ce travail de préparation d'un mémoire d'HDR. Elle a ensuite bien voulu en être rapporteur tout comme Philippe Leray et Stéphane Robin. Je les remercie, ainsi que les autres membres du Jury, Joël Chadœuf et Jean-Michel Marin, pour avoir accepté d'apporter leurs points de vue sur mes activités de recherche.

La recherche est un travail collectif et ce mémoire n'existerait pas sans le soutien de nombreuses personnes. Je tiens à remercier tout particulièrement certaines d'entre elles qui ont accompagnées différentes étapes de mon parcours.

Mes premiers pas dans la recherche ont été au côté de Joël Chadœuf, lors de mon stage de maîtrise, à l'INRA déjà. Je crois bien qu'il a su déclencher la fibre puisque j'ai poursuivi dans ce domaine jusqu'à aujourd'hui encore. Je le remercie pour m'avoir accompagnée dans mes premières collaborations avec les biologistes. Je lui dois également une certaine introduction aux bonnes pratiques de notre métier de chercheur.

Au delà de Joël, c'est toute l'unité de statistiques spatiales de l'INRA d'Avignon que je tiens à remercier pour l'accueil chaleureux et les échanges scientifiques riches dont j'ai pu bénéficier dans mes premières années de Chargée de Recherche à l'INRA. Il a bien fallu que je m'émancipe au bout d'un moment mais j'y retourne toujours avec plaisir et j'espère que les échanges continueront !

Entre le stage et le recrutement il y a eu la thèse. Gilles Celeux et Florence Forbes m'ont fait découvrir les champs de Markov et les méthodes variationnelles. J'y ai pris goût ! Je les remercie pour leurs conseils et leur suivi au cours de ces trois années de thèse.

J'ai eu la chance de croiser dans une école chercheur Alain Franc, autre grand passionné des méthodes variationnelles et du champ moyen, et de me retrouver dans le même train qui nous ramenait chez nous à la fin de la semaine. Une collaboration sur la durée s'est construite suite à ces premières discussions. J'apprécie ses approches "à la physicienne" et sa ténacité dans nos échanges. En général cela converge vers de jolies mathématiques. Le collègue est devenu ami et c'est toujours avec plaisir que les Dubois-Peyrard l'accueillent pour un petit verre de vin et une histoire.

L'autre rencontre, toulousaine cette fois, et qui a largement influencé mes activités de recherche, est celle avec Régis Sabbadin. Comment faire travailler ensemble un informaticien et une statisticienne ? Je remercie Régis pour sa curiosité scientifique et sa patience qui ont permis que cette collaboration nous amène là où nous sommes aujourd'hui. Même si je reste toujours imperméable à la théorie de la complexité, j'ai enfin compris ses retombées. Un autre collègue devenu ami.

La majorité des travaux résumés ici n'auraient pu être réalisés sans les étudiants avec qui j'ai pu travailler. Leurs contributions ont permis de faire avancer de nombreuses pistes de

recherche. J'ai en particulier pris un grand plaisir à co-encadrer Mathieu Bonneau, premier thésard, toujours de bonne humeur. Le plaisir continue avec Julia Radoszycki.

A Toulouse j'ai également découvert une nouvelle unité, MIAT. Je remercie ses membres pour leur accueil il y a maintenant 7 ans. Ils m'ont permis d'y trouver ma place et de partir au travail avec plaisir tous les matins.

Enfin, je termine par Manu qui a eu la patience de relire ce manuscrit dont le sujet est bien loin de l'IHM. Ses retours m'ont permis d'améliorer ce manuscrit ainsi que ma présentation orale. Merci et rendez-vous pour l'étape suivante!

Résumé

De nombreux systèmes réels sont caractérisés par un ensemble d'entités en interaction. Ces interactions peuvent être dues à une proximité spatiale, ou sociale. Elles peuvent aussi correspondre à des relations de cause à effet. Dans tous les cas, elles structurent le système. Dans la famille des modèles stochastiques, les modèles graphiques offrent un cadre de représentation riche pour raisonner sur ces systèmes structurés, que le système soit statique ou dynamique, totalement ou partiellement observé. Ils permettent de modéliser les interactions entre les entités du système et l'incertitude sur ces interactions, pour ensuite formaliser des problèmes de compréhension, diagnostic ou contrôle.

La spécificité d'un modèle graphique tient à la représentation factorisée de la loi de probabilité jointe en un produit de fonctions locales. La représentation en est ainsi simplifiée. Par contre, les tâches d'estimation, inférence et décision deviennent complexes et ne peuvent être résolues de manière exacte pour la plupart des systèmes réels. Une approche classique pour la résolution approchée consiste à travailler à partir de simulations. L'avantage de cette approche est l'existence de propriétés asymptotiques. Le défaut est le temps de calcul des algorithmes de résolution. Cela devient limitant lorsque les opérations doivent être répétées un grand nombre de fois.

Depuis quelques années, une approche alternative est de plus en plus souvent exploitée : l'approche variationnelle, qui trouve ses origines en mécanique statistique. Même si les résultats théoriques sur la qualité des estimateurs sont quasi inexistants, les domaines du machine learning, de la statistique et de l'écologie théorique se sont emparés de ces outils efficaces en pratique, pour résoudre les tâches d'estimation et inférence. En théorie de la décision, par contre, l'utilisation de l'approche variationnelle reste anecdotique.

Les trois tâches d'estimation, inférence et décision ne sont pas indépendantes. Aussi est-il intéressant de développer les méthodes variationnelles pour les trois afin de disposer d'une approche globale cohérente. Les travaux qui sont décrits ici, à l'interface entre statistique et algorithmique, vont dans cette direction et contribuent à démontrer l'intérêt des méthodes variationnelles. Pour l'estimation approchée dans les modèles graphiques à données cachées je propose des applications et des généralisations des méthodes variationnelles pour la mise en œuvre de l'algorithme EM. En inférence, je présente des extensions des méthodes variationnelles afin de ne plus négliger les interactions à longue distance, et je l'applique au calcul des états transients et d'équilibre du modèle de processus de contact. Enfin, en décision, je présente une des premières propositions pour l'exploitation des méthodes variationnelles à la résolution de processus décisionnels de Markov lorsque l'espace d'états et l'espace d'actions sont multidimensionnels. Ces travaux sont complétés par des contributions sur la construction de la structure d'un modèle graphique à partir d'un jeu de données.

Les modèles graphiques sont bien adaptés pour représenter certains systèmes structurés étudiés en écologie ou en épidémiologie, comme les réseaux trophiques, les réseaux de propagation de maladies ou encore les réseaux d'habitats. La plupart des contributions méthodologiques présentées dans ce mémoire ont été mises en œuvre sur des applications dans ces domaines, pour répondre à des questions de compréhension ou de gestion des systèmes biologiques concernés.

Abstract

Many real systems are characterised by a set of elements in interaction. These interactions can be spatial, or social. They can also be the consequence of causality. In all cases, they structure the system. In the family of stochastic models, graphical models provide a rich framework for reasoning on such structured systems. The system under study can be static or dynamic, fully or partially observable. In all these situations, graphical models allow to model interactions between the system's elements and the uncertainty attached to these interactions, and to formalize problems of analysis, diagnosis or control.

Graphical models are characterized by a factored representation of the joint probability distribution as a product of local functions. This representation is therefore simple. However, the three main task of estimation, inference and decision become more complex and cannot be solved exactly for most real systems. A classical approach for approximate resolution is based on simulations of the system. The advantage is the existence of asymptotic properties. The main drawback is the computational time of associated resolution algorithms. This can be a severe limit when a task has to be performed several times.

The last decades, an alternative approach has gained interest among scientists : the variational approach, from statistical physics. Eventhough there are no theoretical results on the quality of the variational methods, they are more and more exploited for estimation and inference in the fields of machine learning, statistics and theoretical ecology because of their good behavior in practice. However, they are almost absent in decision theory.

The three tasks of estimation, inference and decision are not independent. Therefore it is worth developing the use of variational approaches for the three of them, in order to define a coherent and global approach. The research work presented here, at the interface of statistics and algorithmic, participate to this objective and contributes to show the interest of variational methods. For approximate estimation in graphical models with hidden data, I propose applications and generalisations of variational methods to make the EM algorithm tractable. For inference, I present extensions of variational methods which takes into account long range correlations. I apply them to the computation of transient and equilibrium states of the contact process. Finally, for decision I describe one of the first propositions to exploit variational methods for the resolution of Markov decision processes when states and actions spaces are multidimensional.

Graphical models are good candidates to represent structured systems that are studied in ecology or epidemiology, like trophic networks, disease spread network or habitat networks. Most of the methodological contributions presented here have been applied on problems in one of these fields, either for analysis or management of biological systems.

Table des matières

1	Introduction	12
1.1	Motivation	12
1.2	La famille des modèles graphiques probabilistes	13
1.3	Calcul exact et approché dans les MGP	15
1.4	Les méthodes variationnelles pour l'estimation, l'inférence et la décision approchées dans les MGP	16
1.5	Applications en épidémiologie et en écologie	17
2	Plusieurs visions des méthodes variationnelles	18
2.1	Mécanique statistique et énergie libre	18
2.2	Mécanique statistique et approximation de Kikuchi	20
2.3	Machine learning et algorithmes de passage de messages	21
2.4	Statistique et borne inférieure de la vraisemblance	22
2.5	Dynamique des populations et fermeture des moments	22
2.6	Conclusion	23
3	Contributions	25
3.1	Construction du modèle	25
3.1.1	Analyse statistique de données sur grille par tests de permutation	25
3.1.2	Caractérisation du maximum d'entropie sous contrainte de marginales	29
3.2	Estimation dans les modèles graphiques	31
3.2.1	EM, EM variationnel et EM bayésien variationnel	32
3.2.2	Algorithmes EM de type champ moyen pour les champs de Markov cachés	34
3.2.3	Algorithme VBEM pour le modèle de Cox log-Gaussien	38
3.2.4	Applications des MGP à variables cachées en épidémiologie et en écologie	39
3.3	Inférence dans les modèles graphiques	44
3.3.1	Méthodes variationnelles d'ordre 2 pour l'inférence des états d'équilibre et transients d'un modèle SIS sur graphe	45
3.3.2	Rôle de la structure des interactions dans la propagation d'une épidémie	47
3.3.3	Méthode variationnelle d'ordre 1 pour estimer la taille de l'épidémie dans un modèle SIR sur graphe	48
3.4	Décision dans les modèles graphiques	50

3.4.1	Un cadre et un algorithme de résolution approchée pour les Processus Décisionnels de Markov en contexte structuré	52
3.4.2	Modélisation de problèmes de gestion spatialisés en épidémiologie et en écologie	54
3.4.3	Résolution approchée du problème d'échantillonnage adaptatif dans les champs de Markov	58
3.4.4	Application de l'échantillonnage adaptatif à la cartographie en épidémiologie et en écologie	62
3.5	Conclusions	66
4	Projet de Recherche	69
4.1	Enjeux	69
4.2	Objectifs méthodologiques en décision	70
4.2.1	Contrôle d'un processus parfaitement observé : algorithmes de résolution approchée pour les PDM à espaces d'état et d'action factorisés.	70
4.2.2	Echantillonnage d'un processus partiellement observé : méthodes pour l'échantillonnage adaptatif lorsque le phénomène évolue dans le temps.	70
4.2.3	Compromis entre gestion et acquisition d'information dans un processus partiellement observé.	72
4.2.4	Quelques pistes méthodologiques	72
4.3	Objectifs en modélisation et applications	73
5	Curriculum Vitæ détaillé	83
5.1	Etat civil	83
5.2	Parcours professionnel	83
5.3	Projets de recherche	83
5.4	Encadrement	85
5.5	Visites dans laboratoires étrangers	87
5.6	Animation scientifique	88
5.7	Contribution au fonctionnement de collectifs INRA	89
5.8	Activités d'enseignement	90
6	Liste des publications	91
6.1	Articles scientifiques	91
6.2	Chapitres d'ouvrages	92
6.3	Conférences internationales avec sélection sur article long	93
6.4	Conférences nationales ou francophones avec sélection sur article long . . .	94
6.5	Workshops, colloques, congrès (sélection sur résumé ou sans sélection) . . .	95
6.6	Mémoires et rapports de recherche	95

Chapitre 1

Introduction

1.1 Motivation

L'évolution d'une épidémie est le résultat de la combinaison de deux facteurs : les propriétés intrinsèques de la maladie et des individus, et les propriétés du réseau d'interaction entre les individus. En épidémiologie, les questions de compréhension des organisations des populations, de leurs dynamiques, ou de choix de stratégies pour la gestion ou le contrôle de ces populations ne peuvent plus être abordées sans prendre en compte le réseau d'interaction, soit géographique, soit social, qui structure la population étudiée.

Ainsi, en épidémiologie humaine ou animale, les individus se déplacent. Le réseau d'interaction est alors constitué soit de l'ensemble de ces individus, soit d'un ensemble d'entités administratives sur lesquelles les informations sont collectées ([49, 45, 57]. Les liens entre individus ou localités traduisent des contacts privilégiés entre individus (lieu d'habitation, travail, élevage) ou des déplacements (voyage, travail). Ils constituent des chemins pour la propagation d'agents infectieux ou pathogènes.

En épidémiologie végétale, les individus sont statiques : il peut par exemple s'agir d'un verger dont les arbres (les hôtes) sont envahis par un bioagresseur [17]. Dans ce cas, les liens du réseau d'interaction indiquent que deux hôtes sont suffisamment proches pour que l'un soit contaminé par des parasites issus de l'autre. Le réseau aura une forme plus régulière que dans le cas humain ou animal, du fait de la plantation régulière des arbres dans un verger.

L'écologie est un autre domaine où la notion de réseau est centrale dans la description des dynamiques des processus étudiés. Ainsi en écologie des populations, les nœuds du réseau sont des habitats, et un lien entre deux habitats indique qu'il existe un chemin direct pour un déplacement d'individus de l'un vers l'autre [67].

Un dernier exemple, en biologie de la conservation, illustre l'idée que la notion de lien entre deux nœuds d'un réseau d'interaction n'est pas nécessairement symétrique : dans un réseau trophique [114], les nœuds sont les espèces et les liens indiquent des relations proies-prédateurs. Le lien a donc une direction, il pointe depuis la proie vers le prédateur.

Dans tous ces exemples, une information est associée à chaque nœud du réseau, par exemple un état sanitaire, ou une note de sévérité en épidémiologie, une abondance ou simplement une information de présence/absence en écologie. Cette description à l'échelle des individus combinée avec la description du réseau et les connaissances sur les propriétés intrinsèques de la maladie ou la population, donneront une bonne représentation du processus.

Plusieurs objectifs peuvent motiver une telle description du phénomène. Il peut s'agir d'un objectif de compréhension : comment la structure du réseau d'interaction influence-t-elle la propagation d'une épidémie ? Comment des notes de sévérité de maladie sont-elles corrélées dans l'espace et dans le temps ? Il peut s'agir aussi de prédiction : quelle va être l'expansion d'une épidémie à partir de sa zone source ? Enfin, l'objectif peut être un objectif de gestion : quelle stratégie d'isolement et de vaccination adopter lors d'une épidémie ? Comment contrôler l'installation d'une espèce invasive, ou au contraire favoriser celle d'une espèce menacée, sachant que les ressources disponibles sont limitées ?

1.2 La famille des modèles graphiques probabilistes

Pour formaliser les problèmes soulevés par ces questions et les précédentes, les modèles graphiques probabilistes¹ (MGP, [63]) offrent un cadre de représentation riche. Ils permettent de raisonner “facilement” sur un ensemble d'entités élémentaires (individus, espèces, ...) en interaction et de modéliser l'incertitude sur l'état du système complet ou d'un sous système. Dans un MGP, l'état de chaque entité élémentaire est représentée par une variable aléatoire. La probabilité jointe de l'ensemble de ces variables s'exprime de manière factorisée comme un produit de fonctions ne faisant intervenir qu'un sous-ensemble des variables. Plus formellement, soit $X = (X_1, \dots, X_n)$ un ensemble de n variables aléatoires à valeurs respectivement dans² $\Omega_1, \dots, \Omega_n$. On dira que la loi jointe P de X est définie par un modèle graphique s'il existe un ensemble de fonctions positives $\{\Psi_c\}_{c \in \mathcal{C}}$ indicées sur des sous-ensembles de $V = \{1, \dots, n\}$ telles que quel que soit $x \in \Omega = \Omega_1 \times \dots \times \Omega_n$

$$p(X = x) \propto \prod_{c \in \mathcal{C}} \Psi_c(x_c),$$

où x_c désigne l'ensemble $\{x_i, i \in c\}$ et $\mathcal{C}$ est un ensemble de parties de V . Les fonctions Ψ_c jouent un rôle central car leur forme va générer la structure de la corrélation entre les variables. Elles sont souvent désignées comme fonctions potentiel et elles dépendent en général d'un ou plusieurs paramètres. A partir de ces fonctions, il est possible de définir une représentation graphique du MGP, $G = (V, E)$: les nœuds sont les entités et sont indicés sur V . Lorsque deux entités apparaissent dans une même fonction potentiel, elles sont reliées par une arête. Ces arêtes peuvent être orientées ou non, elles forment l'ensemble E . La famille des MGP regroupe une grande variété de modèles. Certains sont des modèles pour raisonner ou prédire, d'autres sont des modèles pour la décision. Certains représentent des phénomènes statiques, d'autres des dynamiques.

- *MGP pour raisonner ou prédire* : dans cette famille, on distingue les réseaux bayésiens [51] qui correspondent au cas où le graphe G est orienté et acyclique, et les champs de Markov [68] lorsque le graphe G est non orienté. Dans un réseau bayésien les fonctions potentiel sont des probabilités conditionnelles. Cette propriété est perdue dans le cas

1. Je précise ici probabilistes car la notion de modèle graphique existe également dans un cadre déterministe, comme par exemple dans le cadre des Problèmes de Satisfaction de Contraintes Pondérées. Le passage entre les deux formalismes est d'ailleurs relativement aisé.

2. Sauf mention contraire, les MGP considérés dans ce mémoire sont à espace d'état discret.

des champs de Markov. Ces modèles statiques peuvent être étendus dans la direction temporelle, conduisant aux réseaux bayésiens dynamiques ([43]) et aux champs de Markov dynamiques.

- *MGP décisionnels* : deux sous-familles de MGP sont disponibles pour modéliser la décision et surtout l’optimiser : les diagrammes d’influence [51] et les processus décisionnels de Markov (PDM) factorisés [11]. Les diagramme d’influence sont des réseaux bayésiens augmentés de nœuds “décision” (déterministes ou non) et de nœuds “utilité”. Ils permettent de représenter des problèmes d’optimisation de la décision sur un système structuré. La notion de temps n’y est pas explicitement représentée. Les PDM factorisés, version décisionnelle des MGP dynamiques, offrent un formalisme pour répondre à la question de la conception de stratégies de gestion d’un système dans la durée. On se trouve alors dans le cadre de l’optimisation séquentielle dans l’incertain, avec des notions de récompense associée à un état et de valeur d’une stratégie. L’action prise au temps t va influencer les probabilité de transition qui gouvernent le processus et l’objectif est de calculer la séquence d’actions optimale pour un critère donné. Les PDM factorisés à horizon fini peuvent être vus comme des cas particuliers des diagrammes d’influence.

Pour un MGP donné, plusieurs tâches peuvent être effectuées : apprentissage de la structure, estimation des paramètres, inférence et décision.

- *Apprentissage* : L’apprentissage de la structure d’un MGP fait appel à un ensemble de techniques statistiques afin d’estimer quels liens sont suffisamment significatifs pour être représentés par une arête dans le modèle. Jusqu’à présent je n’ai pas travaillé sur cette tâche aussi je ne m’étendrai pas plus dessus.
- *Estimation* : Lorsque la structure du MGP est connue (le graphe G associé) ainsi que la forme paramétrique des fonctions potentiel, l’estimation consiste à apprendre ces paramètres à partir de jeux de données complets ou incomplets. Cela relève ici encore des outils de la statistique. Mes travaux sur cette question font l’objet de la section 3.2.
- *Inférence* : Lorsque le modèle est entièrement connu, l’inférence³ regroupe un ensemble de tâches comme le calcul de la probabilité de l’état global le plus probable, ou de la distribution de probabilité d’un sous-ensemble de variables conditionnellement (ou non) à des informations sur d’autres variables du système. Ces questions sont beaucoup étudiées dans les domaines de la statistique et en machine learning. Je présente mes contributions à l’inférence dans les MGP dans la section 3.3.
- *Décision* : Enfin, dans un MGP dans sa version dynamique et décisionnelle, toujours lorsque le modèle est entièrement connu, la tâche de décision pose le problème de trouver la stratégie qui optimise l’espérance de la somme des récompenses au cours du temps. Les méthodes pour traiter ce problème relèvent de l’intelligence artificielle. Je décris mes contributions en décision dans les MGP dans la section 3.4.

Pourquoi existe-t-il une communauté scientifique très active autour des MGP [53, 9, 63, 2] ? Parce que pour la résolution de toutes les tâches présentées ci-dessus, tout calcul exact

3. Cette définition du terme inférence est celle du domaine de l’intelligence artificielle. En statistique, l’inférence inclut l’estimation du modèle.

est hors de portée dès que le modèle dépasse quelques dizaines de nœuds. La complexité est liée non seulement au nombre de nœuds mais aussi à la structure du graphe. Cela pourrait générer une forte limitation à l'utilisation de ces modèles étant donné que dans les exemples présentés plus haut la taille des réseaux d'interaction peut atteindre 100 nœuds pour un verger, de 20 à quelques centaines de nœuds pour un réseau trophique, et plus d'un millier de nœuds en épidémiologie humaine ou animale. La majorité des activités de la communauté des MGP traite de méthodes pour le calcul approché, et de leur application sur de nouveaux problèmes, de plus en plus grands. Un des enjeux est de rajouter la dimension spatiale à des formalismes initialement purement dynamiques et de proposer les outils algorithmiques associés efficaces. Ainsi, dans les cadres classiques pour la décision, l'aspect dynamique est bien représenté, mais l'aspect spatial l'est moins et les techniques de résolution de problème de décision séquentielle ne sont pas adaptées au cas spatial, car trop lentes.

1.3 Calcul exact et approché dans les MGP

Pour chacune des tâches d'estimation, inférence et décision, des familles de méthodes ont été développées pour la résolution approchée du problème exact. Certaines sont spécifiques à la tâche considérée, d'autres, comme les méthodes basées sur la simulation ont été exploitées dans les trois tâches.

L'estimation des paramètres d'un MGP peut être abordée selon deux approches classiques d'estimation : soit fréquentiste avec l'estimateur du maximum de vraisemblance, soit bayésienne avec l'estimateur du même nom. Dans le cas des MGP, la mise en œuvre de ces principes ne peut pas toujours être réalisée de manière exacte. C'est le cas par exemple pour un champ de Markov essentiellement à cause du calcul de la constante de normalisation du modèle. L'existence de données manquantes ou cachées compliquera également les calculs. Trois approches co-existent pour faire de la résolution approchée : les approches basées sur des simulations du MGP [101], celles basées sur la pseudo-vraisemblance [6] et les approches dites variationnelles [112]. Les premières sont gourmandes en temps et en mémoire, mais présentent l'avantage d'offrir des garanties asymptotiques. Les deux suivantes sont plus rapides, mais ne sont que très rarement associées à des résultats théoriques sur leur qualité. L'approximation de la loi jointe du modèle graphique par la pseudo-vraisemblance utilisée seule ne suffit pas toujours pour obtenir un problème d'estimation approchée traitable efficacement. Enfin, les méthodes variationnelles sont de plus en plus utilisées pour leur bon comportement en pratique.

Les méthodes exactes d'inférence reposent pour la plupart sur l'un des deux principes suivant : le principe d'élimination de variables [85] et le principe de séparation/évaluation [87] (uniquement pour résoudre un problème d'optimisation). Ainsi, pour le calcul du mode d'un MGP, dans le premier cas on cherche à sommer sur les variables dans un ordre intelligent de manière à minimiser le nombre de calculs à effectuer et à les mutualiser. Cette approche, qui repose sur les principes de la programmation dynamique non sérielle est développée en intelligence artificielle, pour des MGP mais également pour les MG déterministes. Dans le second cas, l'ensemble des configurations possibles du MGP sont énumérées et on pratique de l'élimination "en masse" grâce à des minorants et des majorants de la probabilité à optimiser. Comme pour l'estimation, les méthodes approchées sont soit par simulation soit de type variationnel. Pour ce problème, les méthodes par passage de messages sont aussi très

efficaces. Il y a quelques années un pont a été établi entre ces dernières et les méthodes variationnelles, ouvrant la porte à diverses extensions et permettant une meilleure compréhension du fonctionnement de ces méthodes.

Enfin, en décision séquentielle dans l'incertain, les algorithmes classiques de résolution exacte d'un PDM font appel soit à la programmation linéaire, soit à la programmation dynamique [97]. Pour des problèmes réels, l'espace d'état du système est trop grand pour pouvoir utiliser ces méthodes. Pour la résolution approchée, nous retrouvons alors les approches par simulation, avec l'ensemble des travaux autour de l'apprentissage par renforcement [108] : il s'agit alors d'apprendre la stratégie optimale à partir de trajectoires simulées. Une autre solution consiste à travailler sur une approximation (par exemple linéaire, [23]) de la fonction de valeur du PDM. Jusqu'à très récemment les méthodes variationnelles n'étaient pas connues dans le domaine de la décision séquentielle dans l'incertain.

1.4 Les méthodes variationnelles pour l'estimation, l'inférence et la décision approchées dans les MGP

Les méthodes variationnelles constituent le fil conducteur de mes activités de recherche. La plupart de mes travaux y font appel, soit exclusivement soit en association avec de la simulation, afin de résoudre les trois tâches citées précédemment. Pourquoi les méthodes variationnelles ?

- Jusqu'à il y a encore quelques années, ces méthodes reconnues en mécanique statistique [18] n'avaient pas encore beaucoup diffusé vers d'autres domaines et il fallait tester leur intérêt pour les MGP.
- L'alternative la plus largement utilisée est la simulation, qui est très coûteuse en temps et/ou mémoire, et demande en pratique de savoir manipuler plusieurs boutons de réglage. Les méthodes variationnelles conduisent à des algorithmes plus rapides et plus simples à mettre en œuvre. Il est vrai que le temps d'exécution d'un algorithme n'est pas toujours une limitation, si la tâche ne doit être effectuée qu'une fois. Cependant, pour résoudre certains problèmes on peut être amené à la répéter un grand nombre de fois. Par exemple, pour sélectionner le meilleur modèle parmi M pour représenter un processus dans chacune des C conditions différentes où il a été observé, $M \times C$ tâches d'estimation doivent être mises en œuvre. De la même manière, les méthodes itératives de résolution de problèmes d'optimisation de la décision requièrent de résoudre à chaque étape des problèmes d'inférence.
- Les trois tâches d'estimation, inférence et décision ne sont pas indépendantes. Les méthodes d'estimation ainsi que celles pour la résolution des problèmes de décision sont gourmandes en calcul de probabilités globales ou marginales. Aussi il était intéressant de développer les méthodes variationnelles dans les trois afin de disposer d'une approche globale cohérente et par exemple de pouvoir réutiliser des calculs approchés en inférence dans des algorithmes pour la décision ou l'estimation.

Les premiers résultats en estimation et inférence que j'ai obtenu en thèse et plus généralement ceux de la communauté [84, 112] des MGP ont été encourageants. J'ai donc naturellement poursuivi dans cette direction. Aujourd'hui les méthodes variationnelles sont

bien présentes sur les deux tâches d'estimation et inférence. Par contre elles restent encore très novatrices en décision, axe sur lequel je me suis plus particulièrement concentré ces dernières années.

1.5 Applications en épidémiologie et en écologie

Les différentes contributions méthodologiques que je présente dans ce document ont pour la plupart été mises en œuvre pour répondre à des questions issues de problématiques en épidémiologie et en écologie. Ces travaux plus finalisés sont une façon de confronter des résultats à des conditions d'utilisation réelles. Ils sont également des moyens de mettre en évidence leurs limites et des sources de nouvelles pistes de recherche à creuser. Comme j'ai essayé de l'illustrer au début de cette introduction, en épidémiologie et en écologie la notion de réseau est prégnante. Il est assez naturel d'utiliser les MGP pour modéliser les processus en jeu. Il faut néanmoins être conscient de deux restrictions. D'une part, tous les processus qui impliquent des propagations à longue distance (dispersion de spores ou de graines par le vent) ne seront pas modélisés efficacement dans un MGP puisque le graphe sous-jacent serait complet. L'aspect factorisé de ces modèles, qui en fait l'intérêt, disparaîtrait. D'autre part, parfois de grosses simplifications de la réalité sont nécessaires (par opposition aux modèles mécanistes) pour pouvoir rester efficace en temps et en espace mémoire. Cependant ce n'est pas forcément rédhibitoire. En particulier lorsqu'il s'agit de concevoir des stratégies de gestion, une description fine des mécanismes n'est pas toujours nécessaire, et travailler sur des variables qualitatives peut suffire. J'ajouterais que de toutes manières les données nécessaires pour construire un modèle plus fin qu'un MGP ne sont pas toujours disponibles. Une fois ces limitations connues, il reste de très nombreux problèmes en épidémiologie et écologie pour lesquels les MGP sont pertinents, comme je l'illustrerai tout au long de ce document.

Chapitre 2

Plusieurs visions des méthodes variationnelles

Si mes travaux font régulièrement appel aux méthodes variationnelles c'est également par un goût pour les jolies mathématiques qui les définissent. Elles ne sont d'ailleurs pas toujours introduites avec les mêmes concepts selon les disciplines, même si des ponts existent. A l'origine les méthodes variationnelles ont été introduites en mécanique statistique. Elles ont ensuite été diffusées, à des rythmes différents, en machine learning, statistique computationnelle et écologie théorique, et portent parfois des noms différents (approximation de Kikuchi, de Kirkwood, méthode de fermeture des moments). Je présente ici différentes visions des méthodes variationnelles, selon la discipline : en mécanique statistique il existe deux entrées vers les méthodes variationnelles, que je présente en 2.1 et 2.2, les deux motivées par le calcul de l'entropie d'un système. La communauté du machine learning s'est intéressée aux méthodes variationnelles pour proposer des algorithmes efficaces d'inférence approchée (sous-section 2.3) dans les modèles graphiques. Les méthodes variationnelles ont aussi été introduites dans la communauté des statistiques pour définir une borne inférieure de la vraisemblance des données observées (sous-section 2.4). Enfin, en écologie théorique, les méthodes variationnelles sont utilisées comme des méthodes de résolution approchée d'un système d'équations différentielles (sous-section 2.5).

2.1 Mécanique statistique et énergie libre

L'approximation variationnelle définit une famille d'approximations dont le principe du champ moyen est le plus simple, mais aussi le plus utilisée car la plus facile à mettre en œuvre. L'approximation en champ moyen se justifie en introduisant un principe variationnel issu de la mécanique statistique : le principe de minimisation de la fonctionnelle énergie libre $F(p)$, également appelée énergie libre variationnelle dans certaines références (voir par exemple [85], [60]). Étant donné $p(x)$ une distribution de probabilité sur l'espace Ω des configurations possibles de $X = \{X_1, \dots, X_n\}$, et une fonction énergie $H(x) = \sum_{c \in \mathcal{C}} H_c(x_c)$ où $\mathcal{C}$ est un sous ensemble des parties de $V = \{1, \dots, n\}$, la fonctionnelle $F(p)$ est définie par

$$F(p) = \mathbb{E}_p[H(X)] - \mathcal{S}(p) ,$$

où $\mathbb{E}_p$ est l'espérance pour la distribution p et $\mathcal{S}(p) = -\mathbb{E}_p[\log p(X)]$ est l'entropie de p . Soit maintenant P_H la distribution du modèle graphique associée à H , c'est à dire

$$P_H(X = x) \propto \prod_{c \in \mathcal{C}} \Psi_c(x_c),$$

avec $\Psi_c(x_c) = \exp(-H_c(x_c))$ ¹. La constante de normalisation de P_H est définie par

$$Z_H = \sum_{x \in \Omega} \prod_{c \in \mathcal{C}} \exp(-H_c(x_c)).$$

L'énergie libre variationnelle se réécrit alors comme

$$\begin{aligned} F(p) &= -\log Z_H + \mathbb{E}_p[\log(\frac{p(X)}{P_H(X)})] \\ &= -\log Z_H + KL(p, P_H), \end{aligned}$$

où $KL(p_1, p_2) = \mathbb{E}_{p_1}[\log(\frac{p_1(X)}{p_2(X)})]$ est la divergence de Kullback-Leibler entre deux distributions p_1 et p_2 de même support. Cette quantité est positive et ne s'annule que si les distributions p_1 et p_2 sont identiques. À partir de ces propriétés, il est facile de voir que la distribution p qui minimise $F(p)$ est P_H . L'énergie libre minimale vaut donc

$$F(P_H) = -\log Z_H.$$

Considérons maintenant, non plus le vrai minimum de l'énergie libre (ou de manière équivalente le minimum de la divergence de Kullback-Leibler), mais le minimum parmi l'ensemble $\mathcal{Q}$ des distributions de probabilité p qui s'expriment comme le produit des marginales d'ordre 1 uniquement (i.e. ne faisant intervenir qu'une seule variable) : $p(x) = \prod_{i \in V} p_i(x_i)$. On se restreint ainsi à des systèmes de variables aléatoires indépendantes. La distribution dans $\mathcal{Q}$ solution de cette minimisation est l'approximation en champ moyen de P_H . En pratique le problème d'optimisation est résolu par la méthode des multiplicateurs de Lagrange afin de prendre en compte les contraintes de normalisation sur les $p_i(x_i)$. La solution s'exprime comme la solution d'un système d'équations de point fixe, que l'on résout par une méthode classique. Notons que les marginales d'ordre 1 obtenues ne sont pas celles de la distribution P_H .

L'approximation en champ moyen est la méthode variationnelle d'ordre 1. Une méthode variationnelle à l'ordre k fournit une approximation de la loi jointe qui ne fait intervenir que des marginales d'ordre au plus k . À l'ordre 2, le problème consiste à identifier, parmi l'ensemble $\mathcal{Q}$ des distributions sur Ω qui s'expriment comme un produit de marginales d'ordre 1 et d'ordre 2 (i.e. faisant intervenir une paire de variables), celle qui minimise la distance de l'énergie libre variationnelle. Lorsque l'on remplace cette fonctionnelle par l'énergie libre de Bethe ([116]), on obtient l'approximation de Bethe, qui s'exprime ainsi

$$p^{Bethe}(x) = \frac{\prod_{(i,j) \in E} p^{Bethe}(x_i, x_j)}{\prod_{i \in V} p^{Bethe}(x_i)^{d_i-1}}, \quad (2.1)$$

1. L'écriture de la distribution d'un MGP sous cette forme faisant intervenir l'exponentielle d'une énergie est en général appelée distribution de Gibbs.

où $G = (V, E)$ est le graphe associé à P_H et d_i est le degré du nœud i dans ce graphe. Cette approximation est exacte lorsque G est un arbre : la loi jointe et les marginales d'ordre 1 et 2 de p^{Bethe} sont celles de P_H . Pour un graphe quelconque, p^{Bethe} est une approximation. En particulier, elle n'est pas normalisée. De plus la résolution analytique du problème de minimisation sous contrainte (i.e. l'expression des marginales de p^{Bethe} en fonction de la distribution P_H) est inaccessible. Nous verrons dans la sous-section 2.3 qu'il existe un algorithme efficace pour calculer un point stationnaire du problème, développé à l'origine en machine learning, de manière totalement indépendante des méthodes variationnelles !

En conclusion, nous pouvons donc voir l'approximation variationnelle comme revenant à résoudre un problème d'optimisation sous contraintes dans le but d'approcher au mieux une distribution complexe par une distribution appartenant à une famille plus simple, s'exprimant en fonction de marginales d'ordre petit, sur laquelle on saura mener des calculs.

2.2 Mécanique statistique et approximation de Kikuchi

Toujours en mécanique statistique, nous retrouvons l'idée d'approcher une distribution jointe par une fonction ne faisant intervenir que des produits de marginales d'ordres raisonnables, avec l'approximation de Kikuchi [60]. Considérons toujours une distribution d'un MGP p , définie sur $\Omega = \Omega_1 \times \dots \times \Omega_n$ de graphe associé $G = (V, E)$. L'idée est d'exprimer la distribution p comme un produit de marginales d'ordres croissants, puis de tronquer en remplaçant par 1 tous les termes d'ordre supérieur à l'ordre choisi ([61, 29]). Comment choisir alors les marginales qui interviennent dans le produit ? Supposons donné un sous-ensemble $\mathcal{R}$ des parties de V , contenant V . Il est possible de construire de manière itérative un ensemble de fonctions positives $C_A(x_A)$ pour tout $A \in \mathcal{R}$, telles que

$$\forall B \in \mathcal{R}, p(x_B) = \prod_{A \in \mathcal{R} \text{ t.q. } A \subset B} C_A(x_A).$$

En particulier,

$$p(x) = \prod_{A \in \mathcal{R}} C_A(x_A).$$

Il est facile de voir que chaque fonction $C_A(x_A)$ s'exprime comme un produit de marginales de la distribution p . Prenons le cas de l'ensemble $\mathcal{R}$ constitué de tous les nœuds de V , de toutes les paires de nœuds de E , et de V . Alors $C_i(x_i) = p(x_i)$ et $C_{ij}(x_i, x_j) = \frac{p(x_i, x_j)}{p(x_i)p(x_j)}$, ce qui conduit à :

$$p(x) = \frac{\prod_{(i,j) \in E} p(x_i, x_j)}{\prod_{i \in V} p(x_i)^{d_i-1}} \times C_V(x).$$

Si l'on tronque à l'ordre 1, on obtient l'approximation en champ moyen, et si l'on tronque à l'ordre 2 on retrouve l'approximation de Bethe.

2.3 Machine learning et algorithmes de passage de messages

Le calcul des marginales d'ordre 1 et 2 dans un MGP est un calcul de type Somme-Produit, bien connu en machine learning ([85]) :

$$\begin{aligned} p(x_i) &= \frac{1}{Z} \sum_{x_{\setminus i}} \prod_{c \in \mathcal{C}} \Psi_c(x_c), \\ p(x_i, x_j) &= \frac{1}{Z} \sum_{x_{\setminus i, j}} \prod_{c \in \mathcal{C}} \Psi_c(x_c), \end{aligned}$$

où $x_{\setminus i}$ désigne le vecteur x privé de x_i . Une manière brutale serait de réaliser un calcul pour chacune des marginales, chaque fois par énumération complète des configurations $x_{\setminus i}$. Une manière plus intelligente de réaliser les calculs consiste à exploiter le fait que certains des calculs dans $\sum_{x_{\setminus i}} \prod_{c \in \mathcal{C}} \Psi_c(x_c)$ peuvent d'une part être effectués indépendamment les uns des autres, d'autre part sont à effectuer dans le calcul de plusieurs marginales et donc peuvent être mutualisés. Ces constats ont conduit à la définition de l'algorithme Sum-Product, de la famille des algorithmes par passage de messages (voir [85, 9, 116]). Chaque nœud va recevoir un message de chacun de ses voisins dans G puis renvoyer vers ses voisins les messages mis à jour. Les messages sont proportionnels à des fonctions potentiel cumulées. Concrètement, pour un MGP pairwise², le message de i vers j est défini ainsi (la constante k sert à la normalisation) :

$$m_{ij}(x_j) \leftarrow k. \sum_{x_i \in \Omega_i} \Psi_{ij}(x_i, x_j) \Psi_i(x_i) \prod_{s \in N(i) \setminus j} m_{si}(x_i).$$

L'ensemble $N(i)$ représente l'ensemble des nœuds reliés à i par une arête dans G (les voisins). Si G est un arbre, il suffit de parcourir le graphe de la racine vers les feuilles puis des feuilles vers la racine et les messages ainsi accumulés au niveau de chaque nœud permettent de calculer les quantités suivantes, qui correspondent exactement aux marginales d'ordre 1 et 2 :

$$\begin{aligned} b_i(x_i) &= k. \Psi_i(x_i) \prod_{s \in N(i)} m_{si}(x_i) \\ b_{ij}(x_i, x_j) &= k. \Psi_{ij}(x_i, x_j) \Psi_i(x_i) \Psi_j(x_j) \prod_{s \in N(i) \setminus j} m_{si}(x_i) \prod_{s \in N(j) \setminus i} m_{sj}(x_j) \end{aligned}$$

Lorsque le graphe n'est pas un arbre, les messages sont mis à jour sur plusieurs itérations, jusqu'à "convergence". Les quantités b_i et b_{ij} fournissent alors des approximations des marginales d'ordre 1 et 2. Des travaux relativement récents ont établi un lien entre les solutions de cet algorithme et les solutions du problème de minimisation de l'énergie libre de Bethe qui définit l'approximation de Bethe (sous-section 2.1). En effet, il a été démontré que les b_i et b_{ij} points fixes de l'algorithme sont des points stationnaires du problème de minimisation [116].

2. La loi jointe d'un MGP pairwise est définie uniquement à partir de fonctions potentiel d'ordre 1 et 2

2.4 Statistique et borne inférieure de la vraisemblance

Considérons maintenant le problème de l'estimation du paramètre θ d'une distribution de probabilité p (de type MGP ou non) à partir d'un jeu de données. Supposons que l'ensemble X des variables du modèle est découpé en deux sous-ensembles disjoints O et H , tels que seul O est observé. H correspond à des données manquantes ou des variables cachées. L'estimateur du maximum de vraisemblance pour une réalisation o des données observées est alors la valeur θ qui maximise la vraisemblance $\mathcal{L}(o|\theta) = \ln p(o|\theta)$. Considérons maintenant une distribution q_h quelconque définie sur le même espace d'état que la distribution des variables H sous p , il est possible par un simple jeu d'écriture de définir une borne inférieure de la vraisemblance ([44]) :

$$\begin{aligned}\mathcal{L}(o|\theta) &= \ln p(o|\theta) \\ &= \ln \sum_h p(h, o|\theta) \\ &= \ln \sum_h q_h(h) \frac{p(h, o|\theta)}{q_h(h)} \\ &\geq \sum_h q_h(h) \ln \frac{p(h, o|\theta)}{q_h(h)} \text{ par l'inégalité de Jensen} \\ &= \sum_h q_h(h) \ln p(h, o|\theta) - \sum_h q_h(h) \ln q_h(h) \\ &\equiv \mathcal{F}(q_h, \theta).\end{aligned}$$

La fonctionnelle $\mathcal{F}(q_h, \theta)$ peut être réécrite en terme de distance de Kullback-Leibler :

$$\mathcal{F}(q_h, \theta) = \mathcal{L}(o|\theta) - KL(q_h(\cdot) \mid p(\cdot \mid o))$$

Ainsi, la meilleure borne inférieure est obtenue en recherchant la distribution q_h qui minimise cette distance. Comme dans la section 2.1, on va restreindre l'exploration à une famille $\mathcal{Q}$ de distributions de probabilité facilement manipulables. Par exemple si H est un vecteur de n variables aléatoires, il est possible de considérer la famille des distributions de n variables indépendantes et de retrouver la définition de l'approximation en champ moyen.

2.5 Dynamique des populations et fermeture des moments

En écologie théorique, les systèmes décrivant des dynamiques de population sont complexes et là aussi parmi les méthodes existantes pour obtenir des descriptions approchées, on retrouve le champ moyen et Bethe. Considérons par exemple le processus de contact sur graphe, de la famille des systèmes de particules en interaction ([47] ou [74, chapitre 6]). Il s'agit d'un modèle dynamique stochastique à temps continu³ utilisé pour modéliser la propagation d'épidémies ou d'individus sur un réseau. A chaque nœud du graphe G est associée une variable $X_i \in \{0, 1\}$ dont l'état varie au court du temps selon les règles suivantes : si i

3. Une définition en temps discret existe aussi

est occupé (état 1), alors il devient vide avec un taux μ , que l'on peut fixer à 1 sans perte de généralité. Si i est vide et que a_i de ses voisins sont occupés, alors il devient occupé à son tour avec un taux βa_i , où β est le taux de contamination par nœud. Pour étudier le comportement du modèle (transition de phase en particulier), il faut établir le système d'équations différentielles ordinaires qui le régit. Ainsi il est possible montrer dans le cas d'un graphe de degré homogène h que

$$\frac{d\rho}{dt} = h\beta p(01) - \rho$$

où ρ est la probabilité qu'un nœud quelconque du graphe soit occupé, et $p(01)$ est la probabilité qu'une paire de nœuds voisins soit dans l'état "(vide, occupé)". L'évolution de $p(01)$ dépend de l'évolution des probabilités marginales sur des triplets de nœuds et ainsi de suite, de sorte que le système complet d'équations décrivant la dynamique du système est bien trop grand pour être étudié tel quel. Le principe des méthodes de fermeture des moments est de tronquer ce système à un certain ordre r et d'approcher toutes les probabilités marginales d'ordre supérieur à r par des expressions ne faisant intervenir que des marginales d'ordre inférieur ou égal à r . Ce choix de l'ordre de l'approximation est à rapprocher du choix de la famille $\mathcal{Q}$ introduite dans les paragraphes précédents. Ainsi, si la troncature est à l'ordre 1, une seule approximation de $p(x_i, x_j)$ est possible : $p(x_i, x_j) = p(x_i)p(x_j)$. Nous retrouvons l'approximation en champ moyen. Si la troncature est à l'ordre 2, plusieurs expressions ont été proposées pour exprimer la probabilité marginale de triplets de variables en fonction des marginales d'ordre 1 et 2 ([28, chapitres 13, 18, 19 & 21]). Parmi elles, nous retrouvons l'approximation de Bethe (2.1). Ainsi, si i, j et k sont trois nœuds du graphe formant une clique, l'approximation de Bethe de $p(x_i, x_j, x_k)$ est $p(x_i, x_j)p(x_j, x_k)p(x_k, x_i)/p(x_i)p(x_j)p(x_k)$. Si une arête n'existe pas dans E , la probabilité de la paire est remplacée par le produit des marginales d'ordre 1. Notons que dans la définition d'une méthode de fermeture des moments n'intervient jamais l'idée de recherche du meilleur représentant de la distribution d'origine par minimisation de la divergence de Kullback-Leibler. Il s'agit plutôt de remplacer un modèle par un autre.

2.6 Conclusion

Les méthodes variationnelles sont donc des méthodes d'approximation utiles pour la résolution de problèmes de natures différentes. Cela va être illustré tout au long de ce mémoire. Ainsi, l'algorithme de résolution approchée d'un PDM sur graphe décrit dans la section 3.4 repose sur l'idée de rechercher la meilleure approximation d'une distribution complexe parmi un sous-ensemble de distributions "plus simples", comme en mécanique statistique (sous-section 2.1). La définition des méthodes variationnelles comme troncature de la décomposition de Moebius d'une distribution jointe (sous-section 2.2) a permis de caractériser la solution du maximum d'entropie sous contraintes de marginales (voir section 3.1). L'algorithme LSDP pour la conception de stratégies adaptatives d'échantillonnage dans les champs de Markov (section 3.4) fait appel à des calculs approchés de marginales pour lequel nous avons exploité l'algorithme de passage de messages (sous-section 2.3). Pour l'estimation approchée dans les champs de Markov cachés et dans le modèle de Cox log gaussien (section 3.2), nous nous sommes appuyés sur la définition de l'approche variationnelle

comme borne inférieure de la vraisemblance des données observées (sous-section 2.4). Enfin, pour décrire l'influence de la structure du réseau dans la dynamique d'un modèle de processus de contact (section 3.3), nous avons approché le système d'équations différentielles régissant cette dynamique grâce à la méthode de fermeture des moments (sous-section 2.5)

Chapitre 3

Contributions

Je présente maintenant une synthèse de mes contributions autour des différentes tâches associées aux MGP : estimation (sous-section 3.2), inférence (sous-section 3.3) et décision (sous-section 3.4). Avant cela, j'introduis un ensemble de résultats autour du travail préliminaire de construction du MGP à partir d'un jeu de données (sous-section 3.1). Une dernière section de synthèse et réflexion (sous-section 3.5) conclue cette partie.

3.1 Construction du modèle

Par construction d'un MGP, j'entends l'étape par laquelle on réalise des choix de modélisation à partir de l'analyse d'un jeu de données. Il s'agit d'évaluer quelle est la structure présente dans ces données afin de la prendre en compte dans le modèle qui sera élaboré. Cela nécessite une phase d'analyse exploratoire des données et de réalisation de tests d'hypothèse. Une autre approche consiste à voir cette étape de construction comme un problème de recherche du modèle qui incorpore l'information présente dans les données, mais pas plus : cela revient à résoudre un problème de maximum d'entropie sous contraintes.

3.1.1 Analyse statistique de données sur grille par tests de permutation

Principaux collaborateurs : Joël Chadœuf (BIOSP - INRA MIA - Avignon), Daniel Nandris et Frédéric Pellegrin (IRD - Montpellier), Gaël Thébaud (UMR BGPI - INRA SPE - Montpellier)

Publications associées¹ : 4, 6, 8, 10.

Les tests de permutation [73, 77]) font référence à des méthodes pour l'exploration d'un jeu de données sans hypothèse sur leur distribution. Ces méthodes, simples à mettre en œuvre, permettent de mettre en évidence des caractéristiques sous-jacentes aux données, qui pourront être traduites dans un modèle de distribution dans une seconde étape. Le principe est le suivant : les observations sont réallouées par permutation puis la valeur d'une statistique calculée sur les observations est comparée à sa valeur calculée sur un grand nombre de permutations, afin de détecter d'éventuels écarts significatifs. Plus précisément,

1. Les numéros des publications se réfèrent à la liste de mes publications, Section 6.

soit $x^{obs} = \{x_1, x_2, \dots, x_N\}$ le vecteur des observations, la première étape consiste à définir l'hypothèse nulle H_0 (en général elle traduit une propriété d'indépendance). L'hypothèse nulle va déterminer les permutations possibles des N valeurs observées. En effet, les permutations sont choisies de sorte que sous H_0 elles aient toutes la même probabilité, qui est aussi celle de x^{obs} . Un grand nombre de permutations est généré parmi celles autorisées et pour chacune la valeur $T(x^{perm})$ de la statistique du test est calculée. Le choix de T est guidé par l'hypothèse alternative H_1 . Si la statistique $T(x^{obs})$ calculée sur x^{obs} est de faible p -valeur dans l'histogramme construit à partir des valeurs calculées sur les permutations ($\{T(x^{perm})\}$) alors on rejette H_0 .

Ce type de tests est largement utilisé dans le cas d'observations indépendantes, par exemple pour tester des égalités de moyennes entre deux échantillons. Dans le cadre de données observées sur grille, la mise en évidence de la structure spatiale des données conduit à un ensemble de tests spécifiques. Le test le plus élémentaire consiste à tester l'indépendance totale des variables aléatoires associées à chaque nœud de la grille. Si cette indépendance globale est rejetée, il est ensuite possible de rechercher une structure purement intra-lignes et pour cela tester l'indépendance entre deux lignes de la grille. Les tests de permutation peuvent également être utilisés pour comparer deux jeux d'observations acquis sur une même grille. Il peut s'agir d'un même processus observé à deux dates successives, ou bien de deux processus observés sur un même site. Se pose alors la question de l'indépendance ou non des deux jeux d'observations. La difficulté dans l'étude de deux images provient de la possible existence d'une structure spatiale propre à chacune. Il faut la prendre en compte dans la construction du test afin de ne pas rejeter H_0 à tort du fait que les permutations ont cassé cette structure. La seconde difficulté vient de la possible présence de données manquantes sur la grille (censure).

J'ai participé à la rédaction d'un article de revue sur la manière d'aborder la vérification d'un certain nombre d'hypothèses classiques du cas spatial dans le cadre des tests de permutation. Cet article ([88]) traite du cas de l'analyse d'une image et de la comparaison de deux images. Ce panorama, à destination de chercheurs modélisateurs en épidémiologie ou écologie, m'a permis de me familiariser avec cet outil. J'ai ensuite participé à l'étude de la question plus particulière de l'analyse de l'indépendance de cas de maladies observés à deux dates successives dans un verger, en présence de censure non liée au phénomène observé (par exemple un arbre mort ou arraché suite à une autre maladie). Nous nous sommes placés dans le cadre suivant : les plantes sont régulièrement espacées le long des lignes et des colonnes de la grille, et une plante peut être dans l'un des état suivants : manquante, saine, observée malade pour la première fois à la date 1 (cas 1) ou observée malade pour la première fois à la date 2 (cas 2). D'autre part, les potentielles inoculations successives d'une même plante ne peuvent pas être observées : une même plante ne peut pas être à la fois un cas 1 et un cas 2. L'hypothèse nulle testée est : la position des cas 2 est indépendante de la position des cas 1. L'hypothèse alternative correspond à un excès d'agrégation ou à l'inverse un excès de régularité. Ainsi la statistique de test qui a été choisie est la fréquence cumulée du nombre de paires de type (cas 1, cas 2) dans un disque de rayon donné. L'utilisation de la fréquence cumulée permet d'accroître la puissance du test et de réduire la variabilité des courbes obtenues. Les tests que nous avons proposés ([109]) tiennent compte de la structure spatiale à chaque date afin de distinguer un écart à l'indépendance dû à cette structure et un écart à l'indépendance dû à une corrélation entre les anciens et les

nouveaux cas. Ainsi, selon les cas (indépendance des cas 2, structuration spatiale des cas 2 uniquement, ou structuration spatiale aux deux dates), des permutations différentes sont utilisées (réallocation au hasard ou au contraire translations). Le problème de la censure apparaît lorsque dans une permutation un cas 2 tombe sur un point de la grille où un cas 1 a été observé ou bien sur un point de la grille qui correspond à un arbre manquant. Ces deux situations sont impossibles dans la réalité. Il faut alors exclure les cas 1 correspondant au moment de calculer les valeurs de la statistique de test.

Je n'ai pas poursuivi d'activité méthodologique sur la définition de nouveaux tests de permutation dans un contexte spatial. Par contre, il s'agit d'un outil que je continue à utiliser et à recommander à des collègues épidémiologistes et à des étudiants en statistique. Ils sont simples et peuvent aider à l'élaboration d'un modèle spécifique de propagation d'une maladie.

Ainsi, nous avons utilisé les tests de permutation pour la comparaison de cas de maladie à deux dates successives pour étudier deux cas de propagation spatio-temporelle d'épidémie aux comportements contrastés : le virus Plum pox (sharka) chez le pêcher et la European stone fruit yellows (ESFY) chez l'abricotier ([109]). Dans le cas de la sharka, l'étude a révélé d'une part une structuration spatiale à chaque date, et d'autre part une corrélation spatiale entre deux dates successives. Un modèle de la dynamique de la sharka devrait donc nécessairement intégrer cette composante spatiale (champ de Markov dynamique ou réseau bayésien dynamique). Dans le cas de l'ESFY, les tests ont permis d'identifier une tendance à l'agrégation intra-rang entre les cas 1 et les cas 2. Par contre il n'y a pas d'évidence pour une structuration spatiale entre cas 1, ou entre cas 2. Si un modèle des dynamiques spatio-temporelles de ESYF devait être construit, il serait donc raisonnable de faire l'hypothèse d'un état initial issu d'une répartition aléatoire et d'une anisotropie des corrélations inter-date.

Ces tests se sont également révélés utiles dans l'étude des causes possibles du développement de la nécrose de l'hévéa, des doutes persistant quant à une transmission par un pathogène ([91]). Des expérimentations ont été mises en place sur plusieurs années et sur plusieurs parcelles afin de pouvoir comparer l'évolution de l'épidémie selon qu'est adoptée ou non une pratique de désinfection systématique du couteau de collecte après chaque arbre (figure 3.1). Nous avons donc pu comparer les cas d'une même parcelle à des dates différentes, ainsi que les cas de parcelles avec et sans désinfection. Notre analyse a confirmé les résultats étiologiques précédemment obtenus par l'IRD : il n'y a pas d'évidence d'une propagation de la maladie par une transmission d'un pathogène (figure 3.2). Néanmoins la structuration spatiale de la maladie est confirmée. Ces résultats renforcent le scénario alternatif d'une maladie physiologique à plusieurs facteurs causée par une accumulation de stress.

Enfin, plus récemment, dans le cadre de la thèse de Mathieu Bonneau, afin d'appliquer l'algorithme LSDP (voir section 3.4) pour proposer des stratégies d'échantillonnage adaptatif des espèces adventices dans une parcelle, nous avons dans un premier temps réalisé un travail de modélisation de la répartition spatiale des adventices par une distribution de champ de Markov. Grâce aux tests de permutation, nous avons au préalable sélectionné parmi l'ensemble des configurations présentes dans les jeux de données (espèce adventice, culture, date, pratique), celles où les adventices présentaient une structure spatiale significative. Pour ces espèces, le choix des fonctions potentiel définissant le champ de Markov a

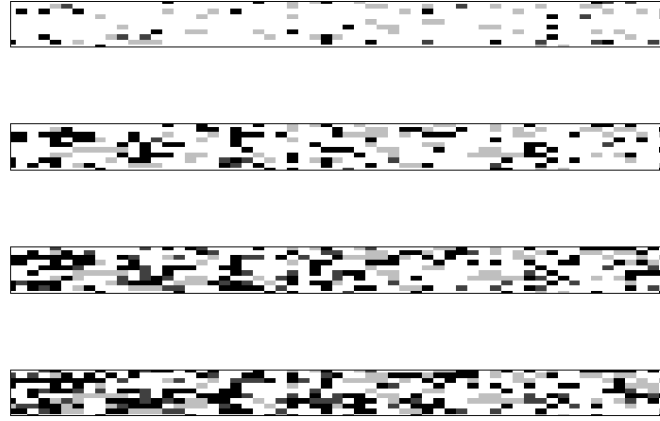


FIGURE 3.1 – Etat sanitaire des arbres à hévéa sur une des parcelles expérimentales sur 4 années successives, codé en niveau de gris : blanc = arbre sain, gris = arbre cassé ou manquant ou atteint d'autre maladie, noir = arbre atteint de la nécrose de l'hévéa.

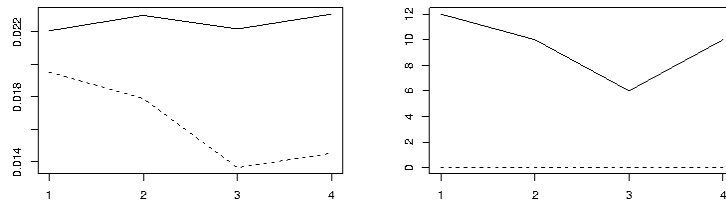


FIGURE 3.2 – Test d'indépendance temporelle, basé sur la proportion de paires formées d'un arbre atteint de la nécrose à la date 2 et d'un arbre dans la même rangée atteint à la date 1. L'axe horizontal représente la distance entre les deux arbres de la paire, la courbe continue représente le quantile à 95% pour la statistique calculée sur les permutations des données et la courbe en points donne la valeur de la statistique calculée sur les données observées. Les deux graphes donnent les résultats pour deux parcelles différentes.

été réalisé dans une étape de sélection de modèle (voir section 3.2).

3.1.2 Caractérisation du maximum d'entropie sous contrainte de marginales

Principaux collaborateurs : Alain Franc (BioGeCo - INRA EFPA - Bordeaux) et Michel Goulard (DYNAFOR - INRA MIA - Toulouse)

Publication associée : 11.

Le travail précédent permet une première exploration donnant des informations sur la structure spatiale dans les données, qu'il faut ensuite savoir traduire en terme de fonctions potentiel d'un MGP. A l'opposé l'approche que je décris maintenant est constructive en ce sens qu'elle permet de définir les fonctions potentiel à partir de fréquences empiriques calculées sur des observations. Ces fréquences empiriques peuvent porter sur une variable ou sur plusieurs, nous les appellerons marginales empiriques. Le problème posé est alors celui de déterminer une loi jointe cohérente avec les marginales empiriques. Une manière classique d'y répondre consiste à le formuler comme un problème de maximum d'entropie sous contraintes ([50]). Cette approche nous assure que l'information supplémentaire contenue dans le modèle solution, en plus de celle apportée par la contrainte, est minimale. Pour certaines contraintes particulières la solution du maximum d'entropie (ME) est bien connue et correspond à des distributions de probabilité classiques : les distributions uniformes sont solutions du ME lorsqu'il n'y a aucune contrainte ; les distributions de Gibbs sont celles qui maximisent l'entropie parmi les distributions jointes de même énergie moyenne ; enfin les distributions gaussiennes peuvent être définies comme solutions du ME sous contrainte d'espérance et de variance. Dans ces trois exemples, la solution du ME sous contrainte est donc connue analytiquement.

Lorsque les contraintes portent sur des marginales de la distribution jointe, si deux marginales portent sur deux sous-ensemble de variables non disjoints alors les contraintes correspondantes vont générer des propriétés d'indépendance conditionnelle (au sens Markovien) dans la solution du ME. Il est assez facile de voir que la structure de la solution du ME sous contrainte de marginales est celle d'un MGP ([115]). Plus précisément, la distribution jointe obtenue est le produit de certaines des marginales, élevées à une certaine puissance. Les exposants sont appelés exposants canoniques ([112]). Dans la plupart des cas, la résolution du ME est impossible analytiquement et les exposants canoniques ne sont donc pas identifiables.

Nous nous sommes posé les deux questions suivantes : pour quelles structures particulières de l'ensemble des marginales empiriques est-il possible d'identifier ces exposants ? Sous quelles conditions un MGP peut-il être obtenu comme la solution d'un ME sous contraintes ? Dans des cas spécifiques, le problème ME sous contrainte de marginales est très simple à résoudre. Si les marginales empiriques sont spécifiées sur des ensembles disjoints de variables, alors la distribution définie comme le produit de ces marginales est solution du ME. Si les marginales sont connues uniquement sur des paires de variables et si ces contraintes n'induisent pas de boucle², alors la solution du ME est un MGP dont le graphe

2. Dans ce contexte on parle de boucle s'il est possible de trouver un chemin partant d'une variable et retournant à la même variable, sans passer deux fois par la même variable intermédiaire (par exemple si les

associé G est un arbre. Les exposants canoniques sont connus ([85]) : les marginales sur les paires interviennent à la puissance 1, et les marginales sur les singletons à une puissance égale au degré du nœud correspondant dans l'arbre moins 1. Afin d'étendre ces résultats, nous avons mis en articulation des résultats connus sur les méthodes de maximum d'entropie ([50]), sur la théorie des graphes ([4]), sur les MGP ([66]) et sur les méthodes variationnelles ([112]). Cela nous a permis de définir une classe de contraintes sur les marginales, la classe des ensembles cohérents de contraintes cordaux et maximaux, pour laquelle nous identifions les exposants canoniques à partir de la structure de l'ensemble des contraintes ([40]).

Considérons le vecteur aléatoire $X = (X_1, \dots, X_n)$, indexé sur $V = \{1, \dots, n\}$, et $x = (x_1, \dots, x_n)$ une réalisation de $X : x \in \Omega_1 \times \dots \times \Omega_n$. Pour tout sous-ensemble I de V je noterai également X_I la collection de variables $\{X_i, i \in I\}$. Alors si $\bar{I} = V \setminus I$, la distribution marginale p_I de X_I est définie comme :

$$p_I(x_I) = \sum_{x_{\bar{I}} \in \Omega_{\bar{I}}} p(x_I, x_{\bar{I}}).$$

Dans le problème ME, avoir une contrainte sur la marginale p_I signifie qu'une marginale empirique a_I a été observée et que l'on contraint la distribution p à satisfaire l'égalité $p_I = a_I$. Considérons maintenant M un ensemble de parties de V , et un ensemble $\{a_I\}_{I \in M}$ de marginales empiriques indexées sur M . Nous supposons sans perte de généralité que M est stable par intersection et inclusion. Si ça n'est pas le cas, il suffit de considérer

$$M' = M \cup \{J : \exists I \in M : J \subset I\}.$$

L'ensemble des marginales indexé sur M est dit cohérent si pour tout I et J dans M et pour tout K inclus dans l'intersection de I et J alors la marginale des variables X_K calculée par intégration de la marginale de X_I est égale à celle calculée par intégration de la marginale de X_J :

$$\forall x_K \in \prod_{k \in K} \Omega_k, \quad \sum_{x_{I \setminus K}} a_I(x_K, x_{I \setminus K}) = \sum_{x_{J \setminus K}} a_J(x_K, x_{J \setminus K}).$$

Nous appellerons alors $a_M = \{a_I\}_{I \in M}$ un ensemble cohérent de contraintes (ECC). Les éléments maximaux de M sont appelés les générateurs de l'ECC.

Il est possible d'associer un graphe $\mathcal{G}_{ECC}$ à tout ECC a_M : les sommets sont les points de V et deux sommets i et j sont reliés par une arête si il existe $I \in M$ tel que $i \in I$ et $j \in I$. Ce graphe est connu sous le nom de graphe des contraintes en Satisfaction de Contraintes ([25]).

La famille des ECC cordaux et maximaux est alors définie comme l'ensemble des ECC tels que le graphe associé $\mathcal{G}_{ECC}$ est cordal et les cliques maximales de $\mathcal{G}_{ECC}$ correspondent aux générateurs de l'ECC (voir figure 3.3). Nous avons démontré la proposition suivante :

Proposition 1 *Soit a_M un ECC maximal et cordal. Si $\mathcal{C}$ et $\mathcal{S}$ sont respectivement l'ensemble des cliques maximales et des séparateurs³ $\mathcal{G}_{ECC}$ alors le MGP défini par*

$$p(x) = \frac{\prod_{c \in \mathcal{C}} a_c(x_c)}{\prod_{s \in \mathcal{S}} a_s(x_s)}$$

marginales des 3 paires suivantes sont connues, (x_1, x_2) , (x_2, x_3) , (x_3, x_1) , alors $(1, 2, 3)$ forme une boucle).

3. Un séparateur est une intersection de deux cliques maximales reliées par une arête dans l'arbre de jonction associé.

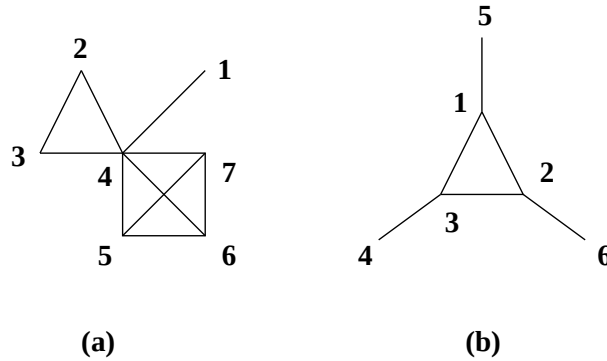


FIGURE 3.3 – Deux exemples de ECC cordaux et maximaux : les générateurs sont (a) $\{\{1, 4\}, \{2, 3, 4\}, \{4, 5, 6, 7\}\}$, et (b) $\{\{1, 2, 3\}, \{1, 5\}, \{2, 6\}, \{3, 4\}\}$.

est une distribution d'entropie maximale sous les contraintes définies par a_M .

L'exposant canonique associé aux cliques maximales est donc égal à 1. Et si un même sous-ensemble S de V est r_S fois séparateur, l'exposant canonique de la marginale correspondante, a_S , est égal à r_S .

Nous avons également montré que

Proposition 2 *Un MGP dont le graphe associé G est cordal est solution du maximum d'entropie sous contraintes de marginales, lorsque les marginales spécifiées sont celles sur les cliques maximales de G .*

Les résultats que nous avons obtenus reposent sur l'aspect théorie de l'information du problème de maximum d'entropie. Ce problème a aussi été défini et étudié en mécanique statistique et la question de l'identification des exposants canoniques n'est pas sans lien avec celle du choix d'un ordre de troncature dans la définition des méthodes variationnelles de la section 2.2. En effet, le résultat d'une décomposition puis troncature de $p(x)$ s'exprime comme un produit de marginales élevées à une certaine puissance. Une perspective aux travaux ci-dessus est d'apporter des éléments pour répondre à la question suivante : sous quelles conditions une approximation variationnelle d'une distribution jointe est-elle exacte ? Intuitivement, si la troncature est réalisée à un ordre inférieur à celui de la clique de plus grande taille dans G , l'approximation ne sera pas exacte. Que se passe-t-il si l'on tronque à un ordre supérieur ? La question reste ouverte.

3.2 Estimation dans les modèles graphiques

Les difficultés computationnelles liées à l'estimation des paramètres d'un MGP sont très bien détaillées dans l'ouvrage [63]. Mes travaux sur cette question se concentrent sur le cas particulier de l'estimation des paramètres d'un MGP lorsque des variables sont cachées ou lorsque des données sont manquantes. La notion de variables cachées (ou latentes) correspond à des variables n'ayant pas nécessairement de réalité physique mais qui permettent de simplifier le modèle de distribution des variables observées. C'est le cas par exemple en segmentation d'image : les variables observées sont des niveaux de gris ou des intensités en chaque

pixel et l'on souhaite regrouper les pixels en classes sur un double critère d'homogénéité des intensités intra classes et de proximité spatiale. La classe associée à chaque pixel est alors une variable cachée. L'objectif de la segmentation d'image est d'affecter une classe à chaque pixel, soit à partir d'une base d'apprentissage (segmentation supervisée) soit en même temps que les caractéristiques des classes sont estimées (segmentation non supervisée). Une variable cachée peut également être utilisée pour modéliser le processus aléatoire qui génère le comportement moyen d'un processus spatial observé. Ainsi dans un processus de Cox log gaussien, l'intensité du processus de Cox est définie comme l'exponentielle d'un champ gaussien. Ce champ aléatoire est la variable cachée du modèle.

La notion de donnée manquante se réfère plutôt au cas d'une variable qui n'est pas observée du fait de contraintes sur le terrain, mais qu'il est possible de rattacher à une mesure concrète, réelle. C'est le cas par exemple lorsqu'un processus spatial est échantillonné en un nombre réduit de points d'une zone d'étude ou lorsqu'une variable est trop coûteuse à observer directement et que l'on cherche à la reconstruire à partir de la connaissance d'une variable qui lui est corrélée. Dans la suite, je ne distinguerai plus les deux sources d'absence d'information car les méthodes d'estimation sont les mêmes. Je parlerai uniquement de variables cachées.

Le calcul de la vraisemblance des paramètres du modèle est plus compliqué lorsqu'il implique des variables cachées, puisque la distribution $p(o)$ des données observées s'exprime comme une marginale de la loi jointe $p(h, o)$ des variables observées et cachées, obtenue par intégration (ou sommation) sur l'espace d'état des variables cachées. Cette marginale n'a en général pas une forme connue.

L'algorithme EM ([26]) est la méthode classique pour calculer l'estimateur du maximum de vraisemblance dans un modèle avec variables cachées. Cet algorithme est à l'origine présenté comme une procédure de maximisation itérative de l'espérance de la logvraisemblance complète du paramètre. Une itération est découpée en deux étapes : une étape de calcul de l'espérance ci-dessus pour la valeur courante des paramètres (étape E) et une étape de mise à jour des paramètres par maximisation de cette espérance (étape M). Les deux étapes E et M de l'algorithme EM peuvent également être vues comme deux étapes de maximisation alternée d'une même fonction des paramètres du modèle et de la distribution des variables cachées ([81]). C'est ce point de vue que je présente ici (section 3.2.1) car le lien est alors immédiat avec les méthodes variationnelles telles que présentées dans la section 2.4. Mes travaux sur l'estimation des paramètres d'un modèle de champ de Markov caché reposent sur cette interprétation (section 3.2.2). De plus, cette vision de EM a récemment été exploitée pour proposer une méthode d'estimation des paramètres dans un cadre bayésien cette fois (algorithme EM bayésien variationnel, VBEM [3]). Comme pour l'algorithme EM, cette procédure est générale et pour chaque modèle il faut l'instancier et faire des choix. J'ai ainsi étudié la mise en œuvre de VBEM pour l'estimation des paramètres du processus de Cox log Gaussien (section 3.2.3).

3.2.1 EM, EM variationnel et EM bayésien variationnel

Les deux étapes de l'algorithme EM peuvent être décrites à partir de la fonctionnelle $\mathcal{F}(q_h, \theta)$ déjà présentée dans la section 2.4 :

$$\mathcal{F}(q_h, \theta) = \mathcal{L}(o|\theta) - KL(q_h(\cdot) \mid p(\cdot \mid o, \theta))$$

Algorithme EM

$$\text{Etape E : } q_h^{(t+1)} = \underset{q_h}{\operatorname{argmax}} \mathcal{F}(q_h, \theta^{(t)}) = \underset{q_h}{\operatorname{argmax}} KL(q_h(\cdot) \mid p(\cdot \mid o, \theta^{(t)}))$$

$$\text{Etape M : } \theta^{(t+1)} = \underset{\theta}{\operatorname{argmax}} \mathcal{F}(q_h^{(t+1)}, \theta) = \underset{\theta}{\operatorname{argmax}} \sum_h q_h^{(t+1)}(h) \ln p(h, o \mid \theta).$$

Dans l'étape E, $q_h^{(t+1)}$ est recherchée parmi l'ensemble des distributions sur le domaine de définition des variables cachées $H = \{h_1, \dots, h_n\}$. La solution de cette étape de maximisation est la distribution qui annule la divergence de Kullback-Leibler $KL(q_h(\cdot) \mid p(\cdot \mid o, \theta^{(t)}))$. Il s'agit de la distribution a posteriori des variables cachées : $q_h^{(t+1)}(h) = p(h \mid o, \theta^{(t)})$. Pour de nombreux modèles, et en particulier les MGP, la résolution exacte des étapes E et M est impossible puisque que cela requiert de calculer toute la distribution $p(h \mid o, \theta^{(t)})$.

Pour s'affranchir de cette complexité calculatoire, il est possible d'appliquer le principe variationnel tel que présenté en 2.4. Cela consiste à réduire l'espace de recherche dans l'étape E à un sous-ensemble $\mathcal{Q}$ de distributions pour lesquelles l'étape M est réalisable. L'algorithme EM est remplacé par l'algorithme suivant :

Algorithme EM variationnel (VEM)

$$\text{Etape E : } q_h^{(t+1)} = \underset{q_h \in \mathcal{Q}}{\operatorname{argmax}} \mathcal{F}(q_h, \theta^{(t)})$$

$$\text{Etape M : } \theta^{(t+1)} = \underset{\theta}{\operatorname{argmax}} \mathcal{F}(q_h^{(t+1)}, \theta)$$

En pratique, le sous-ensemble $\mathcal{Q}$ utilisé est toujours celui associé à l'approximation en champ moyen, car il rend les calculs très aisés : il s'agit de l'ensemble des distributions produits $q_h(h) = \prod_{i=1}^n q_{h_i}(h_i)$. Cela définit l'algorithme suivant :

Algorithme EM en champ moyen

$$\text{Etape E : } q_h^{(t+1)} = \underset{q_h = \prod_{i=1}^n q_{h_i}}{\operatorname{argmax}} \mathcal{F}(q_h, \theta^{(t)})$$

$$\text{Etape M : } \theta^{(t+1)} = \underset{\theta}{\operatorname{argmax}} \mathcal{F}(q_h^{(t+1)}, \theta)$$

L'étape E revient à minimiser la divergence de Kullback-Leibler entre q_h et p . Cela correspond à la résolution du système d'équations de point fixe décrit dans la section 2.1. Dans le cadre de l'estimation des paramètres d'un modèle de champ de Markov caché, d'autres choix peuvent être proposés pour $q_h^{(t+1)}$, comme je l'illustrerai dans la section 3.2.2.

L'approche variationnelle a ensuite été étendue au cas de l'estimation bayésienne. Cette dernière repose en général sur la simulation intensive de lois a posteriori. En 2003 Beal a proposé une alternative originale avec l'algorithme EM bayésien variationnel (VBEM, [3]). Il repose que les deux points clé suivant : *i*) les paramètres sont considérés comme des variables cachées au même titre que les variables H , *ii*) le sous ensemble $\mathcal{Q}$ choisi est celui des distributions $q(h, \theta)$ qui se factorisent en $q_h(h)q_\theta(\theta)$. L'hypothèse est donc faite que les variables cachées et les paramètres sont indépendants conditionnellement aux observations.

Le principe de maximisation alternée d’une fonctionnelle est repris, cette fois-ci avec

$$\mathcal{F}(q_h, q_\theta) = \mathcal{L}(o) - KL(q_h(\cdot)q_\theta(\cdot) \mid p_{h,\theta}(\cdot, \cdot \mid o))$$

L’algorithme VBEM consiste donc en une maximisation de $\mathcal{F}(q_h, q_\theta)$ alternativement en q_h et en q_θ . Les étapes E et M deviennent (dans le cas de variables cachées discrètes) :

Algorithme EM Bayésien Variationnel (VBEM)

Etape VBE : $q_h^{(t+1)} = \frac{1}{Z_h} \exp \int q_\theta^{(t)}(\theta) \ln p(h, o \mid \theta) d\theta$

Etape VBM : $q_\theta^{(t+1)} = \frac{1}{Z_\theta} p(\theta) \exp \sum q_h^{(t+1)}(h) \ln p(h, o \mid \theta)$

Les constantes Z_h et Z_θ sont des constantes de normalisation. Au delà de l’estimation des paramètres du modèle, l’intérêt de VBEM réside également dans le fait qu’il fournit directement un critère pour la comparaison de modèles. Il s’agit de la fonctionnelle $\mathcal{F}(q_h^{VBEM}, q_\theta^{VBEM})$ calculée pour les distributions q_h^{VBEM} et q_θ^{VBEM} fournies par VBEM. Dans [3], il est démontré que dans le cas d’un modèle exponentiel conjugué cette quantité est asymptotiquement de même nature qu’un BIC (modulo de potentiels problèmes de régularité et identifiabilité).

En pratique, VBEM est relativement aisé à mettre en œuvre dans le cas de modèles exponentiels conjugués. Dans les autres cas, l’hypothèse d’indépendance conditionnelle entre H et θ ne suffit pas à réduire suffisamment la complexité calculatoire et il faut proposer des solutions au cas par cas. Nous verrons ainsi dans le cas de l’estimation des paramètres du processus de Cox log gaussien que nous avons dû faire des hypothèses restrictives supplémentaires sur la famille $\mathcal{Q}$ (section 3.2.3).

Un point commun de ces différentes approximations variationnelles pour l’algorithme EM est qu’il est très difficile d’obtenir des résultats théoriques sur l’estimateur obtenu. Quelques résultats ([59]) existent pour l’algorithme VBEM dans le cas de modèles simples (modèles de mélange) ou dans des situations très spécifiques (champs gaussiens sur grille avec données manquantes réparties régulièrement). Néanmoins, VEM et VBEM sont de plus en plus utilisés pour leur bon comportement en pratique et leur facilité de mise en œuvre.

3.2.2 Algorithmes EM de type champ moyen pour les champs de Markov cachés

Principaux collaborateurs : Gilles Celeux (INRIA Futurs) et Florence Forbes (INRIA Rhône-Alpes)

Publications associées : 1, 2, 24, 25, 38, 39, 41, 42, 55, 56, 57.

Afin de traiter le problème de segmentation d’image décrit dans l’introduction de cette partie, il est nécessaire de prendre en compte les dépendances spatiales lors de la création des classes. Les modèles de mélange très utiles en classification non spatiale sont inadaptés dans ce cas. Il est alors classique d’utiliser un modèle de champ de Markov caché discret, associé à un graphe $G = (V, E)$ de type grille régulière : $V = \{1, 2, \dots, n\}$ et l’ensemble des arêtes E correspond aux 4 ou 8 plus proches voisins sur la grille (voir figure 3.4). Le vecteur des variables cachées $H = \{H_1, \dots, H_n\}$ représente les classes associées à chaque nœud de G (appelé pixel dans ce contexte) : $\Omega_i = \{1, \dots, K\}$.

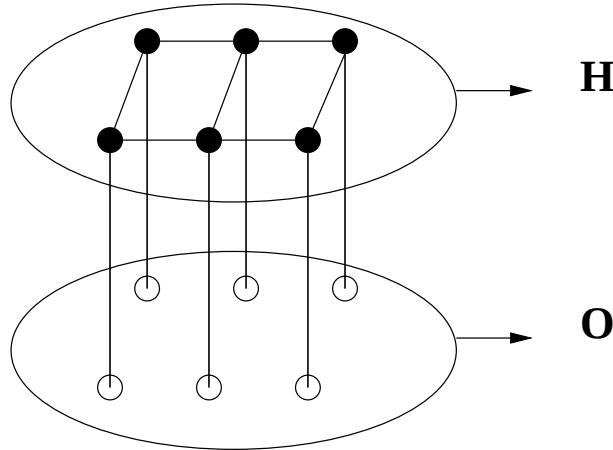


FIGURE 3.4 – Représentation sur une grille régulière des dépendances entre l’ensemble des variables observées (O) et cachées (H) dans un modèle de champ de Markov caché avec voisinage défini par les 4 plus proches voisins.

La distribution jointe de H est celle d’un MGP non orienté :

$$p(H = h) \propto \prod_{c \in \mathcal{C}} \Psi_c(h_c, \beta),$$

où β est un (vecteur de) paramètre(s). Dans le cas pairwise, l’ensemble $\mathcal{C}$ est composé de tous les pixels et de toutes les paires de pixels reliées par une arête dans G . Le voisinage $N(i)$ d’un nœud i est défini comme l’ensemble des nœuds reliés à i par une arête ou, de manière équivalente, l’ensemble des nœuds qui interviennent dans les mêmes fonctions potentiel que i . Dans le modèle de champ de Markov caché, les observations $O = \{O_1, \dots, O_n\}$ sont supposées indépendantes conditionnellement à H , de sorte que la densité conditionnelle $g(o | h)$ se décompose ainsi,

$$g(o | h) = \prod_{i \in V} g(o_i | h_i, \alpha),$$

où la forme des densités conditionnelles $g(o_i | k, \alpha)$ est connue et dépend d’un paramètre α . Le vecteur de paramètres du modèle pour les variables complètes (H, O) est noté $\theta = (\alpha, \beta)$.

Si le modèle de champ de Markov caché permet de mieux formaliser les caractéristiques attendues d’une segmentation d’image qu’un modèle de mélange, l’estimation des paramètres du modèle par EM et le calcul de la segmentation la plus probable ne sont pas appréhendables de manière exacte.

Une première famille de méthodes pour la résolution approchée consiste à alterner une étape de segmentation et une étape d’estimation (eg. [7]). Ainsi, dans l’étape de segmentation, les paramètres du modèle sont connus (cadre supervisé) et dans l’étape d’estimation la segmentation est connue (pas de variables cachées), ce qui simplifie les deux étapes. Cependant, ces approches sont connues pour introduire un biais sur l’estimation des paramètres, du fait qu’elles travaillent avec des segmentations “dures” de l’image, où chaque pixel est affecté à une et une seule classe. Les approches de type mélange, qui permettent de travailler avec des probabilités conditionnelles d’appartenance aux classes, évitent ce problème. Dans

ce cas, dans un premier temps l'algorithme EM (une version approchée en pratique) est utilisé pour estimer les paramètres, et ensuite la segmentation est construite (par exemple par calcul du mode de la distribution conditionnelle des classes). Si l'on utilise l'algorithme VEM, on peut observer en pratique que la version champ moyen de cet algorithme a tendance à conduire à des segmentations qui (même si souvent de bonne qualité, cf figure 3.5) sont surlissées par rapport au résultat attendu. Cela s'explique par le fait que la distribution $q_{h_i}(\cdot)$, solution de l'étape E, peut s'interpréter comme la distribution conditionnelle de H_i sachant que l'état des voisins $N(i)$ du nœud i est fixé à sa valeur moyenne. Cette moyennisation conduit au surlissage.

Néanmoins, le fait de fixer l'état du voisinage à une constante permet de se ramener à un ensemble de variables indépendantes et donc à un modèle de mélange simple pour lequel l'étape M peut être mise en œuvre de manière exacte. Nous avons donc conservé cette idée, et généralisé l'algorithme EM en champ moyen en proposant d'autres façons de fixer cette constante ([15]). La procédure générale proposée consiste à alterner deux étapes : une étape E de génération de l'état constant $h_{N(i)}$ du voisinage, puis une étape M identique à celle de VEM.

Algorithme EM de type champ moyen

Etape E : choix de $\tilde{h}^{(t+1)}$ puis calcul de $q_{h_i}^{(t+1)}(h_i) = p(h_i \mid H_{N(i)} = \tilde{h}_{N(i)}^{(t+1)}, o, \theta)$

Etape M : $\theta^{(t+1)} = \underset{\theta}{\operatorname{argmax}} \mathcal{F}(q_h^{(t+1)}, \theta)$

Cet algorithme peut s'interpréter également comme un algorithme EM pour modèle de mélange fini indépendant au cours duquel le modèle change de manière adaptative à chaque itération en fonction de l'approximation courante du voisinage. Il présente l'avantage de prendre en compte l'information spatiale existant dans le modèle markovien initial, tout en gardant la simplicité du cas indépendant.

Pour l'étape E, nous avons retenu comme alternative au champ moyen une version en champ simulé dans laquelle $\tilde{h}^{(t+1)}$ est une réalisation de la distribution conditionnelle courante des variables cachées (générée grâce à l'échantillonneur de Gibbs). Cet algorithme, appelé EM en champ simulé, a montré de bonnes performances en terme de paramètres estimés. Notamment il ne reproduit pas l'effet de surlissage des segmentations observé avec EM en champ moyen (figure 3.6).

Nous avons également appliqué le principe variationnel en champ moyen, et son extension en champ simulé, au calcul approché de critères de sélection de modèles (BIC et ICL) afin de choisir le nombre de classes K présentes dans la segmentation ([37]). Ici encore, l'approche variationnelle permet d'obtenir de bons résultats en pratique, avec des procédures plus rapides que les approximations par Monte Carlo des mêmes critères.

J'ai réalisé ces travaux lors de ma thèse, co-encadrée par Gilles Celeux et Florence Forbes. Ils constituent ma première expérience des méthodes variationnelles pour les modèles graphiques. Ces travaux sont très cités dans la communauté Analyse d'Image. Ils sont à la base du logiciel SpaCEM3 pour la classification de données complexes⁴.

4. spaCEM³ <http://spacem3.gforge.inria.fr/>

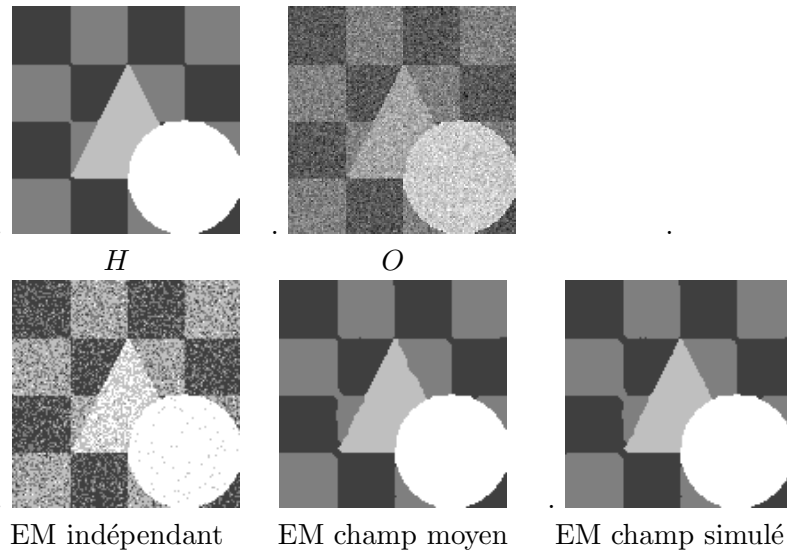


FIGURE 3.5 – Image damier dégradée par un bruit gaussien et segmentations fournies par différentes version de l’algorithme EM.

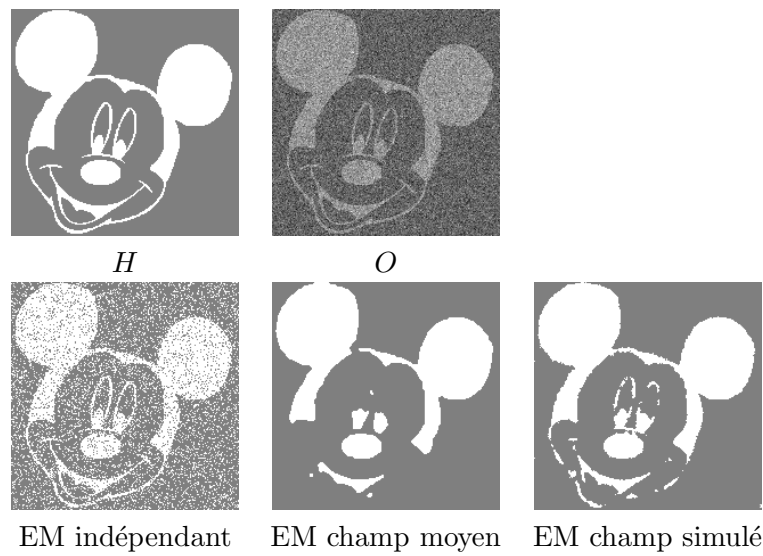


FIGURE 3.6 – Image Mickey dégradée par un bruit gaussien et segmentations fournies par différentes version de l’algorithme EM.

3.2.3 Algorithme VBEM pour le modèle de Cox log-Gaussien

Principaux collaborateurs : Julia Radoszycki et Régis Sabbadin (MIAT - INRA MIA - Toulouse)

Publication associée : 47.

La statistique spatiale est riche en modèles pour représenter les différents types de phénomènes qui peuvent être étudiés. Ainsi, les champs de Markov sont bien adaptés pour représenter des situations où observations et variables cachées sont définies sur un nombre fini de points de la zone d'étude, comme par exemple en chaque pixel d'une image. Lorsque le phénomène étudié peut a priori être observé en tout point de $\mathbb{R}^2$, le modélisateur se tourne plus naturellement soit vers les modèles géostatistiques (par exemple pour représenter la composition des sols) soit vers les modèles de type processus de points (par exemple pour représenter des comptages de graines, d'arbres, ...). Le processus de Cox log gaussien ([79]) fait partie de cette dernière famille. C'est un modèle avec variables cachées : le processus observé O est un processus de Poisson non homogène sur $\mathbb{R}^2$ d'intensité $\lambda_s = \exp(\beta + H_s)$. La structure spatiale qui existe entre les points du processus de Poisson est décrite à travers le processus caché $H = \{H_s\}_{s \in \mathbb{R}^2}$, qui est un champ gaussien de moyenne nulle et de fonction de covariance $C(s_1, s_2)$. La variable O_{A_i} représentant le nombre de points tombés dans la zone A_i , suit une loi de Poisson de paramètre $\Lambda(A_i) = \int_{A_i} \lambda_s ds$. Les hyperparamètres θ du modèle sont composés de β et des paramètres définissant C . Même si ce modèle ne ressemble a priori pas à un MGP, il est possible de s'y ramener. En effet, d'une part en pratique la zone d'étude est discrétisée en une partition d'aires de petite taille $\{A_i\}_{i=1, \dots, n}$ sur lesquelles sont observés les comptages $\{O_{A_i}\}_{i=1, \dots, n}$. Donc un nombre fini de variables est manipulé. D'autre part, la distribution d'un champ gaussien peut s'exprimer comme un produit de fonctions potentiel d'ordre 1 et 2 : les potentiels d'ordre 1 font intervenir le carré des variables ($\psi_i(h_i) = f_i(h_i^2)$) et celles d'ordre 2 le produit des variables ($\psi_{i,j}(h_i, h_j) = f_{i,j}(h_i h_j)$). Le MGP ainsi défini est associé à un graphe G complet puisque dans un champ gaussien chaque variable est a priori corrélée avec toutes les autres. Si l'on se place dans le cas particulier des champs de Markov gaussiens, la matrice de précision (l'inverse de la matrice de covariance) définissant des indépendances conditionnelles est creuse et le graphe associé également.

Ainsi, lorsque se pose la question du choix d'un modèle spatial pour représenter des données spatiale de type comptage sur quadrat, les champs de Markov et les processus de Cox log gaussiens sont deux modèles concurrents. Pour les comparer sur la base d'un critère de type BIC, il est alors nécessaire de disposer de méthodes d'estimation et de calcul de BIC approchées identiques pour les deux familles de modèles. Sinon les critères approchés calculés pour chacun des modèles ne seront pas comparables.

Jusqu'à récemment, l'approche classique pour mettre en œuvre une estimation bayésienne sur le processus de Cox gaussien était l'approche MCMC, avec les limitations classiques (temps de calculs pour atteindre la convergence, nombreux "boutons" à régler). En 2009, les auteurs de [102] ont proposé la méthode INLA (Inverse Nested Laplace Approximation) pour l'estimation dans les modèles gaussiens latents. Elle repose sur deux approximations emboîtées des probabilités a posteriori et conduit à des algorithmes bien plus rapides que par MCMC.

Nous disposons donc de trois approches pour l'estimation bayésienne du processus de

Cox log gaussien : MCMC, INLA et VBEM. Pour estimer un champ de Markov, INLA n'est pas disponible, et MCMC reste trop coûteux puisqu'il faudrait mettre en œuvre autant d'algorithmes que de modèles à comparer. Par contre, l'approche variationnelle est applicable et rapide. Il est donc envisageable de comparer les deux familles de modèles dans ce cadre. Pour cela, nous nous sommes tout d'abord intéressés au problème de la mise en œuvre de l'algorithme VBEM pour le processus de Cox gaussien ([99]). Comme mentionné dans la section 3.2.1, il est facile de calculer ensuite un critère variationnel de type BIC pour ce modèle, à partir des sorties de VBEM.

Les difficultés de mise en œuvre de VBEM sont dues au fait que même avec l'hypothèse d'indépendance entre variables cachées et paramètres ($q_{h,\theta} = q_h q_\theta$), il reste des problèmes calculatoires : *i*) la loi a posteriori de h n'est plus gaussienne et n'a pas de forme connue (nous ne sommes pas dans le cadre d'un modèle exponentiel conjugué), et *ii*) le calcul exact de $\Lambda(A_i) = \int_{A_i} \lambda_s ds$ est impossible.

La solution que nous proposons combine méthode variationnelle et simulation : nous nous restreignons aux distributions q_h de la forme $\prod_{i=1}^n q_{h_i}$, de sorte à se ramener à n distributions unidimensionnelles. Les étapes VBE et VBM sont alors respectivement des étapes de calcul de l'espérance de chaque H_i et de chaque paramètre sous les distributions q_h et q_θ . Nous approchons ces espérances par des estimations à partir de simulations de lois simples (gaussiennes ou uniformes). L'algorithme correspondant, même s'il fait appel à de la simulation, reste bien plus rapide qu'un algorithme MCMC pur (de l'ordre de 100 fois plus rapide). Il fournit une bonne estimation du paramètre β . Dans le cas d'une covariance exponentielle ($C(s_1, s_2) = \sigma^2 \exp(-\alpha \|s_1 - s_2\|)$), le paramètre de variance est bien estimé, mais ce n'est pas toujours le cas pour le paramètre d'échelle α (figure 3.7).

Ce travail a été initié lors du stage de M2R de Julia Radoszycki (2012). Il est encore dans une phase exploratoire et demande à être approfondi. Néanmoins les premiers résultats que nous avons obtenus sont encourageants. Afin d'améliorer la qualité des estimateurs nous envisageons de revenir à une étape VBE complètement variationnelle en approchant les lois a posteriori des données cachées par des distributions gaussiennes selon le principe de l'algorithme EM-EP ([78]). Cela devrait également avoir l'avantage d'accélérer le temps d'exécution de cette étape.

3.2.4 Applications des MGP à variables cachées en épidémiologie et en écologie

Le cadre des MGP avec variables cachées ou données manquantes trouve de nombreuses applications en épidémiologie animale et végétale car les données sont souvent très coûteuses à acquérir. Les ressources disponibles en temps et en argent limitent le nombre de relevés possibles ou ne permettent qu'une observation indirecte du processus d'intérêt. Par ailleurs la plupart des processus étudiés sont spatiaux et pour analyser cette structure spatiale il est souvent plus facile de raisonner sur une segmentation spatiale des observations initiales.

Cartographie de risque en épidémiologie animale

Principaux collaborateurs : David Abrial et Myriam Garrido (Unité d'épidémiologie animale - INRA SA - Clermont-Ferrand)

Publications associées : 14, 15.

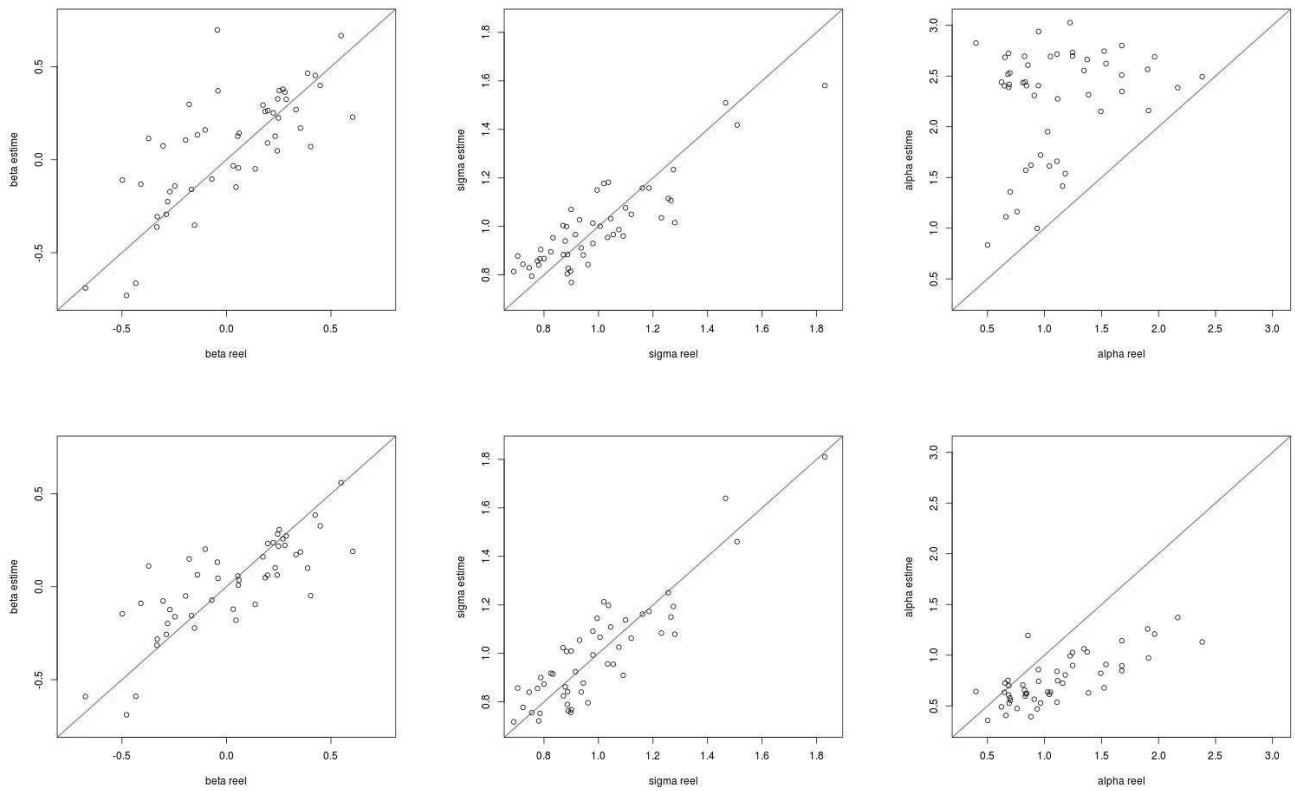


FIGURE 3.7 – Paramètres estimés en fonction du paramètre réel. Première ligne : résultats pour l’algorithme VBEM, deuxième ligne : résultats pour l’algorithme MCMC.

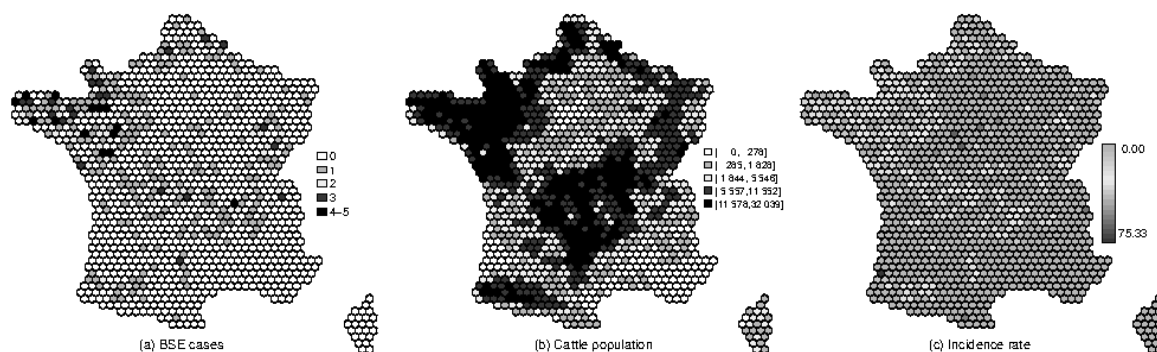


FIGURE 3.8 – Données de cas ESB en France : (a) nombre de cas sur la période d'étude, (b) taille de la population de bovins, (c) taux d'incidence standardisé. Images fournies par M. Charras-Garrido.

En épidémiologie animale, la cartographie du risque consiste à regrouper les entités administratives (échelle à laquelle les informations sont agrégées) dans des classes de risque (e.g. élevé, moyen, faible) à partir des nombres de cas observés (comptages) et de la situation géographique de ces entités. La classification obtenue délimite clairement les zones à risque, sur lesquelles des mesures de protection pourront être appliquées. Ce problème de segmentation de données spatiales se modélise facilement dans le cadre des champs de Markov cachés. Dans les approches classiques (pour la majorité inspirées du modèle proposé dans [8]) la variable cachée correspond au vecteur des risques individuels (valeur continue) en chaque entité. La segmentation est alors obtenue en deux temps : dans une première étape le risque individuel en chaque entité est estimé à partir des données de comptage, puis dans une seconde étape la segmentation à proprement parler est produite à partir de ces risques individuels estimés. Nous avons proposé une approche où les variables cachées représentent directement les classes de risque associées à chaque entité administrative ([19, 20]). Ainsi la segmentation est obtenue automatiquement à partir des données de comptage, sans autre estimation intermédiaire. Cette approche a été mise en œuvre pour la cartographie du risque d'Encephalopathie Spongiforme Bovine (ESB) dans les troupeaux bovins en France (voir figure 3.8). Les paramètres du modèle sont estimés par une version Monte Carlo de EM (MCEM). Dans un travail plus récent sur cette même application, les auteurs ont utilisé l'algorithme EM en champ moyen ([36]). Les classifications obtenues varient avec la forme du modèle de champs de Markov caché (voir figure 3.9). Cependant toutes confirment que l'approche de segmentation automatique proposée est bien adaptée pour localiser les régions à haut risque, celles justement ciblées en pratique pour l'application de mesures de contrôle, et estimer les classes de risque correspondantes.

Cartographie d'une espèce adventice dans une parcelle cultivée

Principaux collaborateurs : Mathieu Bonneau (UBIAT - INRA MIA - Toulouse), Sabrina Gaba (UMR Agroécologie - INRA EA - Dijon) et Régis Sabbadin (UBIAT - INRA MIA - Toulouse)

Publications associées : 19, 49, 50, 51.

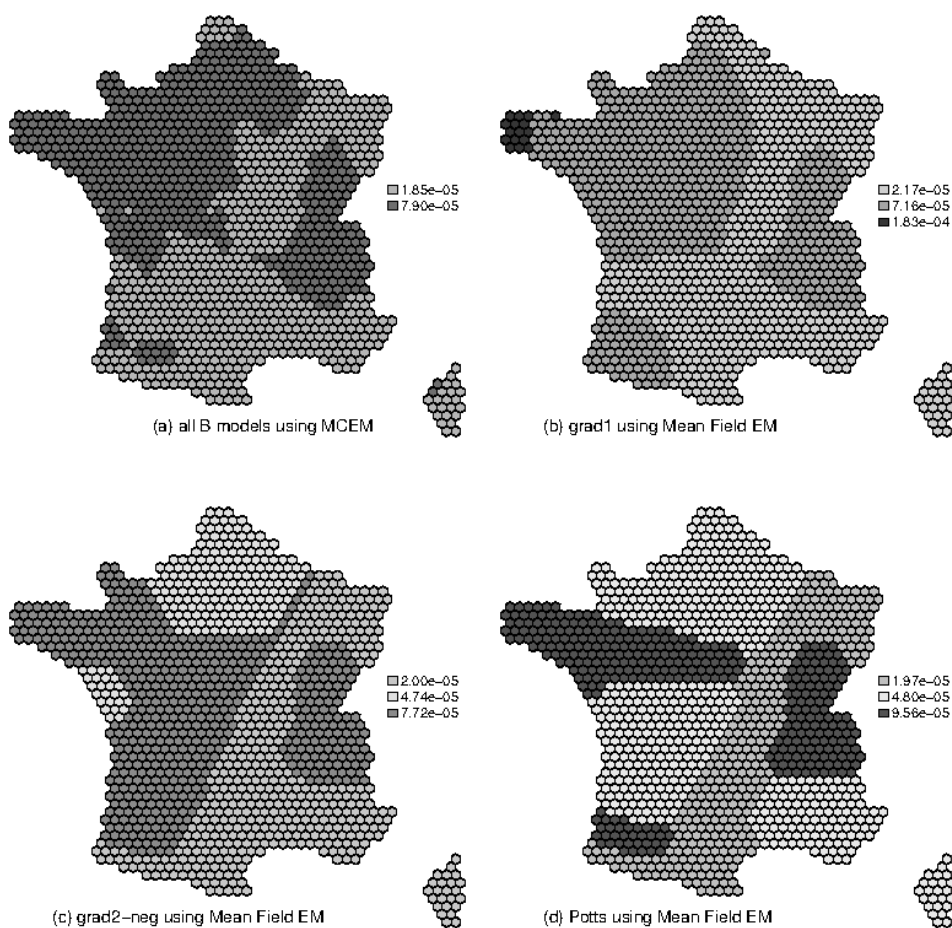


FIGURE 3.9 – Différentes classification du risque d'ESB en fonction du modèle de champs de Markov caché et de la méthode d'estimation utilisée : (a) restauration obtenue quelque soit le modèle avec estimation par MCEM, (b) et (c) restaurations obtenues pour deux modèles à variations spatiales douces et une estimation par EM en champ moyen, (d) restauration obtenue pour le modèle de Potts et une estimation par EM en champ moyen. Images fournies gracieusement par M. Charras-Garrido.

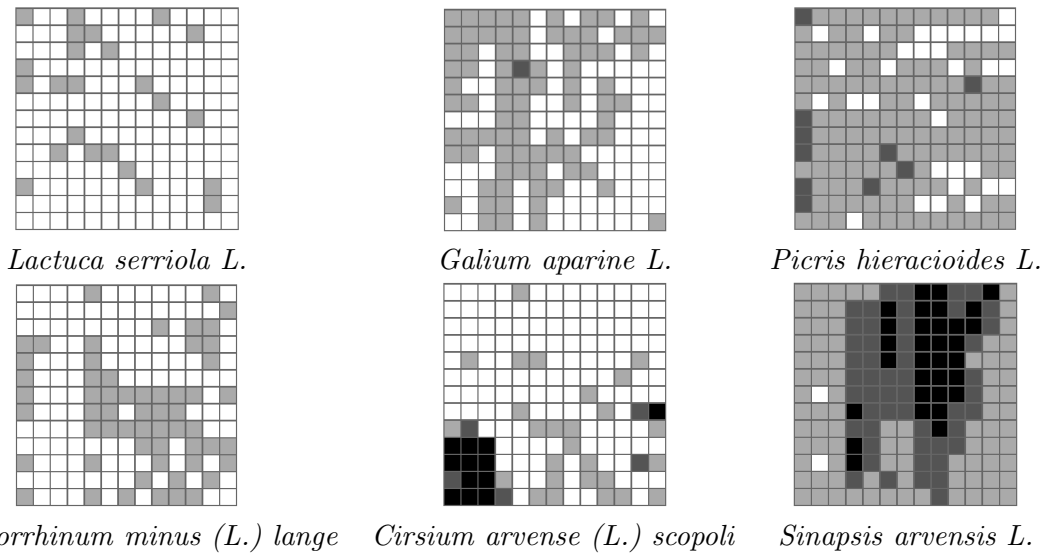


FIGURE 3.10 – Cartes de classes de densité de différentes espèces adventice, à l'échelle d'une parcelle.

Dans le cadre de la thèse de Mathieu Bonneau, nous avons développé un algorithme (LSDP) pour la conception de stratégies d'échantillonnage adaptatif et nous l'avons appliqué à l'échantillonnage des espèces adventices dans une parcelle (voir section 3.4). LSDP travaille à partir d'un modèle de distribution spatiale du phénomène étudié. Cette distribution peut être variable d'une espèce adventice à l'autre (voir figure 3.10). Aussi nous avons proposé différents modèles dans la famille des champs de Markov, correspondant à différentes hypothèses envisageables (isotropie ou non, variation spatiale douce ou abrupte, classes d'abondance a priori homogènes ou non). Dans le cas de données échantillonnées, les paramètres de ces modèles ont été estimés grâce à l'algorithme EM en champ moyen et nous avons utilisé l'approximation variationnelle correspondante de BIC. Cette analyse a montré que le modèle sélectionné par BIC dépend fortement de la taille des quadrats, de l'espèce adventice, de la culture dans la parcelle et de la période de l'année.

Estimation des traits de vie caractérisant la dynamique de la banque de graines adventices, à partir de données d'abondance sur la flore levée

Principaux collaborateurs : Benjamin Borgy (CEFE - Montpellier), Sabrina Gaba (UMR Agroécologie - INRA EA - Dijon), Régis Sabbadin (MIAT - INRA MIA - Toulouse)

Publications associées : 20, 51.

La dynamique des espèces adventices interagit fortement avec les pratiques culturales et l'on peut imaginer que certaines pratiques créent des conditions plus favorables à certaines espèces adventices qu'à d'autres. Pour tester cette hypothèse, l'information sur la flore levée (note d'abondance) telle qu'utilisée pour la cartographie à une date donnée ne suffit plus. Pour comprendre et caractériser les dynamiques temporelles des espèces adventices, il est crucial de disposer également d'informations sur la composition de la banque de graines.

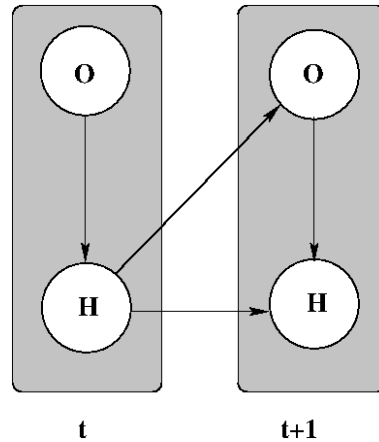


FIGURE 3.11 – Représentation graphique du réseau bayésien dynamique associé au modèle de dynamique des adventices. Les variables cachées H (resp. observées O) représentent les classes d’abondance dans la banque de graines (resp. dans la flore levée).

C’est une des composantes principales de la structure de la communauté. Les graines en dormance peuvent ainsi survivre pendant plusieurs années pour finalement germer lorsque les conditions sont favorables. Ignorer cette source de graines conduirait à un modèle incomplet. Cependant, très peu d’informations sont disponibles sur la composition de la banque de graine, car leur acquisition exige un travail lourd et fastidieux d’échantillonnage, comptage et identification. Nous avons donc développé un modèle de réseau bayésien dynamique composé de deux couches (figure 3.11) : les variables observées, c’est à dire les séries temporelles des notes d’abondance des espèces sur la flore levée, et les variables cachées, les séries temporelles des notes d’abondance des espèces dans la banque de graines. Ce modèle peut aussi être interprété comme un modèle de Markov caché. La dynamique de la banque de graines pouvant être décrite à partir de trois traits de vie (taux de survie, germination et reproduction), nous avons paramétré les lois conditionnelles du modèle avec ces trois traits. Les covariables sont les pratiques culturales. Ici l’estimateur du maximum de vraisemblance a été calculé par une méthode classique d’optimisation. Nous avons constaté que les estimateurs des traits de vie varient effectivement en fonction des pratiques culturales, suggérant des leviers de gestion pour le contrôle des adventices. La suite de ce travail consistera à étendre le modèle, pour l’instant purement temporel, à un modèle spatio-temporel, afin de pouvoir représenter des dynamiques à l’échelle de la mosaïque paysagère. Il est fort probable que pour estimer les paramètres d’un tel modèle, nous devrons utiliser des méthodes approchées et faire appel aux variantes de l’algorithme EM décrites dans la section 3.2.1.

3.3 Inférence dans les modèles graphiques

L’utilisation des méthodes variationnelles pour l’inférence approchée dans les MGP spatiaux ou structurés statiques est maintenant devenue assez classique, avec l’approximation en champ moyen, les algorithmes de passage de messages et leurs extensions. Dans le cas

de modèles dynamiques, c'est en épidémiologie et écologie théorique que l'on retrouve le plus l'usage de ces méthodes (sous le nom de méthodes de fermeture des moments, voir section 2.5), initialement pour des modèles SIS (Susceptible-Infecté-Susceptible) ou SIR (Susceptible-Infecté-Retiré) déterministes puis plus récemment pour les versions stochastiques de ces modèles (eg [35, 39]). Néanmoins la définition de méthodes variationnelles d'ordre supérieur à 1 pour approcher les dynamiques du système étudié reste un sujet de recherche car les plus simples ne conduisent pas toujours à des approximations satisfaisantes et les plus complexes ne peuvent pas toujours être mises en œuvre de manière efficace. Un compromis coût/qualité est à trouver. Ainsi, pour approcher le modèle SIS stochastique sur graphe, nous avons proposé une famille de méthodes variationnelles à l'ordre 2 prenant en compte des corrélations à longue distance et nous avons identifié dans cette famille des candidats qui réalisent ce compromis (sous-section 3.3.1). Ce travail a permis d'étudier le rôle de la structure des interactions dans la propagation d'une épidémie (sous-section 3.3.2). Cette approche pour l'inférence approchée d'un MGP spatio-temporel reste néanmoins trop coûteuse à mettre en œuvre lorsqu'il est nécessaire d'y faire appel de nombreuses fois comme c'est le cas pour la résolution de problèmes d'optimisation de décisions séquentielles en contexte spatial. Dans ce cas, la précision de la méthode d'inférence approchée n'est pas la qualité première attendue. Il nous faut disposer d'une méthode rapide et fournissant une représentation, peut-être grossière, mais suffisamment bonne pour conduire à une bonne approximation de la stratégie optimale. Aussi, nous nous sommes intéressés au calcul approché de l'étendue d'une épidémie dans un modèle SIR par une méthode variationnelle d'ordre 1 : tout en maintenant l'hypothèse d'indépendance, nous avons proposé une famille $\mathcal{Q}$ de distributions candidates plus large que celle associée au champ moyen classique, de sorte à améliorer la qualité de l'approximation tout en restant efficace en terme de temps de calcul (sous-section 3.3.3).

3.3.1 Méthodes variationnelles d'ordre 2 pour l'inférence des états d'équilibre et transients d'un modèle SIS sur graphe

Principaux collaborateurs : Ulf Dieckmann (EEP - IIASA - Autriche), Alain Franc (Bio-GeCo - INRA EFPA - Bordeaux)

Publications associées : 5, 9

Projet : EmerGeom.

Le modèle de processus de contact (voir section 2.5), également connu sous le nom de modèle SIS sur graphe, fait partie de la famille des systèmes de particules en interaction ([47, 74]). Ce modèle est caractérisé par l'existence de deux comportements possibles : soit le processus s'éteint très rapidement, soit au contraire il persiste. En fonction de la valeur du taux de contamination β , le processus adoptera un comportement ou l'autre. Il s'agit d'un phénomène de transition de phase. La valeur du paramètre pour laquelle la transition a lieu s'appelle la valeur critique. Pour étudier la propagation d'une épidémie, une bonne estimation de la valeur critique du modèle est importante car les mesures de contrôle sont prises de sorte à rester en dessous de cette valeur afin de contenir la maladie. L'approximation en champ moyen du processus de contact est connue pour son mauvais comportement près de la transition de phase ([30, 34] [28, Part B], [56]). La principale raison est que l'approxima-

tion en champ moyen ignore les corrélations entre individus. Des approximations à l'ordre 2 ont été proposées en épidémiologie et en écologie théorique pour intégrer les interactions entre voisins, avec parmi les plus classiques l'approximation par paire et l'approximation de Bethe (e.g. [28]). Cependant elles ne permettent pas toujours d'améliorer suffisamment l'estimation de la valeur critique. L'étude analytique de modèles "simples" de la physique statistique, tels que le modèle d'Ising, a montré que l'une des raisons de l'échec des approximations est la persistance, voire l'amplification des corrélations à longue distance lorsque le taux de contamination est proche de la valeur critique ([74, 106, 27]). Les corrélations à longue distance désignent les corrélations qui existent entre deux variables attachées à deux nœuds du graphe qui sont connectés par un chemin de longueur plus grande que 1. Nous avons donc proposé une extension des méthodes variationnelles à l'ordre 2 en intégrant les corrélations non seulement entre paires de variables voisines sur le réseau mais également entre paires de variables éloignées ([89]). Cela reste une méthode variationnelle à l'ordre 2 car les nouvelles approximations ainsi définies ne font intervenir que des marginales sur les paires et les singletons. Nous nous sommes basés sur un ensemble d'approximations qui ont été proposées dans la littérature en écologie théorique comme méthodes de fermeture des moments, et qui ne reposent que sur les corrélations entre voisins (i.e. à distance 1) Puis nous les avons étendues afin d'intégrer les corrélations jusqu'à n'importe quelle distance d . Dans le cas de l'approximation de Bethe, si i, j et k sont trois nœuds du graphe et si l'on note d_{ij} , d_{ki} et d_{jk} les trois distances entre chaque paire de nœuds, nous construisons la forme générale de l'approximation de $p(x_i, x_j, x_k)$ comme

$$p(x_i, x_j)^{d_{ij}} p(x_j, x_k)^{d_{jk}} p(x_k, x_i)^{d_{ki}} / p(x_i) p(x_j) p(x_k).$$

Pour obtenir une approximation ne prenant en compte que les corrélations que jusqu'à la distance d , les probabilités jointes de paires de variables associées à une distance supérieure à d sont remplacées par le produit de leur deux probabilités marginales.

Pour chaque approximation issue de la littérature, nous avons mis en œuvre cette généralisation pour des corrélations jusqu'à la distance 3, sur différents graphes non nécessairement réguliers mais à degré homogène (voir figure 3.12 pour une illustration sur un résultat type).

Nos conclusions sont les suivantes : *i*) pour $\beta \approx \beta_c$ avec $\beta < \beta_c$, les approximations avec corrélations uniquement à distance 1 prédisent mal les trajectoires transientes et l'introduction des corrélations à longue distance améliore la qualité des résultats, *ii*) pour $\beta \approx \beta_c$ avec $\beta > \beta_c$ toutes les approximations surestiment la prévalence dans les états transients et à l'équilibre, *iii*) pour $\beta \gg \beta_c$ toutes les approximations conduisent à une estimation correcte de la prévalence à l'équilibre mais sont moins performantes sur l'estimation des états transients. De plus, notre étude a permis d'exhiber deux approximations conduisant aux meilleures estimations à la fois à l'équilibre et dans les états transients : l'approximation "Power 1" à distance 1 et l'approximation de Bethe normalisée intégrant les corrélations jusqu'à la distance 3.

La limite principale de ce travail est que nous n'avons pu mettre en œuvre les approximations étendues que sur des graphes de degré homogène. Or, à part dans certains cas bien particuliers (épidémie dans un verger par exemple) le graphe d'interaction associé au processus étudié n'a pas cette propriété. Les réseaux sociaux, les réseaux aériens, ou encore les réseaux écologiques ([82, 88]) sont caractérisés par des distributions des degrés plus com-

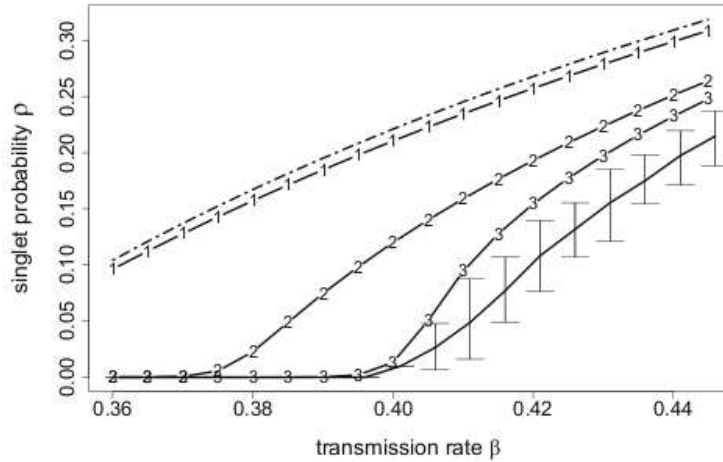


FIGURE 3.12 – Prédiction de l'état d'équilibre ρ (probabilité de l'état S), pour des valeurs du taux de contamination β proches du seuil critique, dans le cas d'une grille régulière à 4 voisins. La courbe continue et les barres d'erreur correspondent à la moyenne et l'intervalle de confiance à 90% obtenus par simulation, les 3 courbes associées aux symboles 1, 2 et 3 correspondent à la prédiction par l'approximation de Bethe normalisée avec corrélations à distances 1, 2 et 3, la courbe en pointillés correspond à la prédiction par une approximation de référence avec corrélation uniquement à distance 1.

plexes. D'un point de vue théorique, les approximations que nous avons étudiées ici peuvent être définies pour des graphes quelconques. Toutefois la mise en œuvre reste lourde car le système d'équations différentielles qui en résulte, même si de taille inférieure au système décrivant le processus exact, est trop grand. En nous limitant aux corrélations à distance 1, nous avons pu dériver et mettre en œuvre l'approximation par paire et l'approximation de Bethe du modèle SIS sur un graphe quelconque ([90]). De manière assez prévisible, nous avons obtenu une amélioration des résultats en champ moyen. Le défi à relever maintenant est de trouver des algorithmes efficaces permettant d'utiliser les approximations variationnelles intégrant les corrélations à longue distance sur ces mêmes graphes.

3.3.2 Rôle de la structure des interactions dans la propagation d'une épidémie

Principaux collaborateurs : Ulf Dieckmann (Evolution and Ecology Program - IIASA - Autriche), Alain Franc (BioGeCo - INRA EFPA - Bordeaux)
Publications associées : 9, 22, 23

Le modèle SIS stochastique sur graphe est très souvent utilisé en écologie pour modéliser des dynamiques de métapopulations ([39]) et en épidémiologie pour modéliser la propagation des épidémies ([34, 86, 31]). Il est reconnu que la structure du graphe représentant le réseau d'interaction entre les individus joue un rôle important dans ces dynamiques ([55, 31, 56, 22]). En particulier, la valeur critique associée à la transition de phase du modèle peut dépendre de caractéristiques géométriques du graphe, comme la distribution

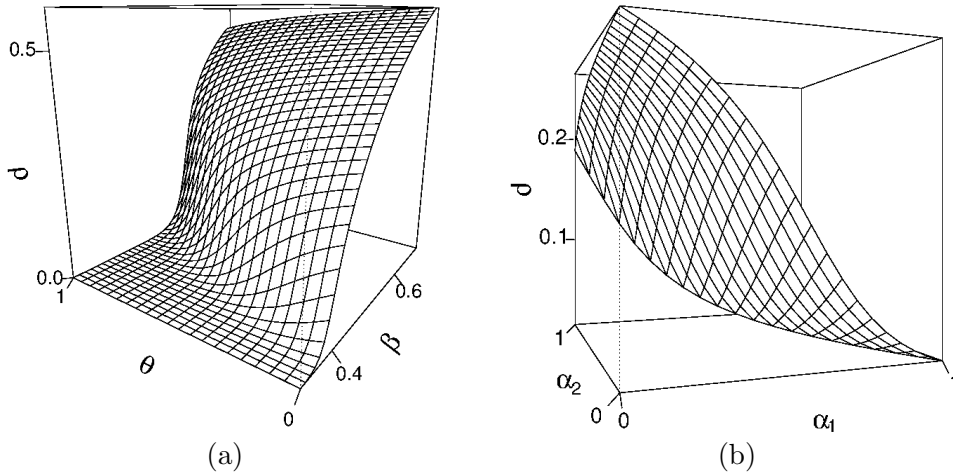


FIGURE 3.13 – Influence de la structure du graphe sur l'état d'équilibre ρ (probabilité de l'état S) pour un graphe quelconque de degré homogène égal à 4 : (a) influence du coefficient de clustering θ en fonction du taux de contamination β (b) influence des 2 coefficients de clustering à distance 2 (α_1 et α_2) dans le cas d'un taux de contamination $\beta = 0.47$.

des degrés ou le coefficient de clustering⁵ ([1, 22]) Cependant, dans [14] les auteurs soulignent que le coefficient de clustering n'est certainement pas suffisant pour résumer à lui seul l'influence du graphe. Aucune des méthodes classiques de fermeture des moments ne permet l'identification d'autres caractéristiques du graphe ayant un rôle important dans la dynamique du processus. L'intérêt des approximations intégrant les corrélations à longue distance proposées dans la section précédente est que les nouveaux paramètres du système d'équations différentielles résultant sont justement des caractéristiques du graphe. Il est donc possible de quantifier leur influence (voir figure 3.13). Nous retrouvons le coefficient de clustering dans ces paramètres, et nous identifions en plus de nouveaux coefficients associés aux corrélations à distance 2 et 3 (appelés coefficients de clustering à distance 2 et 3). Tous ces coefficients sont différentes mesures de la redondance des chemins entre les individus. Pour deux graphes possédant le même nombre d'arêtes, celui qui présente le moins de redondance dans les chemins répartit mieux les chemins entre les individus. Il sera donc plus favorable à la propagation de l'épidémie.

3.3.3 Méthode variationnelle d'ordre 1 pour estimer la taille de l'épidémie dans un modèle SIR sur graphe

Principaux collaborateurs : Régis Sabbadin (MIAT - INRA MIA - Toulouse)

Publication associée : 58.

Dans les travaux précédents le modèle stochastique utilisé pour modéliser des dynamiques épidémiques est le modèle SIS où les individus sont susceptibles ou infectés et re-

5. Le coefficient de clustering est la proportion de triplets i, j, k formant un triangle parmi tous les triplets où au moins deux des trois arêtes possibles sont présentes.

deviennent susceptibles après infection. Dans le cas de maladies qui conduisent au décès de l'individu (ESB, [104]) ou au contraire pour lesquelles l'individu acquiert une immunité permanente (maladies infantiles en épidémiologie humaine ([33])), le modèle SIR est plus adapté puisqu'il existe un troisième état possible, l'état absorbant "retiré". Ce modèle ne présente pas de transition de phase. Dans tous les cas le système converge vers un état où il n'y a plus d'infecté dans la population et où les individus sont soit susceptibles soit retirés. La quantité d'intérêt pour les épidémiologistes et les décideurs est la taille (en espérance sur toutes les trajectoires possibles) de l'épidémie (TE), c'est à dire le nombre moyen d'individus dans l'état retiré à la fin de l'épidémie. Les résultats sur un calcul exact sont rares. Pour un modèle SIR quelconque M. Newman, dans [83], a montré comment calculer la TE numériquement à partir d'outils issus de la théorie de la percolation. Le résultat est en moyenne sur l'ensemble (infini) des graphes aléatoires possédant la même distribution des degrés. Dans [80] l'auteur donne la distribution asymptotique de TE pour un modèle SIR sur un graphe de Bernouilli. Cependant, lorsqu'il s'agit de prendre des décisions afin de contrôler une épidémie dans une population, ces résultats théoriques ne s'appliquent plus et il faut savoir calculer la TE pour le graphe particulier associé à cette population. De plus, pour concevoir une stratégie de contrôle minimisant la TE de nombreux appels à l'algorithme de calcul de cette quantité sont nécessaires. Ce dernier doit donc être rapide.

Afin de satisfaire cette contrainte, nous avons proposé une méthode variationnelle à l'ordre 1 pour le calcul approché de la TE sur un graphe donné. Le modèle SIR considéré est à temps discret (chaque date correspondant à une date de prise de décision), il se définit comme suit. Soit $G = (V, E)$ le graphe des interactions dans la population. A chaque nœud i du graphe G est associée une variable aléatoire X_i à valeurs dans $\Omega_i = \{0, 1, 2\}$. Les trois états représentent respectivement un individu susceptible, infecté ou retiré. Si ρ_{ji} est la probabilité que la maladie soit transmise du nœud j vers un nœud i voisin (i.e. $i \in N(j) = \{k \in V \text{ tq } (k, j) \in E\}$), alors pour une configuration donnée $x_{N(i)}^t$ de $X_{N(i)}^t$ la probabilité d'être infecté conditionnellement à l'état du voisinage est

$$p_i(X_i^{t+1} = 1 | X_i^t = 0, X_{N(i)}^t = x_{N(i)}^t) = 1 - \prod_{j \in N(i), x_j^t = 1} (1 - \rho_{ji}).$$

Nous nous sommes placés dans le cas simple où toutes les autres transitions définissant la chaîne de Markov sont déterministes :

$$p_i(X_i^{t+1} = 2 | X_i^t = 1) = 1, \quad p_i(X_i^{t+1} = 2 | X_i^t = 2) = 1.$$

Nous avons donc fait l'hypothèse que l'unité de temps considérée est au moins égale à la durée d'infection d'un individu. La quantité TE est définie par :

$$TE = E \left[\sum_{i \in V} \mathbb{1}_{[X_i^\infty = 2]} \right].$$

L'approximation en champ moyen pour le calcul de la TE consiste à approcher la distribution spatio-temporelle des n individus par n chaînes de Markov indépendantes. Dans sa mise en œuvre classique il est imposé soit que ces n chaînes soient identiquement distribuées, soit que les distributions temporelles des individus de même degré soient identiques. Dans les

deux cas, ce sont des hypothèses fortes. Afin d'améliorer l'approximation, nous avons défini une approximation en champ moyen dans laquelle chaque individu possède une dynamique différente. De plus nous autorisons la non-stationnarité, afin de compenser la moyennisation sur la dimension spatiale due à l'hypothèse d'indépendance. Ainsi, un élément q de la famille $\mathcal{Q}$, parmi laquelle la meilleure approximation du modèle SIR est recherchée, est défini par n chaînes de Markov indépendantes conditionnellement à l'état de départ :

$$\forall \{x^0, x^1, \dots, x^T\} \in \Omega^{n \times T},$$

$$q(x^0, x^1, \dots, x^T) = q^0(x^0) \prod_{t=1}^T \prod_{i=1}^n q_i^t(x_i^t | x_i^{t-1}).$$

La seule inconnue est $q_i^t(X_i^t = 1 | X_i^{t-1} = 0)$, les autres probabilités de transitions se déduisent de celle-ci ou sont triviales (égales à 0 ou 1). Nous définissons alors l'approximation en champ moyen du modèle SIR comme la distribution q^* dans $\mathcal{Q}$ qui minimise la divergence de Kullback-Leibler entre q et p : $q^* = \arg \min_{q \in \mathcal{Q}} KL(q||p)$, avec

$$KL(q||p) = \sum_{x^0, \dots, x^T} q(x^0, \dots, x^T) \log \frac{q(x^0, \dots, x^T)}{p(x^0, \dots, x^T)}.$$

La résolution de ce problème global est trop complexe et nous avons défini un algorithme itératif qui calcule une approximation en champ moyen locale à chaque pas de temps du processus.

Des comparaisons entre cette solution champ moyen avec d'une part la solution exacte (pour n petit) et d'autre part une estimation par Monte Carlo (pour n grand) ont montré un bon comportement de celle-ci. Néanmoins, comme pour toute approximation en champ moyen, nous avons observé une surestimation de la taille de l'épidémie pour des valeurs basses du taux de contamination. Ce comportement pourrait être amélioré en considérant une méthode variationnelle d'ordre 2 parmi celles présentées précédemment en section 3.3.1, au prix d'un temps de calcul rallongé.

3.4 Décision dans les modèles graphiques

Les travaux que je présente maintenant ne sont plus motivés par l'acquisition de connaissances sur un phénomène ou par sa prédiction. Ils traitent de la gestion et du contrôle de ce système. Dans les agro-écosystèmes, les problèmes de gestion concernent par exemple le contrôle d'une épidémie qui se propage sur une région agricole ou encore le contrôle de l'invasion d'un territoire par une espèce non native. Il peut s'agir aussi de gérer les services (pollinisation, habitat, ...) rendus par un paysage agricole. Dans ces exemples les corrélations entre les composantes du système sont spatiales. Les relations entre les espèces peuvent également être gouvernées par des relations trophiques (qui mange qui). Le système n'est pas spatialisé mais il reste structuré. Tous les acteurs sont convaincus que la prise en compte des interactions existant entre les entités du système pour définir une stratégie de gestion est importante : une stratégie globale (collective) est plus efficace que la somme de stratégies individuelles. Une autre spécificité de ces problèmes de gestion est que l'on souhaite gérer le système non pas à un instant donné mais à moyen ou long terme. Des

décisions sont prises chaque année pour atteindre cet objectif. Nous sommes donc dans le cadre de la décision séquentielle dans l'incertain.

Afin de concevoir des stratégies de gestion des agro-écosystèmes les partenaires agronomes et écologues disposent parfois d'outils pour comparer par simulation-évaluation des stratégies pré-définies. La comparaison est faite sur la base d'un critère à définir (économique ou non). Une autre approche consiste à calculer la stratégie qui optimise ce critère. Cette approche reste moins développée en agro-écologie (même si des travaux existent, voir [58]), surtout lorsque le problème est spatial ou structuré car les outils méthodologiques disponibles ne permettent pas de traiter des problèmes de cette taille. On constate un réel besoin de méthodes pour la conception par optimisation de stratégies de gestion dans le temps des agro-écosystèmes.

Depuis 2008, Régis Sabbadin (chercheur en intelligence artificielle) et moi-même avons développé le thème "Décision Spatialisée" au sein de l'équipe MAD, afin d'une part de proposer de nouveaux cadres et des algorithmes pour la conception de stratégies de gestion par optimisation en contexte spatial, ou plus généralement structuré, et d'autre part de modéliser et résoudre avec nos partenaires agronomes et écologues leurs problèmes de décision. Ce thème se positionne à l'interface entre intelligence artificielle et statistique et nous combinons ces deux domaines pour proposer des méthodes originales. Mes contributions au thème sont l'apport des méthodes variationnelles comme nouvel outil pour la résolution approchée des problèmes de décisions, ainsi que la modélisation et la simulation de systèmes sur réseau d'interaction.

Le cadre des Processus Décisionnels de Markov (PDM, [97]) est bien adapté à la décision séquentielle et à la conception par optimisation. Il est au cœur des travaux du thème "Décision Spatialisée". Un PDM est défini par

- un ensemble d'états, $\mathcal{X}$, qui décrit tous les états possibles du système,
- un ensemble d'actions, $\mathcal{A}$, qui décrit toutes les actions disponibles à chaque instant pour agir sur le système,
- une probabilité de transition qui décrit l'évolution du système conditionnellement à l'action choisie,
- une fonction de récompense, qui quantifie le gain associé à chaque paire état-action du système.

Dans ce cadre, une stratégie est une fonction qui associe une action à chaque état du système, c'est une règle de décision, stationnaire ou non. La (fonction de) valeur d'une stratégie est l'espérance de la somme des récompenses acquises au court du temps si l'on suit cette stratégie à partir d'un état initial donné. Enfin, résoudre un PDM consiste à déterminer la stratégie optimale, c'est à dire celle de plus grande valeur.

Dans les problèmes spatialisés ou structurés décrits ci-dessus, l'espace d'états et l'espace d'actions sont factorisés : si le réseau d'interaction est composé de n nœuds, $\mathcal{X} = \prod_{i=1}^n \mathcal{X}_i$ et $\mathcal{A} = \prod_{i=1}^n \mathcal{A}_i$. En chaque nœud i , une variable est observée (état sanitaire, occurrence) et une action est appliquée (éradication, protection). Pour ces PDM particuliers, la résolution exacte n'est pas possible et les méthodes approchées existantes atteignent leur limite du fait de la taille des problèmes. Il a donc fallu en proposer de nouvelles, de complexité temporelle et spatiale réduites. La première contribution (sous-section 3.4.1) que je présente porte sur un algorithme en champ moyen pour la résolution d'une famille particulière de PDM pour problèmes structurés, celle des Processus Décisionnels de Markov sur Graphe (PDMG).

La seconde concerne un autre cas particulier de problème de décision : l'échantillonnage adaptatif pour la cartographie (sous-section 3.4.3). Dans ce cas les actions sont uniquement cognitives, elles ne modifient pas l'état du système. Des applications de ces travaux méthodologiques en épidémiologie et en écologie sont présentées dans les sous-sections 3.4.2 et 3.4.4.

3.4.1 Un cadre et un algorithme de résolution approchée pour les Processus Décisionnels de Markov en contexte structuré

Principal collaborateur : Régis Sabbadin (MBIAT - INRA MIA - Toulouse)

Publications associées : 13, 29, 31, 43.

Projet : LARDONS

Lorsque l'espace d'état est factorisé, la probabilité de transition du PDM est une probabilité jointe sur les n variables du système. Il serait irréaliste de la stocker et la manipuler comme une grosse fonction de $\mathcal{X} \times \mathcal{X} \times \mathcal{A}$ vers $\mathbb{R}$. Le cadre des PDM factorisés (PDMF, [12, 46]), sous-famille des PDM, était jusqu'à présent le seul cadre disponible pour traiter ce type de problèmes. Dans ce cadre, la probabilité de transition, pour une action donnée, est celle d'un réseau bayésien dynamique quelconque. Des hypothèses d'indépendance conditionnelle sont faites. Dans un PDMF, la fonction de récompense globale est également décomposée en une somme de récompenses locales. La limitation majeure des PDMF réside dans la modélisation de l'espace d'actions $\mathcal{A}$: il reste global. Or, dans beaucoup de problèmes de décision structurés, la décision globale se décompose en un ensemble de décisions locales (une par nœud du réseau le plus souvent). Dans ce cas la taille de l'espace d'action est exponentielle en la taille du réseau et l'espace de stockage et le temps de calcul d'une fonction de valeur le sont aussi. Les algorithmes pour la résolution d'un PDMF, soit exacte comme avec SPUD (Stochastic Programming Using Decision Diagrams [48]) soit approchée comme avec APPRICOD [107]), ne sont alors plus adaptés. Le cas des PDMF avec espace d'actions factorisé a été proposé par [62], mais les algorithmes de résolution disponibles sont spécifiques du cas de variables binaires.

Nous avons considéré un cadre alternatif aux PDMF, le cadre des Processus Décisionnels de Markov sur Graphe (PDMG), dans lequel l'espace d'états et l'espace d'actions sont multidimensionnels de même dimension : $\mathcal{X} = \{\mathcal{X}_1, \dots, \mathcal{X}_n\}$ et $\mathcal{A} = \{\mathcal{A}_1, \dots, \mathcal{A}_n\}$. Dans ce modèle une variable d'état et une action par nœud du graphe sont définies. Comme dans un PDMF la probabilité de transition, pour une action donnée, est celle d'un réseau bayésien dynamique. Par contre il a une structure particulière. La dynamique de la variable d'état au nœud X_i ne dépend que de l'état au pas de temps précédent des variables voisines dans le graphe, et de l'action effectuée au nœud i :

$$p(x^{t+1}|x^t, a^t) = \prod_{i=1}^n p_i(x_i^{t+1}|x_{N(i)}^t, a_i^t), \forall x \in \mathcal{X}, \forall x' \in \mathcal{X}, a \in \mathcal{A}.$$

Cette forme de transition inclut en particulier les modèles SIR et SIS sur graphe, décrits dans la section 3.3. La récompense globale se décompose en une somme de récompenses par nœud et chaque récompense individuelle en i ne dépend que de l'état courant des variables

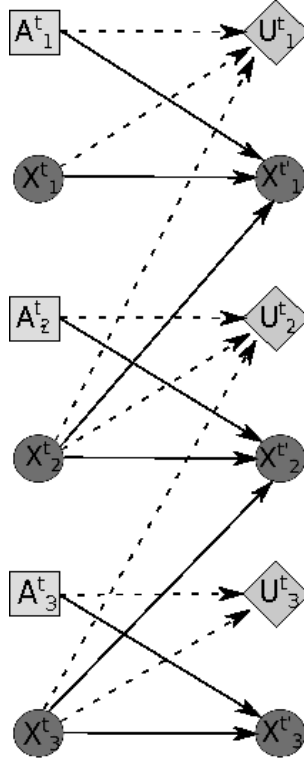


FIGURE 3.14 – Représentation graphique d'un PDMG. Les variables aléatoires sont représentées par des ronds, les actions par des carrés et les récompenses par des losanges. Les flèches pleines représentent les dépendances stochastiques associées aux probabilités de transition locales, et les flèches en pointillés représentent les dépendances déterministes associées aux récompenses. A un pas de temps donné, l'action globale (resp. la récompense globale) se décompose en une action (resp. une récompense) par variable aléatoire.

d'état voisines de i , et de l'action effectuée au nœud i .

$$r(x, a) = \sum_{i=1}^n r_i(x_{N(i)}, a_i), \forall x \in \mathcal{X}, \forall a \in \mathcal{A}.$$

La figure 3.14 donne une représentation graphique d'un PDMG.

Le cadre PDMG est donc plus restrictif que le cadre PDMF dans la forme des modèles de transition et de récompense. L'avantage principal réside dans la prise en compte de la factorisation de l'action globale, qui peut être exploitée pour la construction d'algorithmes de résolution approchée plus rapides.

Ainsi, nous avons proposé un algorithme de résolution approchée d'un PDMG combinant approximation en champ moyen du modèle de dynamique et recherche dans un sous-ensemble restreint de l'ensemble des stratégies ([92]). L'approximation en champ moyen est similaire à celle proposée en 3.3.3 dans le cas particulier du modèle SIR. Le résultat de cette approximation est une approximation de la fonction de valeur (fonction de $\mathcal{X}$ vers $\mathbb{R}$) par une somme de fonctions (une par nœud) ne dépendant chacune que de quelques actions

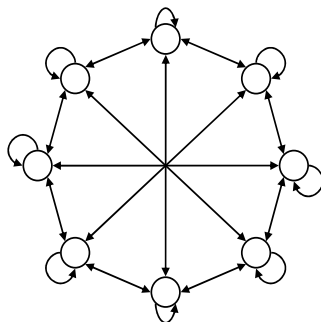


FIGURE 3.15 – Problème de contrôle d’une épidémie dans un parcellaire : structure du graphe utilisé pour tester les performances des algorithmes IPA-CM et PLA.

locales. Ensuite, l’optimum de cette fonction de valeur approchée est recherché uniquement parmi les stratégies locales, pour lesquelles la règle de décision au nœud i ne dépend que des variables d’état en i et dans son voisinage. L’algorithme, appelé IPA-CM est de type Itération de la Politique Approchée (IPA [5]) : il itère une étape d’évaluation de la valeur de la stratégie courante et une étape d’amélioration. Chaque itération est de complexité temporelle quadratique en n et exponentielle en la taille du plus grand voisinage $N(i)$.

Nous avons comparé ([103]) les performances de l’algorithme IPA-CM avec la seule solution algorithmiquement possible pour des problèmes structurés de grande taille : l’algorithme PLA ([38]), basé sur la Programmation Linéaire Approchée. Ces deux procédures ont été évaluées sur des problèmes inspirés de problèmes de décision en épidémiologie végétale et en gestion forestière, et pour des tailles de graphes allant jusqu’à $n = 700$ nœuds. Même si IPA-CM est quadratique et PLA linéaire en n , ce coût en temps de calcul est compensé par une bien meilleure approximation de la fonction de valeur d’une stratégie (figures 3.15 et 3.16). D’autre part si dans de nombreuses situations les qualités des stratégies retournées par IPA-CM et PLA sont comparables, dans des situations où l’influence du voisinage est forte le premier se démarque (figure 3.17).

3.4.2 Modélisation de problèmes de gestion spatialisés en épidémiologie et en écologie

Les hypothèses du cadre PDMG sont suffisamment réalistes pour beaucoup de problèmes de gestion en épidémiologie et en écologie. En effet, dans ces problèmes il y a en général autant de variables d’états que d’actions à réaliser et elles sont naturellement associées aux nœuds du réseau d’interaction : pour le contrôle de la propagation d’une épidémie sur une mosaïque paysagère, l’état sanitaire de chaque parcelle est suivi, et les agriculteurs appliquent un itinéraire technique (culture, type de labour, ...) potentiellement différent dans chaque parcelle. Pour le contrôle d’espèces invasives sur un territoire formé de patches en réseaux, la présence des espèces est surveillée en chaque patch et en chaque patch la décision est prise d’appliquer ou non une mesure d’éradication. La représentation des dynamiques par un produit de probabilités de transitions locales a également tout son sens pour des épidémies ou des espèces qui se propagent “par contact” aux plus proches voisins (par opposition aux dispersions à grande distance). L’hypothèse de récompenses locales associées

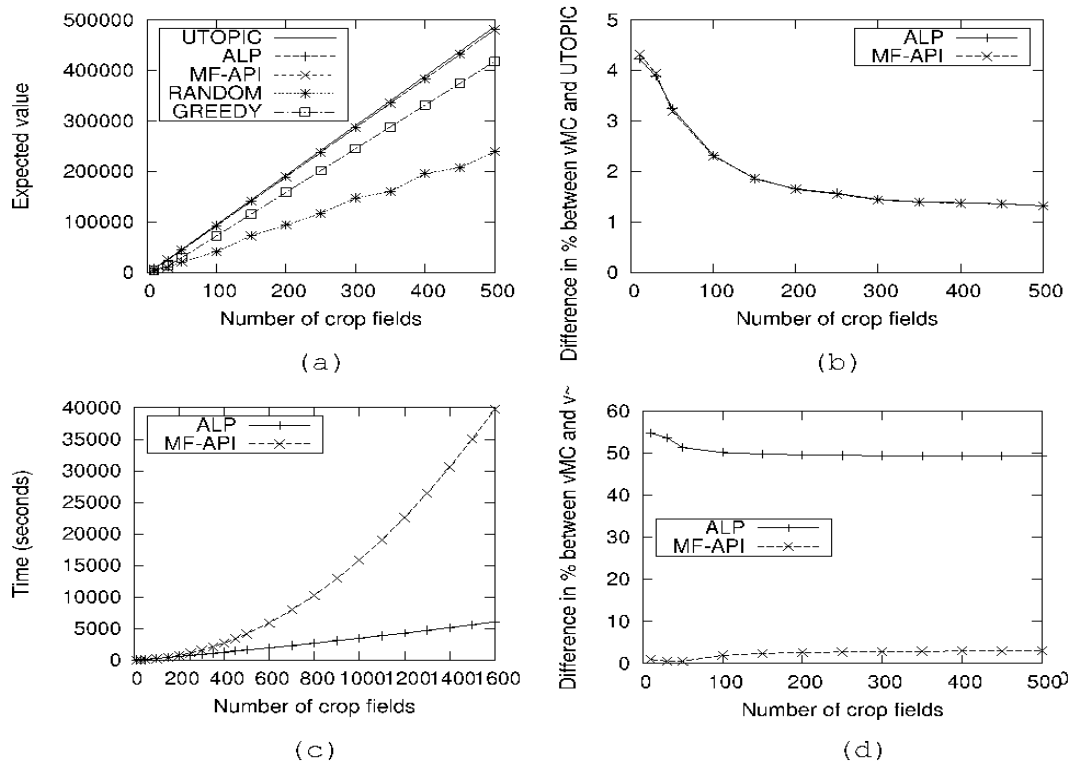


FIGURE 3.16 – Problème de contrôle d’une épidémie dans un parcellaire : (a) valeur moyenne de différentes stratégies et borne supérieure (utopic) de la valeur de la stratégie optimale, (b) différence (en pourcentage) entre les valeurs, estimées par Monte Carlo, des stratégies IAP-CM et PLA et la borne supérieure (utopic), (c) temps de calcul, (d) différence (en pourcentage) entre l’estimation par Monte-Carlo et l’estimation fournie par IAP-CM (resp. PLA) de la valeur de la stratégie IAP-CM (resp. PLA).

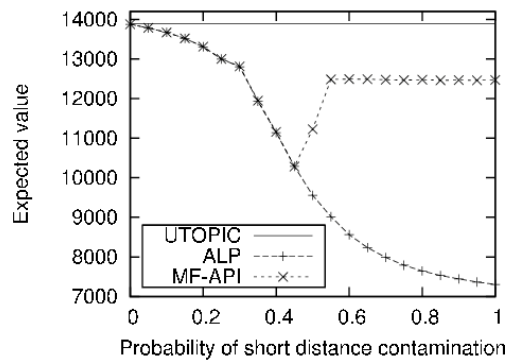


FIGURE 3.17 – Problème de contrôle d’une épidémie dans un parcellaire : influence de la probabilité de contamination à courte distance sur la qualité des stratégies IPA-CM et PLA.

à chaque nœud est aussi réaliste dans de nombreux cas. En épidémiologie végétale, si le critère est économique la récompense globale sera par exemple la somme des coûts des traitements (négatifs) et des revenus issus de la récolte en chaque parcelle. Ainsi comme signalé précédemment, l'algorithme IPA-CM a été mis en œuvre sur deux problèmes de gestion spatialisés : un problème de contrôle d'une épidémie dans une parcelle via des actions de type traitement ou jachère, ainsi que sur un problème de stratégie de coupe de parcelles de forêt face à un risque de tempête. Dans les deux cas il a conduit à des stratégies de bonne qualité. Néanmoins le cadre PDMG a ses limites. Je présente ici trois travaux de modélisation de problèmes de gestion spatialisés ou structurés, un en épidémiologie, deux en écologie, pour lesquels soit le modèle obtenu ne rentre pas dans le cadre, soit l'algorithme IPA-CM s'avère peu performant. Ces constats motivent de nouveaux développements méthodologiques, qui seront discutés dans le chapitre 4 de ce mémoire.

Conception de stratégies collectives pour la gestion du phoma du colza

Principaux collaborateurs : Jean-Noël Aubertot (AGIR - INRA EA - Toulouse) et Régis Sabbadin (MIAT - INRA MIA - Toulouse)

Publications associées : 18, 30, 32

Projet : ADD CEDRE et projet Européen PURE (Pest Use Reduction in Europe).

La conception de stratégies de gestion collective des maladies des cultures est difficile du fait de l'échelle et de la complexité du phénomène de développement et de dispersion des populations de pathogènes. Par certains aspects, le cadre des PDMG est bien adapté pour la modélisation d'un problème de gestion collective : prise en compte des interactions spatiales, décisions spatialisées, possibilités de proposer des modes de gestion intégrée. Cependant, sur d'autres aspects, la modélisation dans le cadre PDMG (à espace d'état discret) demande de faire des approximations par rapports aux modèles de dispersion mécanistes classiques en agronomie : les échelles de temps et d'espace de la dynamique doivent être les mêmes que celle du processus de décision (donc l'année et la parcelle), l'espace d'états de chaque variable locale, $\mathcal{X}_i$ doit être discrétisé en un petit nombre de classes. De plus, pour des pathogènes se dispersant à longue distance, la taille du voisinage dans le réseau d'interaction peut rapidement être trop élevée pour que l'algorithme IPA-MC reste performant en temps. Nous avons illustré ces difficultés de modélisation sur un exemple inspiré d'un problème de contrôle d'une épidémie de phoma du colza [93]. Une des actions possibles sur une parcelle consiste à planter une variété résistante au phoma, le risque étant que cette résistance soit rapidement contournée si trop exploitée. L'objectif est de choisir la stratégie qui optimise le revenu des agriculteurs sur la mosaïque paysagère, cette stratégie devant indirectement optimiser la durée de la résistance. La difficulté majeure a résidé dans la définition et le calcul des probabilités de transition du PDMG à partir du simulateur SIPPOM ([70]) qui est à échelles spatiale et temporelle continues. Les stratégies fournies par IPA-CM à partir de ce modèle ne sont pas meilleures que des stratégies ad-hoc, construites sur le "bons sens". Il est possible que le passage du simulateur vers le modèle de probabilité de transition ait trop dégradé l'information spatiale, ou que les choix d'actions soit trop limités, et que l'on ne soit donc pas dans un domaine où le cadre PDMG est intéressant. Néanmoins il serait intéressant de pouvoir élargir l'espace de recherche au delà des stratégies locales afin d'espérer améliorer les stratégies ad-hoc.

Conception de stratégies de contrôle de deux espèces invasives en compétition

Principaux collaborateurs : Alana Moore (UMR EcoFoG - Guyane), Mickael Mac Carthy (Univ of Melbourne, Australian Centre of Excellence for Risk Analysis, Australia) et Régis Sabbadin (MIAT - INRA MIA - Toulouse).

Les problèmes de gestion d'espèces invasives impliquent parfois plus d'une espèce. Les espèces sont alors en concurrence du fait de la compétition pour la colonisation de nouveaux patchs. Certaines sont meilleures compétitrices que les autres et l'homme n'a pas forcément les moyens de contrôler chacune des espèces. C'est par exemple le cas pour les chats et les chiens sauvages (dingos) en Australie : les dingos dominent les chats et sont également plus faciles à contrôler. Afin de proposer pour ces problèmes des règles de décision prenant en compte la structure du réseau de patchs, nous avons considéré un problème jouet dans lequel l'objectif est de contrôler la propagation de deux espèces invasives, A et B dans les conditions suivantes : l'espèce A domine l'espèce B (B ne peut pas s'installer sur un patch colonisé par A, B disparaît d'un patch si A s'y installe), A est la seule espèce dont on peut contrôler la propagation.

Nous avons réalisé la première étape de modélisation de ce problème sous la forme d'un PDM à espaces d'états et d'actions structurés. Le graphe associé au modèle est celui correspondant au réseau des patchs colonisables par les deux espèces. En chaque nœud, la variable suivie est dans un des état suivants : patch vide, patch occupé par A, patch occupé par B. Les actions possibles sur un patch sont soit l'éradication de A soit "ne rien faire". La probabilité de transition jointe a été définie comme un produit de probabilités de transition locales. Cependant, elles ne sont pas de type PDMG puisque l'on constate une dépendance aux actions sur toutes les parcelles voisines et non uniquement sur la parcelle d'intérêt :

$$p(x'|x, a) = \prod_{i=1}^n p_i(x'_i | x_{N(i)}, a_{N(i)}), \forall x \in \mathcal{X}, \forall x' \in \mathcal{X}, a \in \mathcal{A}.$$

Par contre, la définition de la récompense est bien de type PDMG : nous considérons la somme sur chaque patch du coût d'éradication si elle a eu lieu, plus le coût de présence de l'espèce A ou B le cas échéant.

Nous pouvons d'ores et déjà comparer des stratégies ad-hoc par simulation-évaluation grâce à ce modèle de dynamique. Par contre, ni le cadre PDMF ni le cadre PDMG ne fournissent des algorithmes de résolution approchée approprié. Le premier car il n'exploite pas l'aspect multidimensionnel de l'espace d'action, le second car les transitions locales d'un PDMG ne dépendent que de l'action locale au nœud i .

Optimisation dynamique de la gestion d'un ensemble d'espèces reliées par des interactions trophiques

Principaux collaborateurs : Eve MacDonald-Madden (CSIRO, Univ. of Queensland), Will Probert (CSIRO, Univ. of Queensland) et Régis Sabbadin (MIAT - INRA MIA - Toulouse)
Publications associées : 21, 52.

Projet : MANAE, DyBRES

Un des objectifs en biologie de la conservation est de préserver des espèces à risque. Les budgets alloués à ce type d’actions étant limités, il est important de répartir au mieux les efforts de conservation. Par ailleurs il n’est pas efficace de tenter de gérer séparément chaque espèce puisqu’elles sont reliées par des relations trophiques complexes de type prédateur-proie ([114]). Les espèces peuvent être représentées comme les nœuds d’un réseau trophique (voir figure 3.18 pour un exemple). Comment concevoir une stratégie de gestion d’un groupe d’espèces sur la base des interactions décrites par un réseau trophique ? Cette question est cruciale en biologie de la conservation mais n’avait pas été traitée jusqu’à présent, faute d’outils méthodologiques adaptés. Récemment les auteurs de [76] ont étudié le cas d’un réseau trophique statique. Dans ces travaux, les interactions entre les espèces sont modélisées par un réseau bayésien et différentes stratégies de choix des espèces à protéger sont comparées grâce à ce modèle. Ces stratégies sont basées sur différentes métriques (e.g. [52, 113]) permettant de hiérarchiser les nœuds du réseau, comme par exemple le degré ou la centralité d’un nœud. Ces travaux demandent maintenant à être étendus au cas plus réaliste de réseaux trophiques dynamiques. Pour cela nous avons récemment entrepris de modéliser le problème de décision comme un réseau bayésien dynamique contrôlé et nous comparons les stratégies basées sur les métriques par simulation-évaluation. L’objectif suivant est de proposer des stratégies par optimisation. La difficulté ici par rapport au cadre PDMG réside dans la prise en compte de contraintes globales sur les actions du type “seules k espèces parmi n peuvent être préservées à chaque pas de temps”.

3.4.3 Résolution approchée du problème d’échantillonnage adaptatif dans les champs de Markov

Principaux collaborateurs : Mathieu Bonneau et Régis Sabbadin (MIAT - INRA MIA - Toulouse)

Publications associées : 12, 17, 33, 34, 35, 36, 37, 44, 45, 46.

Projet : LARDONS.

Considérons maintenant un problème de décision un peu particulier où les actions ne modifient pas l’état du système mais ont pour but d’acquérir de l’information : celui du choix de sites d’échantillonnage pour la reconstruction de cartes spatiales. La motivation pour l’étude de cette question se trouve encore une fois en écologie et en épidémiologie : une stratégie de gestion sera d’autant plus efficace que le décideur dispose d’une bonne connaissance des zones de présence de l’espèce ou de la maladie. Cette information peut se représenter sous la forme d’une carte, soit d’occurrence, soit de classe d’abondance (un comptage exact est souvent difficile à obtenir dans ces domaines). Cependant, en général toute la zone d’intérêt ne peut pas être explorée du fait du coût prohibitif de l’observation, et la carte doit être reconstruite/estimée à partir d’un échantillon [24]. Une autre difficulté est la possible imperfection des observations qui devra être prise en compte dans la reconstruction (espèces dont la présence est difficile à observer, espèces pouvant être confondues avec d’autres, ...). La question est alors : où allouer l’effort d’échantillonnage afin d’obtenir les observations les plus informatives, tout en respectant le budget ? L’échantillonnage choisi peut être statique ou adaptatif. Dans le premier cas, l’ensemble des sites à échantillonner est

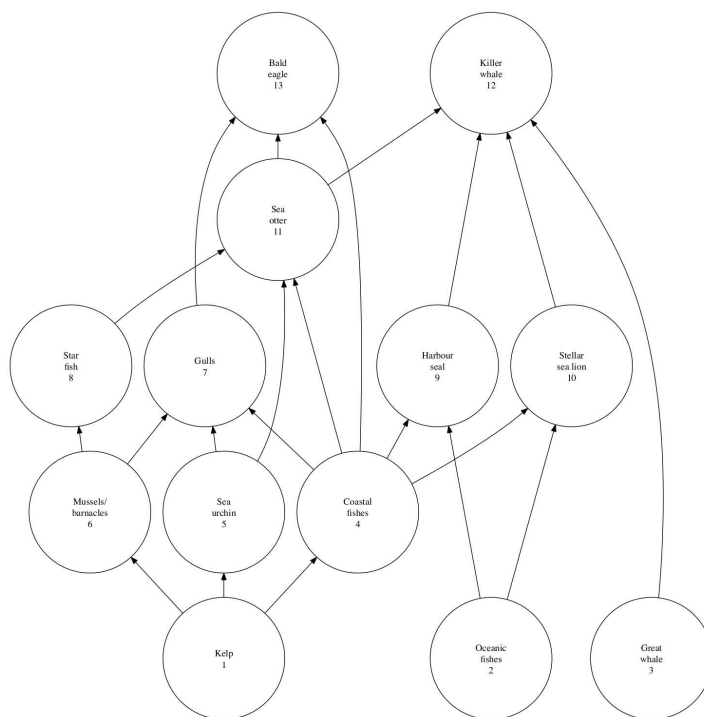


FIGURE 3.18 – Exemple d’un réseau trophique, sous-partie d’un réseau plus large représentant des interactions trophiques d’espèces en Alaska. Image fournie par W. Probert.

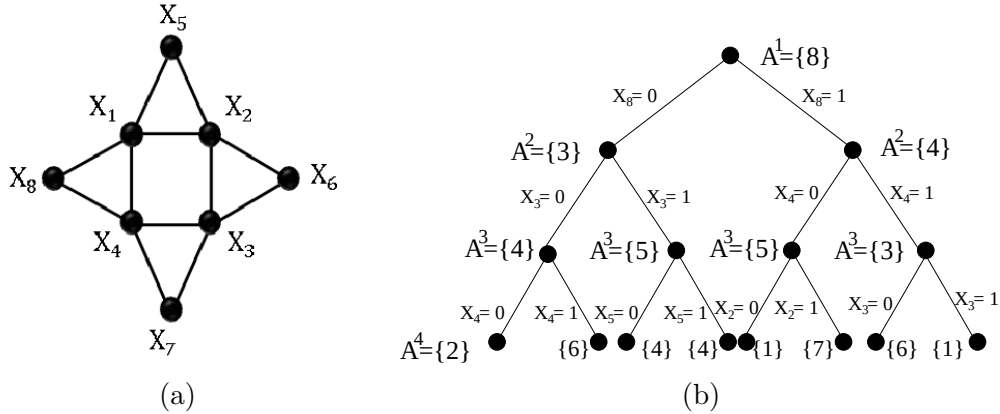


FIGURE 3.19 – Exemple de stratégie adaptative pour l'échantillonnage dans un champ de Markov. (a) : graphe G associée au champ de Markov, (b) : représentation par arbre d'une stratégie adaptative, dans le cas d'une seule observation par étape d'échantillonnage et d'un champ de Markov binaire. Sur les nœuds de l'arbre sont représentées les décisions (indice du nœud échantillonné dans le graphe G) et sur les branches les valeurs observées. La stratégie est adaptative dans le sens où deux réalisations différentes du champ de Markov seront associées à deux trajectoires différentes dans l'arbre.

choisi une fois pour toute avant le début de la campagne d'échantillonnage. Dans le second cas, l'échantillonnage est découpé en plusieurs étapes, l'ensemble est construit au fur et à mesure et les prochains sites échantillonnés sont déterminés en fonction des observations récoltées jusque là ([110] et voir également figure 3.19). Puisque l'espace des stratégies adaptatives inclut l'espace des stratégies statiques⁶, rechercher une stratégie parmi les premières élargit considérablement l'espace des possibles et la meilleure stratégie adaptative ne peut pas faire moins bien que la meilleure stratégie statique. Pourtant l'échantillonnage adaptatif n'est pas toujours choisi sur le terrain car il n'est pas facile de concevoir une stratégie adaptative.

Mes travaux sur le thème de l'échantillonnage adaptatif de données spatiales se découpent en trois contributions décrivent ci-dessous : une modélisation du problème dans le cadre des processus de points et du krigeage, puis une modélisation dans le cadre des champs de Markov avec deux méthodes pour la résolution approchée du problème d'optimisation associé.

Echantillonnage dans une carte binaire et processus de points

La conception de stratégies d'échantillonnage pour la cartographie a été beaucoup étudiée en statistique spatiale, dans le cadre de la géostatistique [13, 41]. Les méthodes existantes reposent en général sur le krigeage ([21]) et les champs gaussiens. Elles sont bien adaptées pour des observations à valeurs continues (pollution, températures). Elles sont moins naturelles dans le cas de données d'occurrence et de données de type classe d'abondance. Nous

6. Une stratégie adaptative fournit des règles de décision dont le résultat dépend des observations passées. Une stratégie statique est un cas particulier où la règle de décision est constante

avons étendu l’approche krigeage et proposé une modélisation du problème de conception par optimisation d’une stratégie adaptative d’échantillonnage pour la reconstruction de carte binaire dans le cadre des processus ponctuels de points ([71]). Le krigeage fournit par définition une prédiction en tout point qui dépend uniquement de la position des points d’échantillonnage et non des valeurs observées. Afin de pouvoir exploiter également ces dernières pour définir la valeur d’une stratégie et donc concevoir un échantillonnage adaptatif, nous avons considéré une version conditionnelle du krigeage. Nous avons ensuite testé l’efficacité d’une heuristique myope en terme de qualité de la carte binaire reconstruite, avec de bons résultats. La difficulté de cette approche reste l’estimation des paramètres du modèle de processus de points, surtout dans le cas de données bruitées.

Echantillonnage dans une carte d’abondance et champ de Markov : une heuristique myope

Nous avons développé une approche de modélisation alternative par champ de Markov. Ce cadre, très classique en analyse d’image, permet de modéliser directement des données de type classe. De plus, il peut s’appliquer à des problèmes d’échantillonnage autres que spatiaux, puisque, contrairement aux champs gaussiens et modèles booléens, la notion de covariance n’est pas nécessairement liée à une notion de distance. Le problème de l’échantillonnage optimal dans les champs de Markov a été abordé par [64]. Dans ces travaux, le cas particulier des champs de Markov définis sur des polyarbres est considéré et des algorithmes sont proposés pour la résolution approchée d’un problème d’échantillonnage statique. Dans le même esprit, nous avons formalisé le problème de la conception d’une stratégie adaptative d’échantillonnage dans un champ de Markov quelconque. Le cas échéant, les observations bruitées sont modélisées avec un champ de Markov caché (CMC). La formalisation requiert tout d’abord de définir une méthode de reconstruction de la carte à partir du résultat de l’échantillonnage et une mesure de la qualité de cette carte. Pour cela nous utilisons les modes des marginales conditionnellement aux observations, bien connues sous le nom de Maximum Posterior Marginals (MPM, [7]), qui fournissent une reconstruction point par point et donnent une indication sur l’incertitude de cette reconstruction. La valeur $V(\delta)$ d’une stratégie est ainsi définie comme l’espérance, sur toutes les trajectoires possibles en suivant δ , de la somme des probabilités conditionnelles des modes. Le calcul de la stratégie de valeur optimale est impossible. Nous avons proposé une première solution heuristique, qui consiste à aller observer à chaque étape de l’échantillonnage, les variables dont l’état reste le plus incertain conditionnellement aux observations déjà acquises [94]. Cette incertitude en un nœud est mesurée par la probabilité du mode local conditionnel. Cette quantité est calculée de manière approchée grâce à un algorithme variationnel de type passage de messages (voir section 2.3) appelé Belief Propagation (BP). Cette heuristique adaptative, appelée BP-max, conduit à de meilleures performances, en terme d’erreur de reconstruction, par rapport à la même heuristique appliquée de manière statique et aussi par rapport un échantillonnage aléatoire. Néanmoins, il s’agit d’une stratégie myope, dans le sens où, à une étape donnée, elle optimise une récompense instantanée, sans tenir compte des étapes à venir. Pour cette raison, elle ne permet pas de gérer les choix d’échantillonnage en fonction d’un budget à respecter.

Echantillonnage dans une carte d’abondance et champ de Markov : l’algo-

rithme LSDP

Dans un second temps, nous avons donc proposé un cadre pour la conception de stratégies adaptatives non myopes d'échantillonnage dans un Champ de Markov. Cette méthode est associée à un algorithme de résolution approchée du problème d'optimisation $\arg \max_{\delta} V(\delta)$, sous contrainte de ne pas dépasser un budget initial. Pour cela, le problème initial a été traduit en un PDM qui ne possède qu'une récompense finale (la somme des probabilités conditionnelles des modes). Puis nous avons exploité les techniques d'apprentissage par renforcement (AR, [108]), classiquement utilisées dans ce cadre. L'AR s'appuie sur l'exploration de paires état-action le long de simulations de trajectoires du PDM pour apprendre les actions optimales. Une application directe des algorithmes d'AR au problème d'échantillonnage n'est cependant pas possible, soit parce que ces algorithmes sont dédiés aux PDM à horizon infini, soit parce que la taille des problèmes traités est trop grande. Néanmoins LSDP s'inspire de l'algorithme LSPI (Least Square Policy Iteration, [65]) pour PDM à horizon infini pour proposer l'algorithme LSDP (Least Square Dynamic Programming, [10, 72]) qui repose sur les principes suivants :

- i)* comme dans LSPI, une approximation paramétrique de la fonction de valeur état-action par une combinaison linéaire d'un ensemble de fonctions de bases mais dont les poids dépendent du temps (puisque en horizon fini la stratégie optimale d'un PDM est non stationnaire),
- ii)* une construction des trajectoires de paires état-action visitées à partir d'un batch ([98]) de simulations du champ de Markov générées off-line,
- iii)* un calcul des poids par moindres carrés et programmation dynamique (car la récompense est uniquement finale).
- iv)* une méthode rapide de calcul des marginales conditionnelles grâce à l'algorithme BP.

L'algorithme LSDP est plus coûteux en temps de calcul que l'heuristique BP-max. Cependant il permet de prendre en compte une contrainte de coût lors de l'exploration des stratégies possibles. Il conduit à des stratégies d'échantillonnage supérieures à celles de BP-max dès lors que le coût n'est pas uniforme dans l'espace. Sur des problèmes jouets d'échantillonnage dans un modèle de Potts sur grille, LSDP s'avère plus performant que les approches classiques d'AR et permet de traiter des problèmes plus grands (figure 3.20).

Dans cette partie de mes travaux, les approximations variationnelles ne sont pas centrales. L'algorithme LSDP repose sur des simulations du champ de Markov. Néanmoins, c'est la combinaison des deux qui a permis de concevoir un algorithme qui fournit des stratégies de bonne qualité tout en restant raisonnable en temps de calcul.

3.4.4 Application de l'échantillonnage adaptatif à la cartographie en épidémiologie et en écologie

Echantillonnage adaptatif pour la cartographie d'une invasion de fourmis "fire ants"

Principaux collaborateurs : Barry Brook (The University of Adelaide – Research Institute for Climate Change and Sustainability, Australie) , Ralph Mac Nally (Monash University, Australia), Régis Sabbadin (MIAT - INRA MIA - Toulouse) et Danny Spring (Monash Uni-

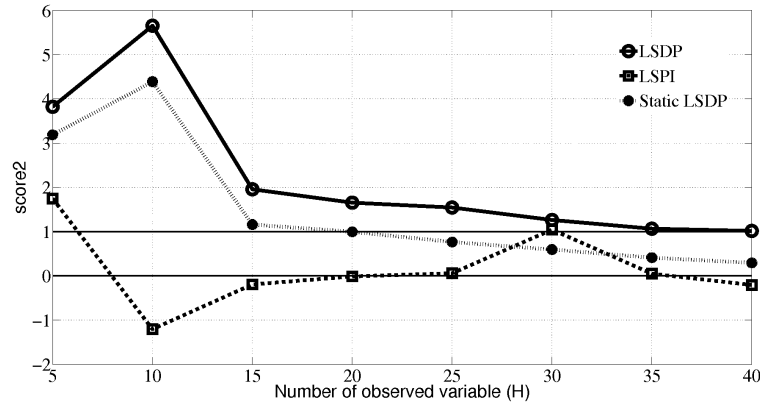


FIGURE 3.20 – Comportement de l’algorithme LSDP dans le cas de l’échantillonnage sur une grille de taille 10×10 . Pour une stratégie δ donnée, $score2(\delta) = \frac{V(\delta) - V(\delta_R)}{|V(\delta_{BP-max}) - V(\delta_R)|}$, où δ_R est une stratégie aléatoire et V désigne la valeur de la stratégie. Une stratégie de score 0 (resp. 1) est une stratégie de même valeur que la stratégie aléatoire (resp. BP-max). L’algorithme LSPI est un algorithme de référence parmi les approches par AR. La stratégie Static-LSDP est obtenue par une variante de LSDP conduisant à des stratégies statiques.

versity, Australia)

Publication associée : 12

Projet : projet de l’Australian Research Council (ARC) «Applying search theory for eradicating invasive species».

La fourmi “fire ant” (*Solenopsis invicta*) est une espèce non native en Australie qui a été découverte pour la première fois près de Brisbane en 2001. Du fait de l’ampleur des dégâts causés au bétail et aux humains (morsures), un programme d’éradication (National Fire Ant Eradication Program) a rapidement été mis en place. Etant donné la surface concernée, une méthode d’échantillonnage est utilisée : l’Adaptive Cluster Sampling (ACS) [110] consiste à réaliser un premier échantillonnage (régulier ou aléatoire) puis à explorer les zones voisines de celles où les fourmis ont été observées, de manière itérative, jusqu’à ce que l’on ne découvre plus de fourmis. Cette méthode d’échantillonnage est adaptative, par contre elle ne prends pas en compte l’objectif de cartographie. La collaboration avec les chercheurs australiens s’est mise en place afin d’étudier d’autres stratégies d’échantillonnage. Le problème présente les caractéristiques suivantes. La zone de surveillance a été divisée en unités d’échantillonnage selon une grille régulière. Sur certaines de ces unités, un traitement pour éradiquer les fourmis a été appliqué. Enfin, la présence des fourmis est détectée via la découverte d’un nid, mais ceux-ci sont difficiles à observer. Il est donc possible que le notateur retourne des faux négatifs. Nous avons modélisé la carte d’occurrence des fourmis et les observations bruitées avec un modèle de champ de Markov caché (figure 3.21), puis nous avons appliqué l’heuristique BP-max ([94]). Comme dans les exemples jouets cités précédemment, nous avons observé une restauration de bien meilleure qualité qu’à partir d’un échantillonnage statique, ou aléatoire, et également une amélioration par rapport à l’échantillonnage ACS. La stratégie LSDP n’a pas été testée sur cette application car elle a

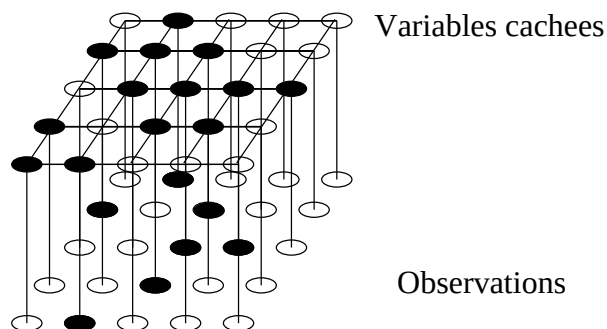


FIGURE 3.21 – Modèle de champ de Markov caché pour la cartographie d’une invasion de fourmis fire ants : en haut les variables cachées et en bas les variables observées. Une variable cachée peut être dans l’état 0 (absence de nid, blanc) ou 1 (présence de nids, noir), une variable observée peut être dans l’état 0 (pas de nid détecté, blanc) ou 1 (nids détectés, noir).

été développée postérieurement à la fin du projet.

Echantillonnage adaptatif pour la cartographie d’une espèce adventice dans une parcelle cultivée

Principaux collaborateurs : Sabrina Gaba (UMR Agroécologie - INRA EA - Dijon), Mathieu Bonneau et Régis Sabbadin (MIAT - INRA MIA - Toulouse)

Publications associées : 16, 48, 49, 50

Projet : Idées’08.

L’heuristique BP-max et l’algorithme LSDP ont été mis en œuvre sur un problème d’échantillonnage d’une espèce d’adventice dans une parcelle. Les données sont des classes d’abondance, et sont supposées non bruitées. Le modèle de champ de Markov issu de l’étude présentée en section 3.2.4 a été utilisé pour représenter la distribution spatiale d’une espèce. Un modèle de coût d’échantillonnage a été développé, basé sur le temps mis par le notateur pour attribuer une note et se déplacer vers le quadrat suivant. Il s’agit donc d’un coût qui dépend à la fois de la position du quadrat et de la valeur de l’observation. A budget équivalent, la stratégie LSDP permet d’observer plus de sites que l’heuristique BP-max. D’autre part ces deux stratégies adaptatives conduisent à de meilleures restaurations de la carte des classes d’abondance que les stratégies adventices classiques (qui sont statiques : en étoile, en W, en Z ou encore régulières, voir figure 3.22). Une parmi ces dernières peut parfois être meilleure pour une carte donnée, mais ce n’est jamais la même et la qualité varie beaucoup d’une carte à l’autre (figure 3.23).

Cette application, comme la précédente, illustre bien l’intérêt d’une stratégie adaptative. La difficulté reste néanmoins que pour les mettre en œuvre l’utilisateur doit disposer d’une estimation de la valeur des paramètres du modèle de champ de Markov. Soit d’autres jeux de données permettent de construire un modèle ad-hoc, soit il est possible d’envisager un premier échantillonnage régulier pour les estimer.

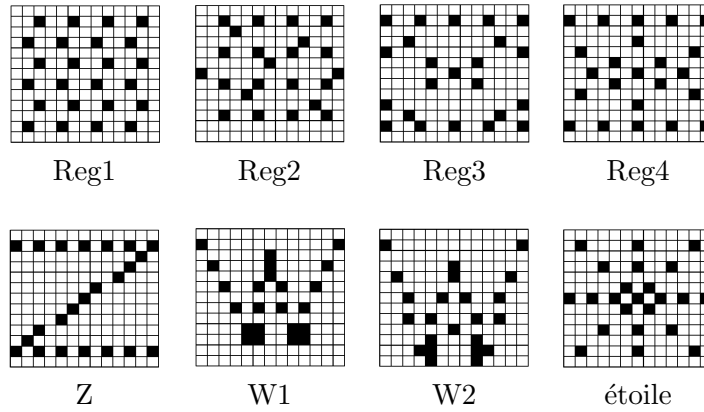


FIGURE 3.22 – Huit échantillonnages statiques classiques, correspondant à 23 quadrats observés.

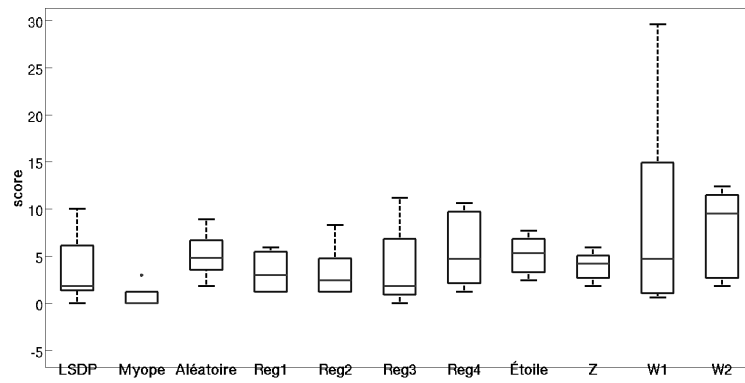


FIGURE 3.23 – Performance des stratégies adaptatives LSDP et BP-max (myope) et des 8 stratégies statiques. Le score utilisé mesure la différence entre nombre de quadrats sur lesquels la classe a été bien estimée, pour la meilleure stratégie et celle considérée.

3.5 Conclusions

Les travaux qui composent ce mémoire d'HDR sont le résultat d'un ensemble de collaborations riches et variées, avec des statisticiens, informaticiens, écologues théoriciens, modélisateurs, agronomes, écologues. Le réseau de ces collaborations est représenté schématiquement sur la figure 3.24. Il a permis des contributions en analyse d'image, intelligence artificielle, physique théorique, statistique computationnelle, systèmes complexes, épidémiologie, foresterie, écologie, ... qui sont résumées dans la table 3.1.

Ces travaux démontrent la variété des domaines d'application des méthodes variationnelles et leur intérêt pour la résolution approchée de problèmes structurés de grande taille. La table 3.1 classe les différents types de modèle graphique que j'ai manipulés, selon les trois axes suivants : structure spatiale explicite, structure temporelle explicite, structure de décision explicite. Pour chacun de ces modèles il est possible de définir une approximation variationnelle et dans tous les cas cela a conduit à des avancées intéressantes en estimation, inférence ou décision. Les méthodes variationnelles sont maintenant une alternative reconnue aux méthodes basées sur la simulation. Cette reconnaissance est forte dans la communauté machine learning, où les chercheurs ont fait progresser les aspects algorithmiques et la validation expérimentale de ces méthodes. Elle reste plus modérée en statistique du fait du manque de résultats théoriques sur la qualité des méthodes variationnelles, cependant cette communauté commence à s'emparer de la question.

L'aspect des modèles graphiques de plus en plus présent dans mes travaux est l'aspect décisionnel. Une des raisons, comme je l'ai déjà mentionnée, est que pour cette tâche les méthodes variationnelles restent encore sous-exploitées. D'autre part, les avancées méthodologiques auxquelles j'ai pu contribuer ont été validées dans différents domaines du vivant, en agronomie, épidémiologie et écologie. Très certainement du fait de ma sensibilité INRA, ces travaux de transfert des méthodes vers les biologistes prennent également une place croissante dans mes activités. Ces dernières années cette double évolution de mon positionnement se concrétise de plus en plus, avec la construction du thème "Décision Spatialisée" au sein de l'équipe MAD, co-animé avec Régis Sabbadin. Rares sont les chercheurs, en France, positionnés sur cette problématique.

	Modèle			Contribution	Application	Publications
	spatial	temporel	décisionnel			
Estimation	X			Algorithme EM de type champ moyen pour les champs de Markov cachés	Cartographie de risque épidémique Cartographie d'espèces adventices des cultures	1, 2, 14, 15, 19, 24, 25, 37, 38, 40, 41, 47, 48, 49, 51, 53, 54
	X			Algorithme VBEM pour le modèle de Cox log gaussien		46
Inférence	X	X		Méthodes variationnelles d'ordre 2 pour l'inférence sur les états transients et à l'équilibre du processus de contact sur graphe	Etude du rôle du réseau d'interaction en épidémiologie	5, 9, 22, 23
	X	X		Approximation variationnelle de la taille de l'épidémie dans un modèle SIR sur graphe		55
Décision	X	X	X	Cadre PDMG et algorithme IPA-CM pour la résolution de PDM à espaces d'état et d'action factorisés	Contrôle d'épidémies Gestion forestière Gestion de la durabilité des résistances	13, 29, 31, 42, 18, 30, 32, 21, 51
	X	X	X	Algorithme LSDP et heuristique BP-max pour la conception de stratégies adaptatives d'échantillonnage dans les champs de Markov	Cartographie d'espèces invasives Cartographie d'espèces adventices des cultures	12, 16, 33, 34, 35, 36, 37, 43, 44, 45 17, 47, 48, 49
Construction	X	X		Caractérisation du maximum d'entropie sous contraintes de marginales		11
	<i>non applicable</i>			Tests de permutation pour l'analyse de données sur grille	Analyses de dynamiques en épidémiologie végétale	4, 6, 8, 10

TABLE 3.1 – Résumé des principales contributions dans les tâches de construction, estimation, inférence et décision pour les modèles graphiques.

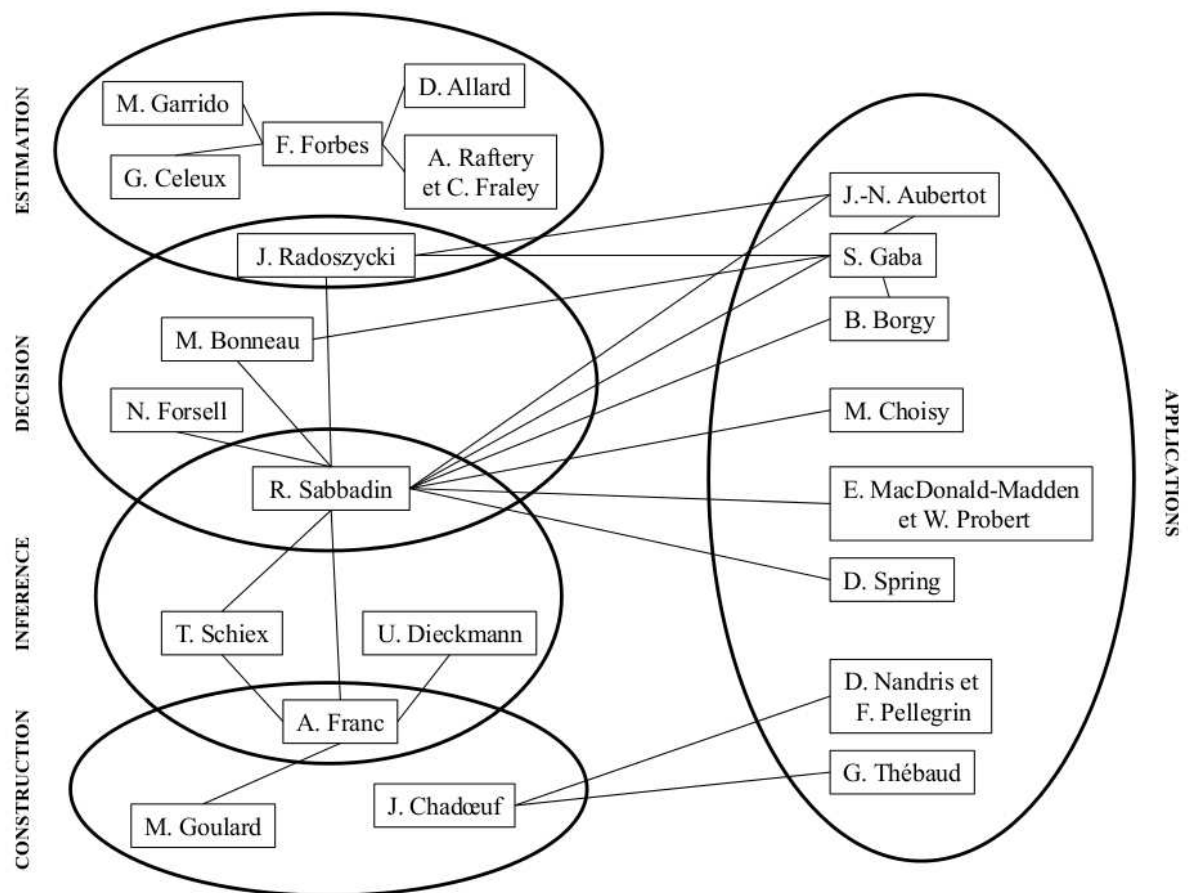


FIGURE 3.24 – Réseau de collaborateurs. Un lien existe entre deux personnes si ces deux personnes et moi-même sommes co-auteurs d’une publication commune.

Chapitre 4

Projet de Recherche

Titre : Modélisation et algorithmique pour la conception de stratégies de gestion dans des processus dynamiques structurés

4.1 Enjeux

La part croissante des travaux finalisés dans mes recherches, dans les domaines de l'écologie et l'épidémiologie, m'ont fait prendre conscience qu'il existe encore très peu d'outils disponibles pour la conception de stratégies de gestion dans ces domaines, du fait de la complexité des phénomènes étudiés. Si certains laboratoires de recherche finalisée intègrent le terme "gestion" dans leur intitulé, bien souvent les recherches de ces laboratoires se concentrent sur l'acquisition de connaissances ou la conception par comparaison de stratégies. La conception par optimisation est rarement abordée, certainement car elle semble encore hors de portée. Pourtant, en écologie, les travaux pionniers du Spatial Ecology Lab (University of Queensland) et du Conservation Decision Lab (CSIRO) à Brisbane sur la mise en œuvre des outils de l'intelligence artificielle pour l'optimisation de l'allocation de ressources ont démontré que l'on peut obtenir des résultats concrets, utiles pour les décideurs ([75, 16, 96, 100]). En épidémiologie végétale ou animale, la conception de stratégies de gestion par optimisation d'un critère est moins développée. Pourtant, comme l'illustrent les quelques applications décrites dans ce manuscrit, les questions de gestion se posent, que ce soit pour contrôler la propagation des adventices, de pathogènes ou de virus, ou pour préserver la durée d'une résistance variétale, la biodiversité dans les zones cultivées, ...

Le cadre de la décision séquentielle dans l'incertain est a priori bien adapté pour représenter ces problèmes où il s'agit d'optimiser au cours du temps une récompense qui dépend de l'état d'un processus dynamique et des actions effectuées pour contrôler ce processus. Néanmoins, le cas de problèmes structurés (que la structure soit liée au spatial ou non), comme on en rencontre en épidémiologie et en écologie, n'a pas encore de solution satisfaisante. La représentation structurée du problème, à travers le réseau d'interactions entre les entités du processus, et entre les entités et les actions, est relativement bien exploitée pour

arriver à une représentation compacte. Par contre, il n'existe pas encore d'algorithme satisfaisant (combinant complexité raisonnable et qualité de l'approximation) pour la résolution approchée. Les algorithmes dédiés à la résolution des PDMF dans leur version avec espace d'action factorisé (PDMF-AF, [62]) et des PDMG ([103]) sont des premiers résultats qui demandent à être améliorés.

Les recherches méthodologiques que j'entends mener pour les prochaines années ont donc pour objectif, d'une part de progresser sur les résultats obtenus en développant des algorithmes pour la résolution approchée de problèmes de décision dans l'incertain lorsque le phénomène contrôlé est structuré, et d'autre part de proposer des modélisations des problèmes de gestion en épidémiologie et en écologie. Le premier axe conduira à des résultats génériques alors que le second apportera des réponses spécifiques à chaque application. Je les décris respectivement dans les sections 4.2 et 4.3. Le schéma de la figure 4.1 résume les questions de recherche associées à ces deux axes et leurs dépendances, ainsi que les projets en cours ou soumis pour les faire progresser (les projets en cours sont détaillés dans la partie CV de ce document).

4.2 Objectifs méthodologiques en décision

4.2.1 Contrôle d'un processus parfaitement observé : algorithmes de résolution approchée pour les PDM à espaces d'état et d'action factorisés.

La suite des développements méthodologiques pour la résolution de problèmes de décision structurés va maintenant consister à lever certaines hypothèses des cadres PDMF-AF et PDMG et surtout des algorithmes de résolution associés, qui comme on a pu le constater sont contraignantes dans la modélisation de problèmes réels. Les algorithmes de résolution des PDMF-AF sont spécifiques au cas de variables binaires et seront très certainement trop coûteux en temps de calcul s'il sont étendus au cas discret fini. Dans le cadre PDMG, une des limitations est l'hypothèse que les fonctions de transition et de récompense en i ne dépendent que de l'action effectuée en i . L'autre est la restriction à la recherche de solutions parmi l'espace restreint des politiques locales, puisqu'il a été démontré qu'une stratégie locale peut être arbitrairement mauvaise. Enfin, quel que soit le cadre, il n'est pas possible actuellement de résoudre efficacement un problème avec contrainte globale sur les actions ou un problème de recherche de stratégies impliquant des actions de contrôle non seulement sur les nœuds mais aussi sur les arêtes du réseau d'interaction. L'objectif de la thèse de Julia Radoszycki, qui a débuté en octobre 2012, sera de proposer des solutions algorithmiques pour lever certaines de ces limites.

4.2.2 Échantillonnage d'un processus partiellement observé : méthodes pour l'échantillonnage adaptatif lorsque le phénomène évolue dans le temps.

Nos travaux d'échantillonnage pour la cartographie s'intéressaient à une photo à l'instant t du processus. Or, en écologie ou en épidémiologie, ce que l'on cherche à observer

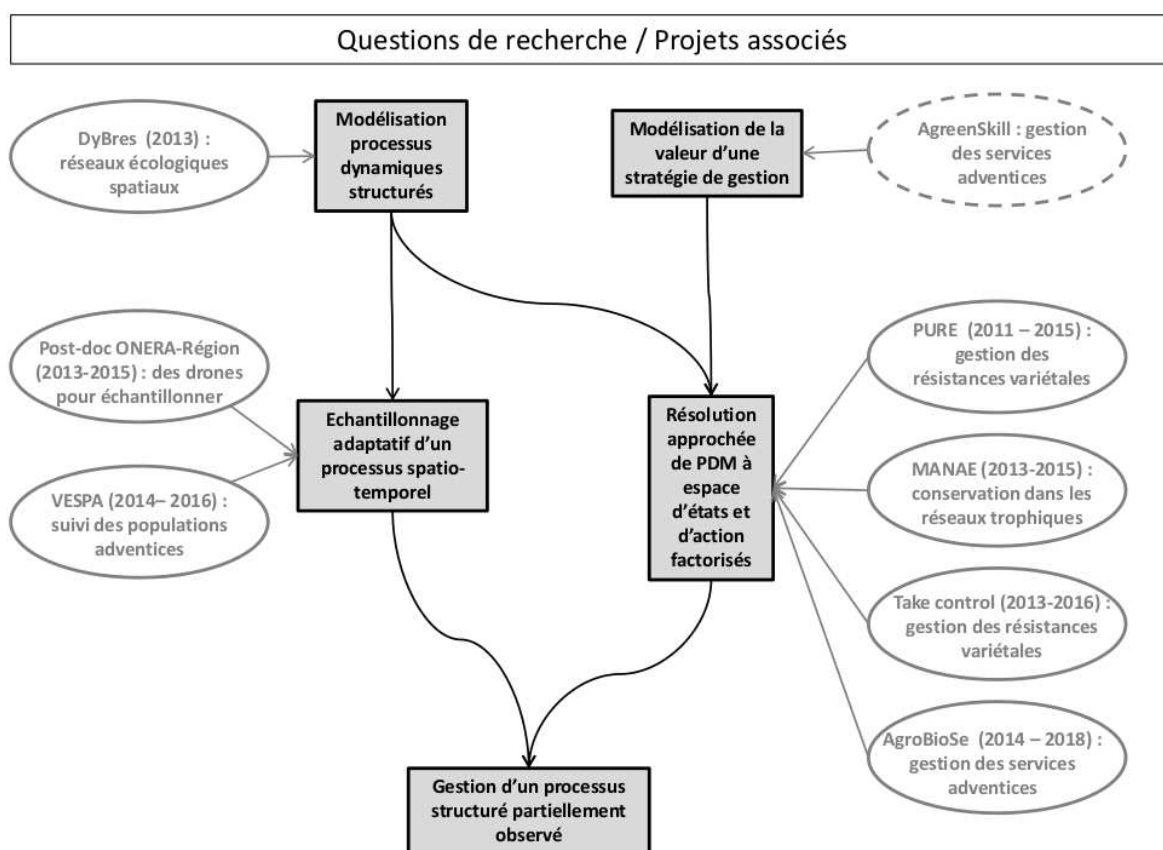


FIGURE 4.1 – Articulation entre les différents axes du projet de recherche, et projets financés associés.

naît, se déplace et disparaît. Les processus sont dynamiques. Le suivi de ces processus demande de savoir répondre aux questions suivantes : combien de sites d’observation sont nécessaires ? Où les positionner ? Faut-il faire évoluer la position de ces sites au cours du temps ? Ou en rajouter ? ... Pour répondre à ces questions, en théorie, l’algorithme LSDP peut être appliqué, puisqu’il s’écrit pour n’importe quel PDM. Le batch généré ne sera plus un batch de cartes mais un batch de trajectoires de cartes. Les simulations et les calculs de marginales, même avec l’approximation Belief Propagation, risquent de conduire à un algorithme trop coûteux en temps. Il faudra envisager d’autres approches que l’optimisation par simulation et disposer de méthodes d’inférence rapides pour des modèles type réseau bayésien dynamique.

4.2.3 Compromis entre gestion et acquisition d’information dans un processus partiellement observé.

A plus long terme, l’objectif est de combiner les résultats de ces deux axes pour aborder le problème plus général de la gestion dans le temps d’un processus structuré partiellement observé. Dans beaucoup de processus contrôlés en écologie et en épidémiologie, la connaissance du processus est imparfaite car l’information est trop coûteuse à acquérir. Il faut décider comment répartir les ressources entre l’acquisition d’information et le contrôle, la répartition optimale ne consistant pas toujours à attendre d’avoir l’information la plus complète avant d’agir. Ainsi, en pratique, en même temps qu’une stratégie de gestion des adventices se met en place à l’échelle de la mosaïque paysagère, il faut choisir les parcelles sur lesquelles la population va être observée. En biologie de la conservation, le réseau trophique reliant les espèces à gérer est en général lui aussi connu imparfaitement : les espèces étant rares, il est difficile de détecter leur présence. Il peut être tout aussi difficile de déterminer “qui mange qui”. Le choix des espèces ou des liens trophiques sur lesquels on alloue un effort d’observation est à faire en même temps que la prise de décision de conservation, sur un budget commun.

La question du contrôle de processus partiellement observés et du compromis entre contrôle et acquisition d’information a été étudiée dans le cas non structuré, avec le cadre des PDM sur croyance (belief-PDM, [54]). Ce sera le point de départ pour aborder le cas des processus structurés.

4.2.4 Quelques pistes méthodologiques

Les méthodes variationnelles vont continuer à avoir une place importante dans la conception d’algorithmes innovants pour la décision, car elles restent encore relativement sous-exploitées dans ce domaine. Elles peuvent être utilisées seules, mais il s’avère aussi que dans certains problèmes, la solution pour atteindre un bon compromis temps/qualité consiste à les coupler avec des méthodes de Monte-Carlo. C’est ce qui a été proposé avec l’algorithme LSDP. Ce constat est également vrai en estimation, avec l’algorithme EM en champ simulé et l’algorithme VBEM pour le modèle de Cox log gaussien. Cette stratégie de combinaison des méthodes variationnelles et MCMC sera certainement à ré-exploiter sur les problèmes décrits ci-dessus.

L’option envisagée dans le cadre de la thèse de Julia Radoszyki, pour proposer de nouveaux algorithmes pour la résolution approchée de PDM à espaces d’état et d’action factorisés,

est de considérer des approximations variationnelles d'ordre supérieur à 1 (le champ moyen, utilisé pour API-CM est d'ordre 1). Une piste est d'étendre l'espace de recherche à des politiques stochastiques (au lieu de déterministes locales comme dans les PDMG), ce qui permet de transformer le problème initial en un problème d'inférence dans un modèle graphique. Dans ce cadre, des travaux récents ont ainsi exploité les algorithmes de passage de message soit pour résoudre de manière exacte des PDM "simples" non factorisés ([111, 42]), soit pour le calcul approché de l'utilité espérée maximale dans un diagramme d'influence ([69]). La transposition au cas des PDM à espaces d'état et d'action factorisés demandera de définir des algorithmes de passage de messages généralisés ([116]) appropriés.

Enfin, les algorithmes de résolution d'un PDM demandent de réaliser de nombreux calculs d'inférence (calcul de marginales). Les algorithmes de passage de message ont été développés initialement pour cela et sont efficaces pour des processus statiques. Pour les processus dynamiques, j'envisage de retourner à la source des méthodes variationnelles, en mécanique statistique, pour étudier la méthode du "path probability" ([29]), qui étend aux processus dynamiques le cadre variationnel tel que décrit en 2.1. Une autre piste, issue du monde déterministe, consiste à voir le modèle graphique représentant le processus comme un problème de satisfaction de contraintes pondérées (Weighted Constraint Satisfaction Problems, WCSP, [105]). Cela permet d'appliquer des algorithmes efficaces de résolution de WCSP pour le calcul de marginales ou de mode d'un modèle graphique. Je suis impliquée dans une collaboration en cours au sein de l'unité UBIAT sur ce sujet.

4.3 Objectifs en modélisation et applications

L'application principale sur laquelle ces développements génériques seront utilisés concerne la gestion et l'échantillonnage des espèces adventices dans la mosaïque paysagère. Les espèces adventices sont nuisibles aux cultures car elles entraînent des pertes de rendement. Cependant elles jouent également un rôle positif dans la préservation de la biodiversité en tant que ressources trophiques ou hôtes de nombreuses espèces ([95, 32]). Un des enjeux en agro-écologie aujourd'hui est de concevoir des stratégies de gestion des adventices qui optimisent un compromis entre rendement et services écologiques rendus.

La seconde application est dans le prolongement des travaux en agronomie sur le contrôle du phoma du colza. En agronomie, un moyen de lutter contre les maladies des cultures est d'utiliser des variétés résistantes. Cependant, une surexploitation conduit rapidement à un contournement de cette résistance, qui devient inutile. Se pose alors la question de la répartition dans l'espace et dans le temps des variétés résistantes et susceptibles afin de préserver la durée de la résistance tout en limitant les pertes de rendements liées aux maladies.

Enfin, en écologie, les outils génériques de l'axe précédent permettront de mieux représenter et résoudre le problème d'allocation de ressources pour la préservation d'espèces en interaction dans un réseau trophique. En particulier, dans cette application, la question du compromis entre gestion et acquisition d'information sera abordée.

Pour ces trois applications, la mise en œuvre des algorithmes pour la conception de stratégies de gestion demandera en amont de construire un modèle de la dynamique du processus contrôlé et de définir le critère qui quantifiera la valeur d'une stratégie. Cela

définit les deux directions de recherche de cet axe :

Modélisation de processus structurés dynamiques. Dans la famille des modèles graphiques, les réseaux bayésiens dynamiques ([43]) sont particulièrement bien adaptés pour modéliser des dynamiques de processus directement exploitables dans le cadre des PDM factorisés. En effet, l'espace d'état est discret fini, le temps est discret et la factorisation de l'espace d'état découle de la structure du réseau bayésien. La modélisation des différentes applications dans ce cadre soulèvera des questions de sélection de modèles et d'estimation des paramètres.

Modélisation de la valeur d'une stratégie de gestion. Comment mesurer la biodiversité? Comment comparer sur une même échelle la production d'une région agricole et la durabilité d'une résistance variétale? ou le coût de protection d'une espèce en voie de disparition et le maintien de cette même espèce? Plus généralement comment modéliser le compromis entre les gains et les coûts associés à une stratégie? Les réponses à ces questions sont multiples et peuvent conduire à des fonctions de récompense d'un PDM de structures différentes, qu'il sera important de savoir traiter dans la construction des algorithmes de résolution.

Bibliographie

- [1] R. Albert and A.-L. Barabasi. Statistical mechanics of complex network. *Review of Modern Physics*, 74 :47–97, 2002.
- [2] D. Barber. *Bayesian Reasoning and Machine Learning*. Cambridge University Press, New York, NY, USA, 2012.
- [3] M.J. Beal. *Variational Algorithms for Approximate Bayesian Inference*. PhD thesis, Gatsby Computational Neuroscience Unit, University College London, 2003.
- [4] C. Berge. *Graphs and hypergraphs*. Elsevier, 1973.
- [5] D. P. Bertsekas and J. N. Tsitsiklis. *Neuro-Dynamic Programming*. Athena Scientific, Belmont, Massachussetts, 1996.
- [6] J. Besag. Statistical analysis of non-lattice data. *The Statistician*, 24 :179–195, 1975.
- [7] J. Besag. On the statistical analysis of dirty pictures. *Journal of the Royal Statistical Society B*, 48(3) :259–302, 1986.
- [8] J. Besag, J. York, and A. Mollié. Bayesian image restoration with two applications in spatial statistics, with discussion. *Annals of the Institute of Statistical Mathematics*, 43(1) :1–59, 1991.
- [9] C. M. Bishop. *Pattern Recognition and Machine Learning*. Springer, 2006.
- [10] M. Bonneau, N. Peyrard, and R. Sabbadin. A reinforcement-learning algorithm for sampling design in Markov random fields. In *Proceedings of the European Conference on Artificial Intelligence*, Montpellier, France, août 2012.
- [11] C. Boutilier, T. Dean, and S. Hanks. Decision-theoretic planning : Structural assumptions and computational leverage. *Journal of Artificial Intelligence Research*, 11 :1–94, 1999.
- [12] C. Boutilier, R. Dearden, and M. Goldszmidt. Stochastic dynamic programming with factored representations. *Artificial Intelligence*, 121(1) :49–107, 2000.
- [13] M. Buesco, J. Angulo, and Alonso F. A state-space model approach to optimum spatial sampling design based on entropy. *Environmental and Ecological Statistics*, (5) :29–44, 1998.

- [14] G. Caldarelli, R. Pastor-Satorras, and A. Vespignani. Structure of cycles and local ordering in complex networks. *European Physical Journal B*, 38 :183–186, 2004.
- [15] G. Celeux, F. Forbes, and N. Peyrard. EM procedures using mean field-like approximations for Markov model-based image segmentation. *Pattern Recognition*, 36 :131–144, 2003.
- [16] I. Chadès, E. McDonald-Madden, M.A. McCarthy, B. Wintle, M. Linkie, and H.P. Possingham. When to stop managing or surveying cryptic threatened species. *Proceedings of the National Academy of Sciences*, 105(37) :13936–13940, 2008.
- [17] J. Chadoeuf, D. Nandris, J. P. Geiger, M. Nicole, and J. C. Pierrat. Modélisation spatio-temporelle d’une épidémie par un processus de Gibbs : Estimation et tests. *Biometrics*, 48 :1165–1175, 1992.
- [18] D. Chandler. *Introduction to Modern Statistical Mechanics*. Oxford University Press, 1987.
- [19] M. Charras-Garrido, D. Abrial, N. Peyrard, and S. Dachian. New classification method for disease mapping based on discrete hidden Markov random fields. *Biostatistics*, 13 :241–255, 2012.
- [20] M. Charras-Garrido, L. Azizi, F. Forbes, S. Doyle, N. Peyrard, and D. Abrial. On the difficulty to delimit disease risk hot spots. *International Journal of Applied Earth Observation and Geoinformation*, 22 :99–105, 2013.
- [21] J.P. Chilès and P. Delfiner. *Modeling Spatial Uncertainty*. Wiley series in Probability and Statistics, 1999.
- [22] V. Colizza, A. Barrat, M. Barthélemy, and A. Vespignani. The modeling of global epidemics : stochastic dynamics and predictability. *Bulletin of Mathematical Biology*, 68, 2006.
- [23] D.-P. de Farias and B. Van Roy. The linear programming approach to approximate dynamic programming. *Operations Research*, 51(6) :850–865, 2003.
- [24] J. de Gruijter, D. Brus, M. Bierkens, and K. Knotters. *Sampling for Natural Resource Monitoring*. Springer, 2006.
- [25] R. Dechter. *Constraint processing*. Morgan Kaufmann, 2003.
- [26] A. Dempster, N. Laird, and D. Rubin. Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society B*, 39 :1–38, 1977.
- [27] R. Dickman and M. Martins de Oliveira. Quasi-stationary simulation of the contact process. *Physica A*, 357 :134–141, 2005.
- [28] U. Dieckmann, R. Law, and J. A. J. Metz, editors. *The geometry of ecological interactions - Simplifying spatial complexity*. Cambridge Studies in Adaptive Dynamics. Cambridge University Press, UK, 2000.

- [29] F. Ducastelle. Variational and mean field formulations of the cluster variational method and the path probability method. *Progress of Theoretical Physics Supplement*, 115 :255–272, 1994.
- [30] R. Durrett and S. Levin. The importance of being discrete (and spatial). *Theoretical Population Biology*, 46 :363–394, 1994.
- [31] V. Eguiluz and K. Klemm. Epidemic threshold in structured scale-free networks. *Physical Review Letters*, 89, 2002.
- [32] D.M. Evans, M. Pocock, J. Brooks, and J. Memmott. Seeds in farmland food-webs : Resource importance, distribution and the impacts of farm management. *Biological Conservation*, 144 :2941–2950, 2011.
- [33] M. J. Ferrari, S. Bansal, L. A. Meyers, and O. N. Bjørnstad. Network frailty and the geometry of herd immunity. *Proceedings of the Royal Society of London Series B—Biological Sciences*, 273(1602) :2743–2748, 2006.
- [34] J. A. N. Filipe and G. J. Gibson. Studying and approximating spatio-temporal models for epidemic spread and control. *Philosophical Transactions of the Royal Society of London*, 353(57) :2153–2162, 1998.
- [35] J. A. N. Filipe and G. J. Gibson. Comparing approximations to spatio-temporal models for epidemics with local spread. *Bulletin of Mathematical Biology*, 63 :603–624, 2001.
- [36] F. Forbes, M. Charras-Garrido, L. Azizi, S. Doyle, and D. Abrial. Spatial risk mapping for rare disease with hidden Markov fields and variational EM. *Annals of applied statistics*. In Press.
- [37] F. Forbes and N. Peyrard. Hidden Markov models selection criteria based on mean field-like approximations. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 25(9) :1089–1101, 2003.
- [38] N. Forsell and R. Sabbadin. Approximate linear-programming algorithm for graph-based Markov decision processes. In *Proceedings of the European Conference on Artificial Intelligence*, pages 590–594, Riva Del Garda, Italy, 2006.
- [39] A. Franc. Metapopulation dynamics as a contact process on a graph. *Ecological Complexity*, 1 :49–63, 2004.
- [40] A. Franc, M. Goulard, and N. Peyrard. Chordal graphs to identify graphical models solutions of maximum of entropy under constraints on marginals. *SIAM Discrete Mathematics*, 24(3) :1104–1116, 2010.
- [41] M. Fuentes, A. Chaudhuri, and D. Holland. Bayesian entropy for spatial sampling design of environmental data. *Environmental and Ecological Statistics*, 14 :323–340, 2007.

- [42] T. Furnstion and D. Barber. Efficient inference in Markov control problems. In *Proceedings of the Annual Conference on Uncertainty in Artificial Intelligence*, Barcelona, Spain, 2011.
- [43] Z. Ghahramani. *Adaptive Processing of Sequences and Data Structures*, chapter Learning Dynamic Bayesian Networks. Berlin : Springer-Verlag, 1998.
- [44] Z. Ghahramani and J. Beal. Graphical models and variational methods. In Manfred Oppert and David Saad, editors, *Advanced mean field methods. Theory and practice*, pages 161–177. MIT press, 2001.
- [45] P. J. Green and S. Richardson. Hidden Markov models and disease mapping. *Journal of the American Statistical Association*, 97 :1055–1070, 2001.
- [46] C. Guestrin, D. Koller, R. Parr, and S. Venkataraman. Efficient solution algorithms for factored MDPs. *Journal of Artificial Intelligence Research*, 19 :399–468, 2003.
- [47] T. E. Harris. Contact interactions on a lattice. *Annals of probability*, 2 :969–988, 1974.
- [48] J. Hoey, R. St-Aubin, A. Hu, and C. Boutilier. SPUDD : Stochastic Planning Using Decision Diagrams. In *Proceedings of the Annual Conference on Uncertainty and Artificial Intelligence*, Stockholm, Sweden, 1999.
- [49] L. Hufnagel, D. Brockmann, and T. Geisel. Forecast and control of epidemics in a globalized world. *Proceedings of the National Academy of Sciences*, 101(42) :15124–15129, 2004.
- [50] E. T. Jaynes. Information theory and statistical mechanics, I. *Physical Review*, 106 :620–630, 1957.
- [51] F. V. Jensen and T. D. Nielsen. *Bayesian Networks and Decision Graphs*. Springer Publishing Company, Incorporated, 2nd edition, 2007.
- [52] F. Jordan. Keystone species and food webs. *Philosophical Transactions of The Royal Society, B*, 364 :1733–1741, 2009.
- [53] M. I. Jordan, editor. *Learning in graphical models*, volume 19 of *Adaptive Computation and Machine Learning series*. MIT Press, 2004.
- [54] L.P. Kaelbling, M.L. Littman, and A.R. Cassandra. Planning and acting in partially observable stochastic domains. *Artificial Intelligence Journal*, 101 :99–134, 1998.
- [55] M. J. Keeling. The effects of local spatial structure on epidemiological invasions. *Proceedings of the Royal Society of London B*, 266 :859–867, 1999.
- [56] M. J. Keeling. The implications of network structure for epidemic dynamics. *Theoretical Population Biology*, 67 :1–8, 2005.

- [57] M. J. Keeling, M. E. Woolhouse, D.J. Shaw, L. Matthews, M. Chase-Topping, D.T. Haydon, S. J. Cornell, J. Kappey, J. Wilesmith, and B.T. Grenfell. Dynamics of the 2001 UK foot and mouth epidemic : stochastic dispersal in a heterogeneous landscape. *Science*, 294 :813–817, 2001.
- [58] J. Kennedy. *Dynamic programming. Applications to agriculture and natural resources*. Elsevier Applied Science Publishers, 1986.
- [59] C. Keribin. Méthodes bayésiennes variationnelles : concepts et applications en neuroimagerie. *Journal de la Société Française de Statistiques*, 151(2), 2010.
- [60] R. Kikuchi. A theory of cooperative phenomena. *Physical Review*, 81(6) :988–1003, 1951.
- [61] R. Kikuchi. CVM entropy algebra. *Progress of Theoretical Physics Supplement*, 115 :1–26, 1994.
- [62] K-E. Kim and T. Dean. Solving factored MDPs with large action space using algebraic decision diagrams. In *Proceedings of the Pacific Rim International Conference on Artificial Intelligence*.
- [63] D. Koller and N. Friedman. *Probabilistic Graphical Models : Principles and Techniques*. MIT Press, 2009.
- [64] A. Krause and C. Guestrin. Optimal value of information in graphical models. *Journal of Artificial Intelligence Research*, 35 :557–591, 2009.
- [65] M. Lagoudakis and R. Parr. Least-squares policy iteration. *Journal of Machine Learning Research*, 2003.
- [66] S. L. Lauritzen. *Graphical Models*. Clarendon Press, 1996.
- [67] M. A. Leibold, M. Holyoak, N. Mouquet, P. Amarasekare, J. M. Chase, M. F. Hoopes, R. D. Holt, J. B. Shurin, R. Law, D. Tilman, M. Loreau, and A. Gonzalez. The metacommunity concept : a framework for multi-scale community ecology. *Ecology Letters*, 7 :601–613, 2004.
- [68] S. Z. Li. *Markov random field modeling in image analysis*. Springer-Verlag, 2001.
- [69] Q. Liu and A. Ihler. Belief propagation for structured decision making. In *Proceedings of the Annual Conference on Uncertainty in Artificial Intelligence*, August 2012.
- [70] E. Lô-Pelzer, L. Bousset, M.H. Jeuffroy, M. U. Salam, X. Pinochet, M. Boillot, , and Aubertot J.N. SIPPOM-WOSR : a simulator for integrated pathogen population management to manage phoma stem canker on winter oilseed rape. I. description of the model. *Field Crops Research*, 118 :73–81, 2009.
- [71] R. Sabbadin M. Bonneau, N. Peyrard. Echantillonnage spatial basé sur le krigeage pour la reconstruction de carte d’occurrence. In *Conférence en Reconnaissance de Formes et Intelligence Artificielle*, Caen, France, 2010.

- [72] R. Sabbadin M. Bonneau, N. Peyrard. Reinforcement-learning for sampling design in Markov random fields. In *Proceedings of the International Conference on Computational Statistics*, Limassol, Chypre, 2012.
- [73] B.F.J. Manly. *Randomization, bootstrap and Monte Carlo methods in biology*. Chapman and Hall, London, 1997.
- [74] J. Marro and R. Dickman. *Nonequilibrium Phase Transition in Lattice Models*. Monographs and texts in statistical physics, Collection Alea Saclay. Cambridge University Press, Cambridge, UK, 1999.
- [75] T.G. Martin, I. Chadès, P. Arcese, P.P. Marra, H.P. Possingham, and D.R. Norris. Optimal conservation of migratory species. *PLoS One*, 2(8) :e751, 2007.
- [76] E McDonald-Madden, R. Sabbadin, P. W. J. Baxter, I. Chadès, and H. P. Possingham. Making food web theory useful for conservation decisions. In prep., 2013.
- [77] P. Mielke and K. Berry. *Permutation methods. A distance function approach*. Springer, New York, 2001.
- [78] T. Minka. *A family of algorithms for approximate bayesian inference*. PhD thesis, MIT, 2001.
- [79] J. Møller, A. R. Syversveen, and R. P. Waagepetersen. Log gaussian cox processes. *Scandinavian Journal of Statistics*, 25 :451–482, 1998.
- [80] P. Neal. SIR epidemics on a Bernoulli random graph. *Journal of Applied Probability*, 40(3) :779–782, 2003.
- [81] R. M. Neal and G. E. Hinton. A view of the EM algorithm that justifies incremental, sparse and other variants. In *Learning in graphical models*, pages 355–368. Dordrecht, Kluwer Academic Publishers, 1998.
- [82] M. E. J. Newman. The structure and function of complex networks. *SIAM Review*, 45 :167–256, 2003.
- [83] M. J. E. Newman. Spread of epidemic disease on networks. *Physical Review E*, 66, 2002.
- [84] Manfred Opper and David Saad, editors. *Advanced Mean Field Methods. Theory and Practice*. Neural Information Processing Series. MIT Press, 2001.
- [85] G. Parisi. *Statistical Field Theory*. Addison-Wesley, 1988.
- [86] R. Pastor-Satorras and A. Vespignani. Epidemic spreading in scale-free networks. *Physical Review Letters*, 86 :3200–3203, 2001.
- [87] J. Pearl. *Heuristics : Intelligent Search Strategies for Computer Problem Solving*. Reading, MA : Addison-Wesley, 1984.

- [88] N. Peyrard, A. Calonnec, F. Bonnot, and J. Chadœuf. Explorer un jeu de données sur grille par tests de permutation. *Revue de Statistique Appliquée*, 53(1), 2005.
- [89] N. Peyrard, U. Dieckmann, and A. Franc. Long-range correlations improve understanding of the influence of network structure on contact dynamics. *Theoretical Population Biology*, 73(3) :383–94, 2008.
- [90] N. Peyrard and A. Franc. Cluster variation approximations for a contact process living on a graph. *Physica A*, 358 :575–592, 2005.
- [91] N. Peyrard, F. Pellegrin, J. Chadœuf, and D. Nandris. Statistical analysis of the spatio-temporal dynamics of rubber bark necrosis : no evidence of pathogen transmission. *Forest Pathology*, 36 :360–371, 2006.
- [92] N. Peyrard and R. Sabbadin. Mean field approximation of the policy iteration algorithm for graph-based Markov decision processes. In *Proceedings of the European Conference on Artificial Intelligence*, pages 595–599, Riva Del Garda, Italy, 2006.
- [93] N. Peyrard, R. Sabbadin, E. Lô-Pelzer, and J. N. Aubertot. A graph-based Markov decision process framework applied to the optimization of strategies for integrated management of diseases. In *American Phytopathological Society and Society of Nematologist joint meeting*.
- [94] N. Peyrard, R. Sabbadin, D. Spring, B. Brook, and R. Mac Nally. Model-based adaptive spatial sampling for occurrence map construction. *Statistics and Computing*, 2011.
- [95] S.G. Potts, J.C. Biesmeijer, C. Kremen, P. Neumann, O. Schweiger, and W.E. Kunin. Global pollinator declines : trends, impacts and drivers. *Trends in Ecology and Evolution*, 25(6) :345–353, 2010.
- [96] W. J. M. Probert, C. E. Hauser, E. McDonald-Madden, M. C. Runge, P. W. J. Baxter, and H. P. Possingham. Managing and learning with multiple models : objectives and optimization algorithms. *Biological Conservation*, 144 :1237–1245, 2011.
- [97] M. L. Puterman. *Markov Decision Processes*. John Wiley and Sons, New York, 1994.
- [98] E. Rachelson, F. Schnitzler, L. Wehenkel, and D. Ernst. Optimal sample selection for batch-mode reinforcement learning. In *Proceedings of the International Conference on Agent and Artificial Intelligence*, Rome, Italy, 2011.
- [99] J. Radozsycki, N. Peyrard, and R. Sabbadin. Algorithme VBEM pour le processus de cox log gaussien. In *Actes des Journées de Statistique*, Toulouse, France, 2013.
- [100] T.J. Regan, I. Chadès, and H.P. Possingham. Optimally managing under imperfect detection : a method for plant invasions. *Journal of Applied Ecology*, 48(1) :76–85, 2011.
- [101] C. Robert and Casella G. *Monte Carlo Statistical Methods*. Springer texts in Statistics. Springer, 2004.

- [102] H. Rue, S. Martino, and N. Chopin. Approximate bayesian inference for latent gaussian models by using integrated nested Laplace approximations. *Journal of the Royal Statistical Society Series B*, 71(2) :319–392, 1999.
- [103] R. Sabbadin, N. Peyrard, and N. Forsell. A framework and a mean-field algorithm for the local control of spatial processes. *International Journal of Approximate Reasoning*, 53(1) :66–86, 2012.
- [104] M. Salathé and J. Jones. Dynamics and control of diseases in networks with community structure. *PLoS Computational Biology*, 6(4), 2010.
- [105] T. Schiex, H. Fargier, and G. Verfaillie. Valued constraint satisfaction problems : hard and easy problems. In *Proceedings of the International Joint Conference on Artificial Intelligence*, Montreal, Canada, 1995.
- [106] R. Snyder and R. Nisbet. Spatial structure and fluctuations in the contact process and related models. *Bulletin of Mathematical Biology*, 62 :959–975, 2000.
- [107] R. St-Aubin, J. Hoey, A. Hu, and C. Boutilier. APRICODD : Approximate policy construction using decision diagrams. In *Proceedings of the Conference on Neural Information Processing Systems*, Denver, USA, 2000.
- [108] R. S. Sutton and A.G. Barto. *Reinforcement Learning : An Introduction*. MIT Press, 1998.
- [109] G. Thébaud, N. Peyrard, S. Dallot, Calonnec A., and G. Labonne. Investigating disease spread between two assessment dates with permutation tests on a lattice. *Phytopathology*, 95 :1453–1461, 2005.
- [110] S. Thompson and G. Seber. *Adaptive sampling*. Series in Probability and Statistics. Wiley, New York, 1996.
- [111] M. Toussaint and A. Storkey. Probabilistic inference for solving discrete and continuous state Markov decision processes. In *Proceedings of the International Conference on Machine Learning*, 2006.
- [112] M. J. Wainwright and M. I. Jordan. Graphical models, exponential families, and variational inference. *Foundations and Trends in Machine Learning*, 1 :1–305, 2008.
- [113] S. Wasserman and K. Faust. *Social Network Analysis*. Cambridge University Press, 1994.
- [114] R. J. Williams and N. D. Martinez. Simple rules yield complex food webs. *Nature*, 404 :180–183, 2000.
- [115] J. Williamson. Maximising entropy efficiently. *Electronic Transactions in Artificial Intelligence*, 7, 2002.
- [116] J. Yedidia, W. Freeman, and Y. Weiss. Constructing free energy approximations and generalized belief propagation algorithms. *IEEE Transactions on Information Theory*, 51(7) :2282–2312, 2005.

Chapitre 5

Curriculum Vitæ détaillé

5.1 Etat civil

Nom : Peyrard Nathalie
Nationalité : française
Date de naissance : 04/11/1975
Situation familiale : mariée, 2 enfants

5.2 Parcours professionnel

actuellement : **CR1 INRA**, centre de Toulouse, Unité de Mathématique et Informatiques Appliquées de Toulouse (MIAT), équipe Modélisation des Agro-écosystèmes et Décision (MAD).
2006 - 2007 : **CR2 INRA**, centre de Toulouse, Unité de Mathématique et Informatiques Appliquées de Toulouse (MIAT), équipe Modélisation des Agro-écosystèmes et Décision (MAD).
2003 - 2006 : **CR2 INRA**, centre d'Avignon, Unité Biostatistique et Processus Spatiaux (BioSP).
2002 - 2003 : **Post-doctorante IRISA Rennes**, projet Vision Spatio-Temporelle et Apprentissage (VISTA)
1998 - 2001 : **Doctorante INRIA Rhône-Alpes**, projet IS2, et **Monitrice** à l'université J. Fourier, Grenoble 1.

5.3 Projets de recherche

▷ Projets nationaux

2014 - 2018 : **AgroBioSe**, projet ANR appel Agrobiosphère.
Titre : Biodiversité et services écosystémiques en agro-écosystèmes céréaliers intensifs : utilisation des concepts de l'agro-écologie pour atteindre les objectifs ECOPHYTO 2018.
Partenaires : CNRS-Chizé, UMR Agroécologie (INRA Dijon), Centre d'Ecologie Fonctionnelle et Evolutive (CNRS, Montpellier), MIAT (INRA Toulouse), Laboratoire de

Mathématiques d'Avignon, ONCFS, CNERA.

Rôle : leader du workpackage "Modelling and optimising sustainable weed management strategies".

2014 - 2016 : VESPA, projet Pour et Sur le Plan Ecophyto 2018.

Titre : Valeur et optimisation des dispositifs d'épidémiosurveillance dans une stratégie durable de protection des cultures.

Partenaires : UMR Agroécologie (INRA Dijon), MIAT (INRA Toulouse), Laboratoire d'Economie Appliquée de Grenoble, UMR Structure et Marchés Agricole, Ressources et Territoires (UMR INRA/Agrocampus Ouest, centre de Rennes), unité Sciences en Société (UR INRA IFRIS, Centre de Versailles).

Rôle : responsable du volet sur la conception de réseaux de surveillance des adventices.

2013 - 2016 : TakeControl, projet du métaprogramme INRA Sustainable Management of Crops Health (SMaCH).

Titre : Deployment strategies of plant quantitative resistance to take control of plant pathogen evolution.

Equipes partenaires : unités GAFL, PV et BioSP (INRA Avignon) et MIAT (INRA Toulouse).

Rôle : responsable scientifique de MIAT.

2013 : DyBRES, projet de l'appel à Idées 2012 du Réseau National des Systèmes Complexes (RNSC).

Titre : Dynamique de la Biodiversité sur des Réseaux Ecologiques Spatiaux.

Partenaires : Vincent Calgagno (INRA, ISA Sophia Antipolis), Frédéric Faure (Univ. Joseph Fourier, Institut Fourier, Grenoble), Annick Lesne (CNRS, LPTMC Paris), Stéphanie Manel (Univ. Marseille, AMAP Montpellier-LPED Marseille), François Massol (CNRS, CEFV Montpellier), François Munoz (Univ. Montpellier II, AMAP Montpellier), Sandrine Petit (UMR Agroécologie INRA Dijon), Estelle Pitard (CNRS, Laboratoire Charles Coulomb), Nathalie Peyrard (MIAT INRA Toulouse).

2010 - 2014 : Ficolof, projet ANR blanc.

Titre : Filtrage par cohérences locales fortes pour les réseaux de fonctions de coûts.

Equipes partenaires : GREYC (Univ. Caen), LIRMM (Univ. Montpellier 2), MIAT (INRA Toulouse).

2010 - 2014 : LARDONS, projet ANR blanc.

Titre : Learning And Reasoning for Deciding Optimally using Numerical and Symbolic information.

Equipes partenaires : GREYC (Univ. Caen), Lip6 (Paris 6), LAMSADE (Univ. Paris Dauphine), MIAT (INRA Toulouse).

2009 - 2010 : Projet Idées'08, projet de l'appel à Idées 2008 du Réseau National des Systèmes Complexes (RNSC).

Titre : Développement d'une méthode originale pour la résolution de processus décisionnels de Markov partiellement observables et spatialisés : application à la gestion de communautés d'adventices.

Partenaires : Iadine Chadès (CSIRO), Sabrina Gaba (UMR Agroécologie- INRA Dijon- EA), Nathalie Peyrard et Régis Sabbadin (MIAT - INRA - Toulouse).

2007 - 2008 : DECMIRI, projet du GdR ComEvol.

Titre : Dynamique, évolution et contrôle des maladies infectieuses sur réseau d'inter-

action.

Partenaires : Marc Choisy (GEMI - IRD, Montpellier), Alain Franc (BioGeCo - INRA - Bordeaux), Nathalie Peyrard et Régis Sabbadin (MIAT - INRA - Toulouse).

Rôle : responsable du projet.

2004 - 2005 : EmerGeom, projet de l'action transversale INRA "Epidémiologie et Risques Emergents".

Titre : Risque d'émergence d'une maladie et géométrie des interactions entre hôtes : apport des modèles mathématiques et de la théorie des graphes.

Partenaires : Alain Franc (BioGeCo - INRA - Bordeaux) et Nathalie Peyrard (MIAT - INRA - Toulouse).

▷ Projets internationaux

2013 - 2015 : MANAE, projet de collaboration entre MIAT et le Center of Excellence in Environmental Decision en Australie, soutenu par le département MIA de l'INRA.

Titre : Managing Networks in Agro-Ecology.

Partenaires : Iadine Chadès (CSIRO, Australia), Eve MacDonald-Madden, Will Probert and Tara Martin (CSIRO, University of Queensland, Australia), Nathalie Peyrard et Régis Sabbadin (MIAT - INRA - Toulouse).

2011 - 2015 : PURE, projet européen.

Titre : Pest Use Reduction in Europe, workpackage « Integrated Pest Management Design and Assessment Methodology ».

Collaborations principales : Jean-Noël Aubertot (AGIR - INRA - Toulouse) et Régis Sabbadin (MIAT - INRA - Toulouse).

2006 - 2009 : Projet ARC, projet de l'Australian Research Council.

Titre : Applying search theory for eradicating invasive species.

Partenaires : Barry Brook (The University of Adelaide – Research Institute for Climate Change and Sustainability, Australia), Ralph Mac Nally (Monash University, Australia), Nathalie Peyrard et Régis Sabbadin (MIAT - INRA - Toulouse), Danny Spring (Monash University, Australia).

5.4 Encadrement

▷ Encadrement de stages Bac +3 et Bac +4

2008 : Amine Rafik, Master 1 Informatique et Gestion de Toulouse 3.

Titre : Développement d'un simulateur de dynamiques épidémiques sur réseau. Application au contrôle de maladies infantiles dans une population structurée en classes d'âge.

Co-encadrants : Régis Sabbadin (MIAT - INRA - Toulouse) et Marc Choisy (GEMI - IRD, Montpellier).

2007 : Michel Laviron, 2ième année SupAero.

Titre : Gestion durable de résistances variétales à l'aide des processus décisionnels de Markov sur graphe : mise en œuvre informatique.

Co-encadrants : Régis Sabbadin (MIAT - INRA - Toulouse) et Jean-Noël Aubertot (AGIR - INRA - Toulouse).

2006 : Emilie-Anne Guerch, Master 1 Ingénierie Mathématique de Toulouse 3.

Titre : Gestion durable de résistances variétales à l'aide des processus décisionnels de Markov sur graphe : modélisation.

Co-encadrants : Régis Sabbadin (MIAT - INRA - Toulouse) et Jean-Noël Aubertot (AGIR - INRA - Toulouse).

2004 : **Benoit Allemand**, 3ième année IUP Génie Mathématique et Informatique d'Avignon.

Titre : Classification non supervisée de données géostatistiques.

Co-encadrant : Denis Allard (BioSP - INRA - Avignon).

▷ **Encadrement de stages Bac +5**

2013 : **Geoffroy Janvier**, Master 2 Professionnel Ingénierie Mathématique de Toulouse 3.

Titre : Réalisation d'une toolbox Matlab sur les Processus Décisionnels de Markov sur Graphe.

Co-encadrants : Régis Sabbadin et Marie-José Cros (MIAT - INRA - Toulouse).

2012 : **Julia Radoszycki**, Master 2 Recherche Probabilités et Statistiques de Toulouse 3, et 5ième année INSA.

Titre : Algorithme VBEM pour l'estimation du processus de Cox log gaussien.

Co-encadrant : Régis Sabbadin (MIAT - INRA - Toulouse).

Poursuite : Julia Radoszycki poursuit en thèse, co-encadrée par Régis Sabbadin, Sabrina Gaba et moi-même.

2012 : **Gabriel Sirvent**, Master 2 professionnel Intelligence Artificielle, Robotique et Reconnaissance des formes de Toulouse 3.

Titre : Optimisation de l'échantillonnage dans les Markov Logic Networks.

Co-encadrant : Régis Sabbadin (MIAT - INRA - Toulouse).

2009 : **Usman Farrok**, Master 2 Recherche Intelligence Artificielle de Toulouse 3.

Titre : Optimization of spatial sampling for map construction.

Co-encadrant : Régis Sabbadin (MIAT - INRA - Toulouse).

2009 : **Mathieu Bonneau**, Master 2 Professionnel Ingénierie Mathématique de Toulouse 3.

Titre : Échantillonnage spatial optimal basé sur les modèles booléens et le krigeage.

Co-encadrant : Régis Sabbadin (MIAT - INRA - Toulouse).

Poursuite : Mathieu Bonneau a poursuivi en thèse, co-encadré par Régis Sabbadin, Sabrina Gaba et moi-même.

2008 : **Benjamin Borgy**, Master 2 Recherche Biodiversité, Ecologie et Evolution de Toulouse 3.

Titre : Modélisation de stratégies de contrôle du phoma du colza dans le cadre des Processus Décisionnels de Markov sur Graphe.

Co-encadrants : Régis Sabbadin (MIAT - INRA - Toulouse) et Jean-Noël Aubertot (AGIR - INRA - Toulouse).

Poursuite : Benjamin Borgy a réalisé une thèse dans l'UMR Agroécologie de Dijon, la collaboration se poursuit autour de la modélisation de dynamiques adventices.

2007 : **Lamine Conde**, M2 Recherche Probabilités et Statistiques de Toulouse 3.

Titre : Extension du cadre des Processus Décisionnels de Markov sur Graphe pour l'optimisation de la structure de réseaux d'interaction dans les processus spatio-temporels.

Co-encadrant : Régis Sabbadin (MIAT - INRA - Toulouse).

▷ **Encadrement de thèses**

2012 - 2015 : **Julia Radoszycki**.

Titre : Méthodes de conception par optimisation de stratégies de gestion de phénomènes spatio-temporels : application à la gestion de communautés de plantes adventices.

Co-encadrants : Régis Sabbadin (MIAT - INRA - Toulouse) et Sabrina Gaba (UMR Agroécologie - INRA Dijon).

Financement : contrat doctoral, ED Mathématique, Informatique et Télécommunications de Toulouse (MITT).

2009 - 2012 : **Mathieu Bonneau**.

Titre : Echantillonnage adaptatif optimal dans les champs de Markov. Application à l'échantillonnage d'une espèce adventice.

Co-encadrants : Régis Sabbadin (MIAT - INRA - Toulouse) et Sabrina Gaba (UMR Agroécologie- INRA Dijon).

Financement : INRA.

Poursuite : Mathieu Bonneau est actuellement en post-doc dans le département Resource management and Geography de l'université de Melbourne (Australie).

▷ **Encadrement de post-doctorants**

2012 : **Will Probert**, 6 mois.

Titre : Optimal management of food-webs.

Co-encadrants : Eve MacDonald-Madden (CSIRO, Univ. of Queensland) et Régis Sabbadin (MIAT - INRA - Toulouse).

Financement : Australian Centre of Excellence in Environmental Decision.

2010 : **Alana Moore**, 5 mois.

Titre : Control of competitive contact processes on graphs : application to the control of interacting invasive species on a network of sites.

Co-encadrants : Mickael MacCarthy (Univ. of Melbourne, Australian Centre of Excellence for Risk Analysis, Australia) et Régis Sabbadin (MIAT - INRA - Toulouse).

Financement : INRA / Australian Centre of Excellence for Risk Analysis.

2013 - 2015 : (en cours de recrutement)

Titre : Design of mathematics and artificial intelligence tools for the mapping and sustainable management of crops pests using autonomous UAVs.

Co-encadrants : Florent Teichtel (ONERA) et Régis Sabbadin (MIAT - INRA - Toulouse).

Financement : ONERA - Région Midi-Pyrénées.

5.5 Visites dans laboratoires étrangers

1999 et 2000 : **University of Washington - USA**, 2 séjours de 3 mois au sein du département de statistique.

Sujet : Champs de Markov cachés pour l'analyse d'images IRM.

Collaborateurs : Adrian Raftery (University of Washington - USA) et Florence Forbes

(MISTIS - INRIA Rhône-Alpes).

2005 : **IIASA - Autriche**, séjour de 3 semaines au sein du projet Evolution and Ecology Program.

Sujet : Méthodes variationnelles d'ordre 2 pour l'inférence des états d'équilibre et transients d'un modèle SIS sur graphe.

Collaborateurs : Ulf Dieckmann (Evolution and Ecology Program - IIASA - Autriche), Alain Franc (BioGeCo - INRA - Bordeaux).

2013 : **CSIRO/University of Queensland - Australie**, visite d'une semaine.

Sujet : Optimal management of food-webs

Collaborateurs : Iadine Chadès (CSIRO), Eve MacDonald-Madden (CSIRO, Univ. of Queensland) et Régis Sabadin (MIAT - INRA - Toulouse).

5.6 Animation scientifique

▷ Organisation de workshops internationaux

2008 : co-organisation de l'édition 2008 du workshop **Spatial Statistics and Image Analysis in Biology** (SSIAB) (Toulouse, mai 2008).

2010 : co-organisation du workshop **Graphical models for reasoning on biological systems : computational challenges**, satellite à la European Conference on Complex Systems (ECCS) (Lisbonne, septembre 2010).

2012 : co-organisation du workshop **Algorithmic Issues for Inference in Graphical Models (AIGM)**, satellite à la European Conference on Artificial Intelligence (ECAI) (Montpellier, août 2012).

2013 : co-organisation de l'édition 2013 du workshop **AIGM** (Paris, septembre 2013).

▷ Membre de comités de programme nationaux

2012 : membre du comité de programme des **IV journées francophones des sciences de la conservation** (Dijon, mai 2012).

▷ Relecture d'articles

Article dans revue scientifique : Journal de la Société Française de Statistique, Bernoulli, Statistics and Computing, IEEE Transactions on Image Processing, Transactions on Pattern Analysis and Machine Intelligence, Computer Vision and Image Understanding, Journal of Electronic Imaging, Signal Processing, Theoretical Population Biology.

Article dans conférence : AAAI Conference on Artificial Intelligence 2010 et 2011, Uncertainty and Artificial Intelligence 2010, 2011, 2012, 2013, International Joint Conference on Artificial Intelligence 2013.

Article dans workshop : Counting 2008.

▷ Evaluation de projets ANR

2008 : 1 projet ANR non thématique blanc.

2010 : 1 projet ANR du programme blanc SVSE 6.

▷ **Animation de groupe de travail**

2006 - ... : animatrice principale du **réseau MSTGA** (Modélisation Spatio-Temporelle sur Graphe et Approximations, <http://carlit.toulouse.inra.fr/MSTGA/>). Il s'agit d'un des réseaux méthodologiques du département MIA, auquel participent des chercheurs INRA, INRIA et universitaires. Ce réseau est financé par MIA et par le Réseau National des Systèmes Complexes (RNSC).

▷ **Jury de recrutement**

2013 : concours CR2 INRA en Mathématique et Modélisation.

2012 : concours IR INRA en Calcul Scientifique.

2012 : concours Maître de Conférence AgroParisTech en Statistique Computationnelle.

2009 : concours CR2 INRA en Modélisation et Bio-informatique.

▷ **Jury et comité de thèse**

2013 - ... : membre du comité de thèse de Julien Stoeher. Thèse intitulée “ Choix de modèles pour les champs de Gibbs cachés”, et réalisée sous la direction de Jean-Michel Marin (Université Montpellier 2), Lionel Cucala (Université Montpellier 2) et Pierre Pudlo (Université Montpellier 2/INRA).

2013 - ... : membre du comité de thèse de Julie Aubert. Thèse portant sur l'analyse statistique de données métagénomiques, et réalisée sous la direction de Stéphane Robin (INRA MIA) et Sophie Schbath (INRA -MIA).

2012 : : examinatrice de la thèse de Steven Volant. Thèse intitulée “Modélisation semi-paramétrique sous dépendance markovienne” et réalisée sous la direction de Stéphane Robin (Unité AgroParisTech/INRA - Paris) et M.-L. Martin-Magniette (Unité AgroParisTech/INRA et URGV - Paris).

2009 -2011 : membre du comité de thèse de Lamiae Azizi. Thèse intitulée “Champs aléatoires de Markov cachés pour la cartographie du risque en épidémiologie”, et réalisée sous la direction de Florence Forbes (INRIA Rhône-Alpes) et Myriam Garrido (Unité d'épidémiologie animale - INRA - Clermont-Ferrand).

▷ **Réflexives**

Par ailleurs, de 2004 à 2009 j'étais membre du groupe des co-animateurs des séminaires “Réflexives”, organisés par Marie-Claude Roland à l'INRA. Ces séminaires ont pour objectifs d'aider les étudiants en thèse à construire, situer et présenter leur sujet de thèse. Ils reposent sur une animation originale autour de binômes encadrant/docteurant.

▷ **Expertise**

2012 - ... : membre du groupe VTH (Variétés Tolérantes aux Herbicides) du comité de surveillance biologique du territoire.

5.7 Contribution au fonctionnement de collectifs INRA

2012 - ... : membre élue du Conseil de Service de l'unité MIAT.

2012 - ... : membre élue du Conseil de Gestion du département MIA.

2006 - 2011 : membre élue du Conseil Scientifique du département MIA.

5.8 Activités d'enseignement

2012 - 2013 : INSA 5ième année, filière Génie Mathématique et Modélisation.

Cours, TD et TP sur les champs de Markov (12h).

2003 - 2004 : IUP Génie Mathématique et Informatique d'Avignon. Cours,

TD et TP de Probabilités et Simulation Stochastique (50 h par an).

1998 - 2001 : Deug Science de la vie, Science de la terre de Grenoble I. Cours

et TD de Statistiques Descriptives et Probabilités (64 h par an).

Chapitre 6

Liste des publications

6.1 Articles scientifiques

Articles publiés

- (1) G. Celeux, F. Forbes, N. Peyrard, EM Procedures Using Mean Field-Like Approximations for Markov Model-Based Image Segmentation, *Pattern Recognition*, 36, pages 131-144, 2003.
- (2) F. Forbes, N. Peyrard, Hidden Markov Models Selection Criteria based on Mean Field-like approximations, *Transactions on Pattern Analysis and Machine Intelligence*, 25(9), pages 1089-1101, 2003.
- (3) N. Peyrard, P. Bouthemy, Motion-based selection of relevant video segments for video summarization, *Multimedia Tools and Applications Journal, Special Issue : Video Segmentation for Semantic Annotation and Transcoding*, 26, pages 255-274, 2005.
- (4) N. Peyrard, A. Calonnec, F. Bonnot, J. Chadœuf, Explorer un jeu de données sur grille par tests de permutation, *Revue de Statistique Appliquée*, 53(1), pages 59-78, 2005.
- (5) N. Peyrard and A. Franc, Cluster variation approximations for a contact process living on a graph, *Physica A*, 358, pages 575-592, 2005.
- (6) G. Thébaud, N. Peyrard, S. Dallot, A. Calonnec, G. Labonne, Investigating disease spread between two assessment dates with permutation tests on a lattice, *Phytopathology*, 95, pages 1453-1461, 2005.
- (7) F. Forbes, N. Peyrard, C. Fraley, D. Georgian-Smith, D. Goldhaber, A. Raftery, Model-Based Region-Of-Interest Selection in Dynamic Breast MRI, *Journal of Computer Assisted Tomography*, 30(4), pages 576-687, 2006.
- (8) N. Peyrard, F. Pellegrin, J. Chadœuf, D. Nandris, Statistical analysis of the spatio-temporal dynamics of rubber Bark Necrosis : no evidence of pathogen transmission, *Forest Pathology*, 36, pages 360-371, 2006.
- (9) N. Peyrard, U. Dieckmann, A. Franc, Long-range correlations improve understanding the influence of network structure on per contact dynamics, *Theoretical Population Biology*, 73(3), pages 383-394, 2008.
- (10) G. Thébaud, N. Peyrard, S. Dallot, A. Calonnec, G. Labonne, J. Chadœuf, Testing the spatial association of disease patterns between two dates in orchards. *Acta Horticulturae*, 781, pages 255-260, 2008.

- (11) A. Franc, M. Goulard, N. Peyrard, Chordal graphs to identify graphical models solutions of maximum of entropy under constraints on marginals, *SIAM Discrete Mathematics*, 24(3), pages 1104 - 1116, 2010.
- (12) N. Peyrard, R. Sabbadin, D. Spring, B. Brook, R. Mac Nally. Model-based adaptive spatial sampling for fire ants invasion map construction, *Statistics and Computing*, 2011 (version papier 23(1), pages 29-42, 2013).
- (13) R. Sabbadin, N. Peyrard, N. Forsell, A framework and algorithms for the local control of spatial processes, *International Journal of Approximate Reasoning*, 53(1), pages 66-86, 2012.
- (14) M. Charras-Garrido, D. Abrial, N. Peyrard, S. Dachian, New classification method for disease mapping based on discrete Hidden Markov Random Fields, *Biostatistics*, 13, pages 241-255, 2012.
- (15) M. Charras-Garrido, L. Azizi, F. Forbes, S. Doyle, N. Peyrard, D. Abrial, On the difficulty to delimit disease risk hot spots, *International Journal of Applied Earth Observation and Geoinformation*, in press.
- (16) Publication commune équipe MAD, Décision dans les agro-écosystèmes, *Revue d'Intelligence Artificielle*.

Articles en révision

- (17) M. Bonneau, N. Peyrard, R. Sabbadin, Reinforcement learning based design of sampling strategies under cost constraints in Markov Random Fields, *Computational statistics and Data Analysis*.
- (18) P. Tixier, J.-N. Aubertot, G. Caron-Lormier, S. Gaba, N. Peyrard, J. Radoszycki, R. Sabbadin, Modelling interaction networks for managing agricultural services, Thematic Volume in *Advances in Ecological Research*.

Articles en préparation

- (19) M. Bonneau, S. Gaba, N. Peyrard, R. Sabbadin, How to characterise the weed spatial structure using density classes? *Ecological Modelling*.
- (20) B. Borgy, S. Gaba, N. Peyrard, R. Sabbadin, X. Reboud, Weeds dynamics buried in the seed bank : the use of hidden Markov model to predict life history traits. *Revue non encore déterminée*.
- (21) W. Probert, E. MacDonald-Madden, N. Peyrard, R. Sabbadin, Managing ecological food webs. *Revue non encore déterminée*.

6.2 Chapitres d'ouvrages

- (22) A. Franc, N. Peyrard, B. Roche, *Propagation d'agents pathogènes dans les réseaux*. Dans : Introduction à l'épidémiologie quantitative des maladies infectieuses. Guégan J.F. et Choisy M. eds, 2009.
- (23) A. Franc, N. Peyrard, *Rôle de la géométrie du réseau d'interaction dans l'émergence d'une maladie*. Dans : Les maladies émergentes chez le végétal, l'animal et l'homme. Enjeux et stratégies d'analyse épidémiologique, Editions QUAE, 2010.

6.3 Conférences internationales avec sélection sur article long

Articles publiés

- (24) G. Celeux, F. Forbes, N. Peyrard, Mean field approximation principle and Markov random field model-based image segmentation. *Second European Conference on Highly Structured Stochastic Systems*, Pavie - Italie, septembre 1999.
- (25) G. Celeux, F. Forbes, N. Peyrard, Mean field approximation principle for parameter estimation in hidden Markov models. *First European Conference on Spatial and Computational Statistics*, Ambleside - Royaume Uni, septembre 2000.
- (26) N. Peyrard, P. Bouthemy, Content-based video segmentation using statistical motion models, *British Machine Vision Conference (BMVC'02)*, Cardiff - Royaume Uni, septembre 2002.
- (27) N. Peyrard, P. Bouthemy, Motion-based selection of relevant video segments for video summarization, *International Conference on Multimedia and Expo*, Baltimore - USA, juillet 2003.
- (28) N. Peyrard, P. Bouthemy, Detection of meaningful events in videos based on a supervised classification approach, *International Conference of Image Processing*, Barcelone - Espagne, septembre 2003.
- (29) N. Peyrard, R. Sabbadin, Mean Field Approximation of the Policy Iteration Algorithm for Graph-based Markov Decision Processes, *European Conference on Artificial Intelligence (ECAI'06)*, Trentino - Italie, août 2006.
- (30) N. Peyrard, R. Sabbadin, E. Lô-Pelzer, J. N. Aubertot, A Graph-based Markov Decision Process framework applied to the optimization of strategies for integrated management of diseases, *American Phytopathological Society and Society of Nematologist joint meeting*, San Diego - USA, juillet 2007.
- (31) M. Choisy, N. Peyrard, R. Sabbadin, A probabilistic decision framework to optimize the dynamics of a network evolution : application to the control of childhood diseases, *European Conference on Complex Systems (ECCS'07)*, Dresde - Allemagne, octobre 2007.
- (32) N. Peyrard, R. Sabbadin, J. N. Aubertot, A Markov Decision Process for integrated pest management, *International Congress on Modelling and Simulation (MODSIM'07)*, Christchurch - Nouvelle Zélande, décembre 2007.
- (33) N. Peyrard, R. Sabbadin, U. Farrokh Niaz, Decision-theoretic Optimal Sampling with Hidden Markov Random Fields, *European Conference on Artificial Intelligence (ECAI'10)*, Lisbon - Portugal, août 2010.
- (34) N. Peyrard, R. Sabbadin, D. Spring, R. Mac Nally, B. Brook. Spatial sampling in HMRF mapping problems : static and adaptive algorithms, *European Conference on Complex Systems (ECCS'10)*, Lisbonne - Portugal, septembre 2010.
- (35) M. Bonneau, N. Peyrard, R. Sabbadin. A Reinforcement-Learning Algorithm for Sampling Design in Markov Random Fields, *European Conference on Artificial Intelligence (ECAI'12)*, Montpellier - France, 2012.
- (36) M. Bonneau, N. Peyrard, R. Sabbadin. Reinforcement-learning for sampling design in Markov random fields. *International Conference on Computational Statistics (COMPSTAT'12)*, Limassol - Cyprus, 2012.

Article soumis

- (37) R. Sabbadin, N. Peyrard. Optimal Information Gathering Problems in Markov Logic Networks. *International Conference on Data Mining*, Dallas - USA, 2013.

6.4 Conférences nationales ou francophones avec sélection sur article long

Articles publiés

- (38) G. Celeux, F. Forbes, N. Peyrard, Approximation du champ moyen et segmentation d'images. *Journées de Statistique*, Grenoble - France, mai 1999.
- (39) G. Celeux, F. Forbes, N. Peyrard, Critères pour la sélection d'un modèle de champ de Markov caché. *Journées de Statistique*, Fès - Maroc, mai 2000.
- (40) G. Piriou, N. Peyrard, P. Bouthemy, J.-F. Yao, Détection d'événements dans les vidéos à l'aide de modèles probabilistes de mouvement. *Traitement et Analyse de l'Information Méthodes et Applications (TAIMA'03)*, Hammamet - Tunisie, septembre 2003.
- (41) G. Celeux, F. Forbes, N. Peyrard, Modèle de Potts avec champ externe et algorithme EM pour la segmentation d'image, *Congrès Francophone de Reconnaissance des Formes et Intelligence Artificielle (RFIA'04)*, Toulouse - France, janvier 2004.
- (42) N. Peyrard, D. Allard, F. Forbes, Comparaison de deux modélisations pour la classification de données géostatistiques, *Journées de Statistique*, Pau - France, juin 2005.
- (43) N. Peyrard, R. Sabbadin, Itération de la politique approchée pour Processus Décisionnels de Markov sur graphe, *Congrès Francophone de Reconnaissance des Formes et Intelligence Artificielle (RFIA'06)*, Tours - France, janvier 2006.
- (44) M. Bonneau, N. Peyrard, R. Sabbadin, Echantillonnage spatial basé sur le krigeage pour la reconstruction de carte d'occurrence, *Congrès francophone de Reconnaissance des Formes et Intelligence Artificielle (RFIA'10)*, Caen - France, janvier 2010.
- (45) M. Bonneau, N. Peyrard, R. Sabbadin, Echantillonnage spatial basé sur le krigeage pour la reconstruction de cartes d'occurrence, *Conférence de la société Française de Recherche Opérationnelle et Aide à la Décision (ROADEF'10)*, Toulouse - France, février 2010.
- (46) M. Bonneau, N. Peyrard, R. Sabbadin. Echantillonnage adaptatif optimal dans les champs de Markov discrets. *Journées Francophones de Planification, Décision et Apprentissage pour la conduite de systèmes (JFPDA'11)*, Rouen - France, juin 2011.
- (47) J. Radozsycki, N. Peyrard, R. Sabbadin, Algorithme VBEM pour le processus de Cox log gaussien. *Journées de Statistique*, Toulouse - France, mai 2013.

6.5 Workshops, colloques, congrès (sélection sur résumé ou sans sélection)

Note : sont listées ici uniquement les communications ne correspondant à aucune publication dans une revue ou une conférence.

- (48) M. Bonneau, S. Gaba, N. Peyrard, R. Sabbadin. An adaptive sampling method using hidden Markov random fields for the mapping of weeds at the field scale. *Workshop of the EWRS Working Group on Weeds and Biodiversity*, Dijon - France, mars 2011.
- (49) M. Bonneau, S. Gaba, N. Peyrard, R. Sabbadin. Weeds Sampling for Map Reconstruction : a Markov Random Field Approach. *French-Danish workshop in Spatial Statistics and Image Analysis in Biology (SSIAB'12)*, Avignon - France, mai 2012.
- (50) M. Bonneau, S. Gaba, N. Peyrard, R. Sabbadin. Échantillonnage d'une espèce adventice sur une parcelle cultivée pour la cartographie de classes d'abondance. *Journées francophones des sciences de la conservation*, Dijon - France, mai 2012.
- (51) B. Borgy, S. Gaba, N. Peyrard, R. Sabbadin, X. Reboud. Identification des traits d'histoire de vie des espèces adventices par l'utilisation d'un modèle de Markov caché et des séries de flore levée. *Journées francophones des sciences de la conservation*, Dijon - France, mai 2012.
- (52) W. Probert, E. McDonald-Madden, N. Peyrard, R. Sabbadin. Computational issues surrounding the dynamic optimisation of management of an ecological food web. *ECAI workshop on algorithmic issues for inference in graphical models (AIGM'12)*, Montpellier - France, août 2012.
- (53) A. Franc, N. Peyrard, R. Sabbadin, T. Schiex, An example of a link between artificial intelligence and statistical physics : exact computation of the partition function of a graphical model on a series-parallel graph. *Workshop on algorithmic issues for inference in graphical models (AIGM'13)*, Paris - France, septembre 2013.
- (54) J. Radoszycki, N. Peyrard and R. Sabbadin, Evaluation of stochastic policies for factored Markov decision processes. *Workshop on algorithmic issues for inference in graphical models (AIGM'13)*, Paris - France, septembre 2013.

6.6 Mémoires et rapports de recherche

- (55) N. Peyrard, Algorithme SEM pour l'estimation paramétrique d'un processus booléen de segments, mémoire de DEA, 1998.
- (56) N. Peyrard, Convergence of MCEM and Related Algorithms for Hidden Markov Random Field, Rapport de Recherche Inria, no. 4146, mars 2001.
- (57) N. Peyrard, Approximations de type champ moyen des modèles de champ de Markov pour la segmentation de données spatiales, thèse de l'Université J. Fourier, octobre 2001.
- (58) N. Peyrard, R. Sabbadin, Evaluation of the expected epidemic size of a spatial SIR process, Rapport de recherche de l'unité MIAT, RR-2012-1, 2012.