



SPIM

Thèse de Doctorat



UFC

école doctorale **sciences pour l'ingénieur et microtechniques**
UNIVERSITÉ DE FRANCHE-COMTÉ

MODÉLISATION DYNAMIQUE DE LA DENSITÉ DE POPULATION VIA LES RÉSEAUX CELLULAIRES ET OPTIMISATION MULTIOBJECTIF DE L'AUTO-PARTAGE

 **Laurent MOALIC**

Soutenue le : 12/12/2013

A Belfort

Jury

Rapporteur

TALBI El-Ghazali, Professeur, Université de Lille

Rapporteur

NOEL Thomas, Professeur, Université de Strasbourg

Examineur

CUNG Van Dat, Professeur, INP Grenoble

Directeurs de thèse

SPIES François, Professeur, UFC

CAMINADA Alexandre, Professeur, UTBM

LAMROUS Sid, Maître de Conférence, UTBM

Remerciements

Je tiens à remercier ici les personnes qui ont influencées, de près ou de loin, le travail présenté dans cette thèse.

Merci tout d'abord aux rapporteurs de ce travail, El-Ghazali Talbi et Thomas Noël, pour leur lecture attentive. L'exercice est certainement délicat mais ils ont accepté d'apporter leur éclairage avisé sur ce manuscrit. Merci également à Van Dat Cung, examinateur, qui a bien voulu présider le jury.

Merci à l'équipe qui m'a accompagné au cours de ces années de thèse et bien au-delà. Alexandre Caminada, qui a toujours su être là pour m'épauler et éclairer mon travail scientifique. Sid Lamrous, avec qui les moments partagés auront été très enrichissants. J'adresse également mes profonds remerciements à François Spies, qui m'a donné l'opportunité de réaliser cette thèse.

Merci à mes collègues, avec qui les discussions nombreuses sont toujours fructueuses. Jean-Noël Martin, qui bien plus qu'un voisin de bureau aura été une source d'inspiration. Oumaya Baala, Frederic Lassabe, Mohammad Dib, Jun Hu et bien d'autres... merci à eux. Enfin, merci à Alexandre Gondran pour son enthousiasme permanent, pour les moments partagés et les nouvelles pistes à explorer...

Table des matières

Introduction générale	1
I Modélisation dynamique de la densité de population via les réseaux cellulaires	7
1 Modélisation de la mobilité, état de l'art	9
1.1 Introduction	10
1.2 Taxinomie des modèles	11
1.2.1 Modèles de mobilité aléatoires	12
1.2.2 Modèles de mobilité déterministes	16
1.2.3 Modèles de mobilité hybrides	17
1.3 Cas de la téléphonie comme marqueur de mobilité	27
1.3.1 Le projet CAPITAL [Systems 1997]	29
1.3.2 Les travaux du MIT : une source de référence	29
1.3.3 Autres travaux utilisant la téléphonie mobile	33
1.4 Synthèse et perspectives	35
2 La téléphonie mobile : un marqueur de présence	39
2.1 Introduction	40
2.2 La téléphonie mobile : une nouvelle source de données	41
2.2.1 Structure et fonctionnement d'un réseau cellulaire	41
2.2.2 Quelles données sont disponibles ?	45
2.3 Un nouveau modèle de distribution de la population	47
2.3.1 Mobilité et caractérisation socio-économique du territoire	48

2.3.2	Modèle de distribution spatio-temporelle des mobiles	55
2.4	Résultats et analyse	58
2.4.1	Environnement d'analyse	58
2.4.2	Constitution des données de référence	60
2.4.3	Expérimentation et résultats	64
2.4.4	Construction d'une <i>Signature Territoriale</i>	70
2.4.5	Généralisation de la <i>Signature Territoriale</i>	73
2.5	Synthèse	73
II	Optimisation combinatoire multiobjectif de l'auto-partage	77
3	Optimisation combinatoire multiobjectif : état de l'art	79
3.1	Introduction	80
3.2	Notions de base et définitions	81
3.2.1	Optimisation mono-objectif	81
3.2.2	Formulation d'un problème d'optimisation multiobjectif . . .	82
3.2.3	Espace de décision vs espace objectif	82
3.2.4	Optimisation combinatoire et métaheuristiques	83
3.2.5	Dominance et Pareto optimalité (le problème de la relation d'ordre)	85
3.3	Métaheuristiques pour les problèmes multiobjectifs	87
3.3.1	Classification des principales approches	87
3.3.2	Les approches par transformation mono-objectif	87
3.3.3	Les approches non-Pareto	92
3.3.4	Les approches Pareto	93
3.3.5	Les approches hybrides	101

3.4	Evaluation des algorithmes multiobjectifs : une approche à base d'indicateurs	103
3.4.1	Indicateur de contribution	104
3.4.2	Indicateur ε -additif	105
3.4.3	Indicateur hypervolume	105
3.4.4	Indicateur de distance générationnelle	108
3.5	Synthèse	108
4	Algorithmes mémétiques et optimisation multiobjectif	111
4.1	Introduction	112
4.2	Schéma directeur d'un algorithme mémétique	113
4.2.1	La recherche locale dans un contexte multiobjectif	114
4.2.2	L'opérateur génétique dans un contexte multiobjectif	115
4.2.3	L'algorithme mémétique, une recherche locale améliorée ?	115
4.3	MEMO : un algorithme mémétique multiobjectif	117
4.3.1	NSGA-II comme opérateur génétique	117
4.3.2	L'opérateur de recherche locale	121
4.3.3	Déroulement de l'algorithme mémétique	123
4.4	Optimisation appliquée à la planification urbaine : localisation de stations d'auto-partage	125
4.4.1	Introduction	125
4.4.2	Enjeu de la localisation des stations	126
4.4.3	Description formelle du problème	127
4.5	Application de l'algorithme mémétique MEMO	134
4.5.1	Introduction des indicateurs de comparaison	135
4.5.2	Définition d'un ensemble de référence avec NSGA-II, PLS et FLSMO	137

4.5.3	Définition du paramétrage de MEMO	149
4.5.4	Analyse des solutions de MEMO	152
4.5.5	Analyse de la progression de l'algorithme MEMO	157
4.6	De l'optimisation à l'aide à la décision multiobjectif	158
4.6.1	Contexte de l'aide à la décision	158
4.6.2	La plate-forme geoLogic	159
4.6.3	Aide à la décision multiobjectif	161
4.6.4	Les projets s'appuyant sur la plate-forme geoLogic	164
4.7	Synthèse	169
Conclusion générale et perspectives		173
Liste des publications		177
Bibliographie		179

Liste des tableaux

2.1	Temps moyen passé par catégorie	53
2.2	Répartition de la population par thème	53
4.1	Rôles des opérateurs selon le type d'algorithme	116
4.2	NSGA-II : croisement aléatoire et mutation aléatoire	139
4.3	NSGA-II : croisement aléatoire et mutation premier voisin améliorant	140
4.4	NSGA-II : croisement aléatoire et mutation meilleur voisin améliorant	141
4.5	NSGA-II : croisement élitiste et mutation aléatoire	142
4.6	NSGA-II : croisement élitiste et mutation premier voisin améliorant .	143
4.7	NSGA-II : croisement élitiste et mutation meilleur voisin améliorant	144
4.8	Résultats de la recherche locale PLS	146
4.9	Résultats de la recherche locale FLSMO	147
4.10	Résultats de MEMO avec croisement aléatoire	151
4.11	Résultats de MEMO avec croisement élitiste	152
4.12	Synthèse des résultats obtenus avec les algorithmes présentés	153

Table des figures

1.1	Taxinomie des modèles de mobilité	11
1.2	Exemple de mouvements obtenus avec le modèle "Random Walk" . .	13
1.3	Diagramme d'état du modèle de chaîne de Markov à une dimension (extrait de [Bar-Noy 1994])	14
1.4	Trois trajectoires de 2000 pas du modèle "Normal Walk" pour σ $=15^\circ, 30^\circ$ et 60° (extrait de [Tuan 2003])	15
1.5	Exemple de motif de déplacement d'un mobile utilisant le modèle de mobilité Gauss-Markov (extrait de [Camp 2002])	16
1.6	Séquence de cellules d'un profil usager (extrait de [Erbas 2001]) . . .	19
1.7	Méthodes de détermination des chemins (extrait de [Kammann 2003])	20
1.8	MBMM : grille des probabilités de transition	21
1.9	Graphe du centre ville de Stuttgart (extrait de [Tian 2002])	27
1.10	Terrain simulé par un graphe de Voronoï (extrait de [Jardosh 2005])	28
1.11	Vue de Rome en temps réel	30
1.12	Illustration des vues de la mobilité à Singapour	32
1.13	Illustration du réseau de bus actuel (en jaune) et des préconisations de nouvelles lignes	34
2.1	Evolution du nombre d'abonnements téléphoniques dans le monde .	40
2.2	Maillage et PPC	43
2.3	Couverture radio	43
2.4	Génération d'un événement de communication	44
2.5	Génération d'un événement de transfert intercellulaire	45
2.6	Cartographie temporelle de la distribution de la population belfortaine	47

2.7	Process complet de la pondération spatio-temporelle du territoire . .	49
2.8	Organisation du sursol	50
2.9	Table de profilage temporel du sursol	51
2.10	Module de distribution des mobiles	55
2.11	Distribution des mobiles à un instant t sans prise en compte des propriétés du sursol	58
2.12	Distribution des mobiles à un instant t sans et avec prise en compte des poids d'attraction spatio-temporelle	59
2.13	Plate-forme geoLogic	60
2.14	Informations vectorielles : SIG shapefile	60
2.15	Zonage et attractivité des zones à 17h	61
2.16	Différence cumulée à 6h30, 10h et 22h	62
2.17	Présence à 3 périodes de temps	63
2.18	Comparaison EMD / modèle en zone d'activité industrielle sans habitat	65
2.19	Comparaison EMD / modèle en centre ville avec habitats et commerces	66
2.20	Comparaison EMD / modèle en secteur périurbain 1 avec habitat et très peu de commerce	67
2.21	Comparaison EMD / modèle en secteur périurbain 2 avec habitat et très peu de commerce	68
2.22	Symétrie du modèle en secteur périurbain 1 avec habitat et très peu de commerce	69
2.23	Symétrie du modèle en secteur périurbain 2 avec habitat et très peu de commerce	69
2.24	Signature Territoriale pour le centre ville (courbes hautes) et la zone d'activité industrielle (courbes basses) - pour deux mardis et jeudis différents	71
2.25	Signature Territoriale pour les zones périurbaines - pour deux mardis et jeudis différents	72

3.1	Espace de décision / espace objectif	82
3.2	Pareto dominance et incomparabilité	86
3.3	Classification des méthodes de résolution d'un MOP	88
3.4	Approche par agrégation des objectifs	89
3.5	Approche par methode ε -contrainte	90
3.6	Approche par but	92
3.7	Distance de <i>crowding</i> : moyenne des côtés de l'hypercube	95
3.8	Indicateur hypervolume I_H	106
3.9	Décomposition d'hypervolume de dimension 3 par HSO (extrait de [While 2006])	108
4.1	Sélection par tournoi binaire	119
4.2	Croisement élitiste : chaque parent fournit ses meilleurs gènes	120
4.3	Sélection et croisement suivis d'une recherche locale	124
4.4	Illustration des sites pouvant accueillir une station	128
4.5	Stations et zone de capture	129
4.6	Implantation des stations selon l'attraction vs les flux	131
4.7	Equilibre des flux entrants / flux sortants à l'instant t pour la station s	132
4.8	Deux fonctionnements d'une station au cours de la journée	133
4.9	Problème multiobjectif de maximisation	135
4.10	Evolution des indicateurs hypervolumes moyens de 6 configurations NSGA-II	145
4.11	Parcours PLS vs FLSMO	148
4.12	Progression des algorithmes de référence	149
4.13	Premières itérations des algorithmes de référence	149
4.14	Comparaison de deux ensembles pour I_C	154

4.15 Solutions PLS et MEMO projetées sur f_1-f_2 ; f_1-f_3 et f_2-f_3	155
4.16 Solutions NSGA-II et MEMO projetées sur f_1-f_2 ; f_1-f_3 et f_2-f_3 . . .	156
4.17 Progression de PLS, NSGA-II et MEMO	157
4.18 Plate-forme geoLogic	159
4.19 Vues issues de la plate-fome geoLogic	160
4.20 Outil d'aide à la décision de geoLogic	162
4.21 Filtres sur l'espace des critères	163
4.22 Filtres sur l'espace de décision	163
4.23 Territoire Mobile : analyse d'un réseau de bus	165
4.24 ecoTaxi : optimisation d'infrastructures de stations d'échange automatique de batteries	166
4.25 PUMDP : adéquation entre offre et demande pour une mobilité à vélo	167

Liste des algorithmes

1	NSGA-II	95
2	PLS	99
3	HSO : Hypervolume by Slicing Objectives	107
4	Algorithme mémétique générique MOGLS	114
5	NSGA-II	118
6	Local Search : FIHC	122
7	MEMO	124
8	Fast Local Search for Multiobjective Problems	147

Introduction générale

Contexte

Ce travail de thèse s'est déroulé sous la tutelle du laboratoire FEMTO de l'Université de Franche-Comté.

Les problèmes de transport et la recherche opérationnelle (RO) constituent deux domaines souvent liés. De nombreux problèmes cités en référence au sein de la communauté RO concernent en effet des applications liées au transport, telles que les problèmes de tournées de véhicules ou encore les problèmes de logistique urbaine. Réciproquement, en matière de transport, de nombreuses décisions sont prises en s'appuyant sur des outils d'optimisation, tel que pour planifier l'affectation des bus et chauffeurs lors du graphissage de lignes de bus. Ce lien étroit s'explique essentiellement par deux aspects. D'une part, en matière de transport les problèmes sont souvent très complexes. Il y a généralement de nombreuses données à prendre en compte et l'espace de recherche est immense tant il y a de solutions potentiellement réalisables. C'est le cas par exemple avec les problèmes de tournées de véhicules où la taille de l'espace de recherche croît de façon exponentielle avec le nombre de sites à visiter. D'autre part, les problèmes de transport ont des conséquences financières considérables. Les investissements en jeu en matière de construction d'infrastructures routières par exemple, ou de déploiement de services de transport en commun, s'élèvent très rapidement à plusieurs millions d'euros.

De ce lien étroit entre transport et recherche opérationnelle nous identifions deux enjeux majeurs qui doivent permettre d'aboutir à une prise de décision. Le premier concerne l'identification des besoins, à quelle demande doit-on répondre. Le deuxième concerne la réponse à apporter, quelle solution doit-on retenir et mettre en oeuvre. Chacun de ces enjeux est abordé dans ce mémoire de thèse : le premier, à travers une nouvelle approche qui permet d'appréhender les besoins en matière de service de mobilité ; le deuxième à travers de nouveaux outils et algorithmes d'optimisation et d'aide à la décision.

Motivations et objectifs

L'élaboration de modèles permettant d'identifier la façon dont les personnes se déplacent sur un territoire doit constituer un élément de réponse essentiel au premier enjeu. En effet, des modèles à même de refléter les besoins en services de mobilité au cours de la journée et d'un jour à l'autre font partie intégrante du processus

de décision qui précède la réalisation d'aménagements divers. Des modèles existent pour aborder ce problème, tel que le modèle à quatre étapes ou encore des modèles dits désagrégés. Ils permettent de construire les grands flux de déplacements réalisés sur un territoire. Toutefois, ces modèles supposent de disposer de données d'entrée acquises à travers des campagnes de collectes, des comptages ou des enquêtes. Les enquêtes ménages/déplacements constituent actuellement souvent l'élément de référence utilisé par les collectivités pour prendre des décisions d'aménagement.

Les téléphones portables qui équipent aujourd'hui une très large majorité de personnes fournissent un nouveau moyen pour percevoir les déplacements de la population. Bien que les téléphones mobiles existent depuis un certain temps, seul un nombre limité de travaux récents tentent de les exploiter pour en tirer une information utile sur les déplacements des personnes. Il y a pourtant là une formidable source d'informations appuyée sur un échantillon de personnes jamais atteint dans les enquêtes. De plus, cette donnée continuellement renouvelée offre des possibilités d'actualisation des connaissances selon les besoins. Les réseaux de téléphonie mobile apportent en ce sens une réponse aux enjeux posés en matière de compréhension et de modélisation de la mobilité.

En effet, les enquêtes, très coûteuses, ne sont actualisées que très rarement (en général tous les 10 ans). De plus, elles sont basées sur un échantillon très réduit (environ 1% de la population), et fournissent une vue figée, un instantané des déplacements. La téléphonie mobile vient pallier ces trois problèmes : l'information est actualisée en permanence, l'échantillon disponible concerne la quasi-totalité de la population, et ces données fournissent une vue dynamique des déplacements, au cours de la journée mais aussi pour chaque jour de la semaine, pour tout type de semaine. Par contre, les données téléphoniques souffrent de deux inconvénients principaux : la faible fréquence des relevés d'événements et l'approximation géographique de la position. La fréquence des relevés est faible car les événements relevés par le réseau cellulaire ne sont pas générés en continu mais uniquement lorsque le mobile est utilisé (appel, SMS, etc.), ou lorsque le mobile se déplace (changement de zone de localisation). La localisation est approximée car, autant il est possible de savoir à quelle station radio est rattaché un téléphone, autant il n'est pas possible de savoir où le mobile se trouve dans la zone de couverture de l'antenne. Des modèles adéquats sont alors à inventer pour palier ces problèmes et bénéficier des avantages offerts.

Le deuxième enjeu que nous avons identifié concerne les problèmes de décision. Les problèmes rencontrés dans le domaine du transport, et plus généralement dans de nombreuses situations réelles, relèvent d'une complexité telle que des algorithmes exacts ne peuvent les résoudre dans un temps raisonnable. Le concept même de temps raisonnable est très variable en fonction du problème. Il peut aller de la seconde (voire moins) à plusieurs semaines. Dès lors qu'une approche exacte ne peut s'appliquer, les approches de type heuristique fournissent souvent une alternative intéressante. On distingue deux niveaux d'heuristiques selon leur aptitude à être

généralisées. Les heuristiques simples et les métaheuristiques. Les heuristiques simples sont dédiées à un problème et n'ont pas vocation à être génériques, même s'il est parfois possible de les adapter à d'autres types de problème. Les métaheuristiques définissent un cadre général qui leur permet d'être appliquées à plusieurs types de problème. Les algorithmes génétiques et les Recherches Tabou font partie de cette classe.

Par ailleurs, il est fréquent que les problèmes rassemblent plusieurs objectifs à optimiser simultanément, généralement contradictoires entre eux. L'optimisation multiobjectif s'adresse à ce type de problème en proposant différentes approches, par transformation du problème en un problème mono-objectif par exemple. Mais les algorithmes les plus efficaces prennent en considération toutes les fonctions objectif simultanément à travers les approches dites Pareto. Le résultat de ces approches est alors non pas une solution unique, qui soit la meilleure pour tous les objectifs, mais un ensemble de solutions de meilleur compromis. Une telle prise en compte de différents objectifs en même temps entraîne bien entendu une complexité accrue. Une des principales difficultés rencontrées avec ce type d'approches vient du fait que cet ensemble n'est pas connu lorsque l'on aborde un problème réel et c'est le cas du problème de transport que nous avons traité.

Dans le cadre de ce mémoire de thèse, nous nous intéressons à la chaîne complète qui permet de guider une décision multicritère dans le domaine de l'aménagement du territoire et du transport. Nous traitons ainsi les deux phases principales impliquées dans le processus de décision : la modélisation des déplacements de la population en s'appuyant sur la téléphonie mobile d'une part, et l'élaboration d'une métaheuristique hybride pour résoudre un problème d'optimisation multiobjectif d'autre part. Les traces que laissent les téléphones portables peuvent être vues comme des empreintes électromagnétiques. Nous proposons dans ce travail d'exploiter ces données pour construire une information sur la façon dont la population se répartit dans l'espace et dans le temps. Puis l'accent est mis sur la résolution du problème de la localisation de stations pour un service d'auto-partage électrique en contexte multicritère. Nous présentons en particulier la démarche qui nous a permis de modéliser le problème puis de concevoir une heuristique nouvelle, mêlant théorie évolutionnaire et recherche locale à travers l'élaboration d'un algorithme mémétique. Les deux phases principales de modélisation et d'optimisation ont été expérimentées et validées dans un contexte réel.

Organisation du mémoire

Ce mémoire de thèse est organisé en deux parties, organisées autour des deux enjeux présentés. Chaque partie est composée de deux chapitres que nous présentons brièvement.

Partie I Modélisation dynamique de la densité de population via les réseaux cellulaires.

Le premier chapitre est dédié à l'état de l'art sur les modèles de mobilité, et de la téléphonie mobile comme marqueur de mobilité. Nous y présentons tout d'abord une taxinomie des modèles issus de la littérature autour de trois familles : les modèles aléatoires, les modèles déterministes et les modèles hybrides. Les modèles hybrides sont eux-mêmes regroupés selon trois classes, les modèles agrégés, les modèles désagrégés, et enfin les modèles dits activité-centré. Ces derniers modèles qui s'appuient sur un chainage d'activités sont particulièrement adaptés à modéliser finement les besoins en mobilité. La téléphonie mobile est également présentée comme une nouvelle source de données permettant de marquer une présence sur un territoire. Nous présentons les travaux ayant attiré à ce domaine et qui laissent entrevoir le potentiel que peuvent apporter ces données.

Le deuxième chapitre présente le modèle de distribution spatio-temporelle des mobiles que nous proposons. Nous nous attachons dans un premier temps à présenter le fonctionnement d'un réseau cellulaire ainsi que les données pouvant servir à marquer la présence sur un territoire. Puis nous décrivons le fonctionnement du modèle proposé. L'accent est mis sur la façon dont la description du sursol est utilisée pour définir une attraction spatio-temporelle du territoire. Enfin, une phase de validation étayée sur un cas d'étude réel est présentée associant une collectivité et un opérateur téléphonique.

Partie II Optimisation combinatoire multiobjectif de l'auto-partage.

Le chapitre 3, premier chapitre de cette deuxième partie, décrit un état de l'art des métaheuristiques combinatoires permettant de résoudre des problèmes multiobjectifs. Dans un premier temps, nous présentons les principaux concepts permettant d'appréhender le domaine des problèmes multiobjectifs ; puis les principaux modèles sont décrits, organisés autour de quatre grandes familles. La première regroupe les approches par transformation du problème multiobjectif en un problème mono-objectif. La deuxième définit les approches dites non-Pareto, qui considèrent tous les objectifs mais séparément les uns des autres. La troisième famille, dite Pareto, regroupe les algorithmes qui s'appuient sur la relation de dominance pour construire des populations de solutions. La dernière famille consiste en une hybridation de mécanismes issus des familles précédentes. Ce chapitre se termine sur une présentation des mécanismes permettant d'évaluer des ensembles de solutions, en se basant sur des indicateurs de qualité.

Enfin, le dernier chapitre de ce mémoire de thèse présente un nouvel algorithme combinatoire hybride incorporant une recherche locale dans un schéma évolutionnaire en contexte multiobjectif. Les principales caractéristiques des algorithmes mémétiques sont dans un premier décrites. Puis nous développons l'algorithme proposé en montrant

comment les opérateurs génétiques et la recherche locale s'articulent. Une phase d'expérimentation est alors proposée pour résoudre le problème de localisation de stations pour un service d'auto-partage électrique, que nous avons formalisé sous forme d'un problème d'optimisation multiobjectif combinatoire. Les résultats obtenus sont comparés aux résultats issus de trois autres algorithmes. Enfin, la dernière partie de ce chapitre présente la plate-forme d'aide à la décision geoLogic que nous avons développée notamment pour le problème de l'auto-partage.

Ce mémoire se termine par une conclusion qui revient sur les principales contributions apportées dans cette thèse. Finalement, une ouverture est proposée avec la présentation de plusieurs perspectives qui nous semblent particulièrement intéressantes à poursuivre.

Première partie

Modélisation dynamique de la densité de population via les réseaux cellulaires

Modélisation de la mobilité, état de l'art

Ce chapitre permet de situer le contexte dans lequel s'inscrit cette première partie de notre mémoire de thèse, à savoir la modélisation dynamique de la mobilité. En effet, pour situer la contribution apportée dans le chapitre suivant, ce chapitre présente d'abord un échantillon représentatif de modèles existants sous forme d'une taxinomie : aléatoires, déterministes, ou hybrides. L'accent sera mis sur les modèles hybrides, qui apportent le meilleur niveau de réalisme et s'apparente de fait au modèle que nous proposerons dans le chapitre 2. Puis, sur la lignée de nos travaux, les récentes contributions de la littérature qui utilisent la téléphonie mobile comme marqueur de mobilité sont présentées.

Sommaire

1.1	Introduction	10
1.2	Taxinomie des modèles	11
1.2.1	Modèles de mobilité aléatoires	12
1.2.2	Modèles de mobilité déterministes	16
1.2.3	Modèles de mobilité hybrides	17
1.3	Cas de la téléphonie comme marqueur de mobilité	27
1.3.1	Le projet CAPITAL [Systems 1997]	29
1.3.2	Les travaux du MIT : une source de référence	29
1.3.3	Autres travaux utilisant la téléphonie mobile	33
1.4	Synthèse et perspectives	35

1.1 Introduction

Une connaissance précise de la façon dont la population se déplace sur un territoire apparaît aussi complexe à déterminer que primordiale. Les modèles de trafic et plus généralement de mobilité développés au cours de ces soixante dernières années reposent sur diverses théories, impliquant notamment la prise en compte des activités réalisées au cours de la journée ou encore l'occupation du sol.

A l'origine, les besoins en infrastructures très coûteuses, construites dans le cadre de projets à long terme, ont poussé les collectivités à disposer de prévisions de trafic. Selon les besoins, différents modèles ont été proposés. A l'échelle macroscopique, un flux entre deux villes (trafic interurbain) peut être estimé par un modèle dit "gravitaire" [PREDIT 1999]. Dès lors que l'on s'intéresse à du trafic intra-urbain, un niveau de complexité plus important est à intégrer. D'autre part, les priorités en termes d'urbanisation évoluent en fonction des contextes économiques, sociaux et environnementaux. Ces modèles ont dû être affinés et repensés pour répondre à cette complexité d'une urbanisation instable.

C'est tout d'abord aux États-Unis dans les années 1950, avec l'apparition de l'informatique, que les premiers modèles capables de représenter le trafic urbain sont élaborés [PREDIT 1999]. Apparaît alors une première génération de modèles, dits modèles "classiques" ou "à 4 étapes". Ils permettent de réaliser une prévision de trafic suivant une procédure linéaire découpée en quatre étapes : la génération, la distribution, le choix du mode et l'affectation. Ils sont calibrés pour une ville déterminée à partir de données de comptages. Par leur approche agrégée, ces modèles permettent des prévisions globales de trafic entre groupements de zones, notamment aux heures de pointe, afin de dimensionner les infrastructures et ainsi planifier les investissements.

Les besoins en planification à court terme ont donné naissance à une approche "désagrégée" basée de plus en plus sur l'individu ou sur le ménage. Des évolutions du modèle classique en ce sens ont donc été proposées. Mais pour aller plus loin, ces nouveaux modèles désagrégés également appelés "à choix discret", sont capables de prédire le mode de déplacement au niveau individuel. Ces derniers sont calibrés à partir d'enquêtes de plus en plus fines.

Dans ce chapitre, on se propose de parcourir les principaux modèles de mobilité rencontrés dans la littérature, qui en fonction de leurs caractéristiques permettent de répondre à des enjeux de nature diverse. Les modèles seront présentés autour d'une taxinomie en trois familles : les modèles aléatoires, les modèles déterministes, et les modèles hybrides. Alors que les modèles aléatoires sont généralement appréciés pour leur simplicité, dès lors que le niveau de réalisme souhaité devient un critère prépondérant les modèles hybrides sont préférés. Outre les modèles classiques, une tendance très récente vise à utiliser la téléphonie mobile pour en déduire une information

de répartition de la population sur un territoire. Bien qu'il n'y ait pas encore de modèles à proprement parler dans les travaux présentés dans la littérature, il y a là une nouvelle source de données à très fort potentiel. Aussi, nous introduirons à travers plusieurs projets les travaux ayant été effectués ces dernières années sur cette thématique. Cela permettra de positionner notre approche qui sera présentée dans le chapitre suivant.

Ce chapitre est organisé autour de deux sections principales. La première sera consacrée à l'énumération et l'explicitation des différents modèles rencontrés dans la littérature, des plus arbitraires où l'aléatoire joue un grand rôle, aux plus réalistes tenant compte du type de mobile, de la géographie ainsi que des activités des personnes. La deuxième section permettra de recenser de façon précise les travaux ayant porté sur l'utilisation de la téléphonie mobile pour en déduire une information sur la façon dont les utilisateurs de portables se déplacent.

1.2 Taxinomie des modèles

Les modèles existants présentent des caractéristiques bien distinctes. On peut identifier différentes familles de modèles selon que l'on s'attache davantage à simuler le déplacement des individus, ou bien à simuler les effets de groupe. De plus, tous les modèles n'accordent pas la même importance à la topologie du terrain. Cette section présente une taxinomie des différentes approches de modélisation selon le contexte d'étude (cf. Figure 1.1).

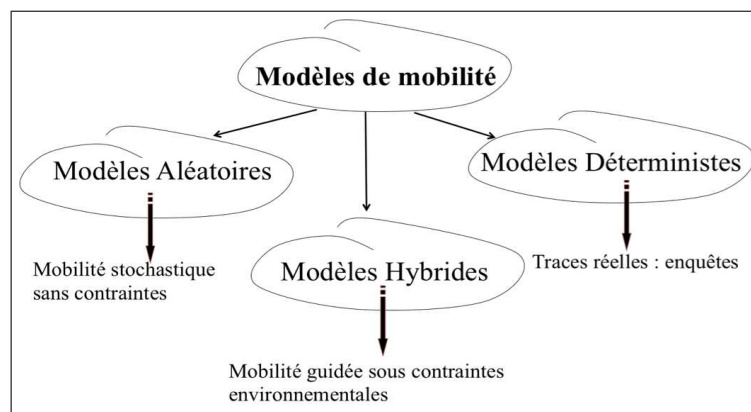


FIGURE 1.1 – Taxinomie des modèles de mobilité

Les modèles de mobilité aléatoires simulent des mouvements et directions arbitraires sur une zone de terrain. Ces modèles ont pour avantage la simplicité d'implémentation. En radio mobile, ils sont souvent utilisés pour simuler et anticiper sur le pire des scénarios par rapport à l'occupation spatio-temporelle du terrain. Ces modèles

demeurent peu réalistes du fait qu'il est peu probable que les mobiles se dispersent aléatoirement dans tout un secteur comme un bâtiment ou une ville. A l'échelle individuelle, les trajectoires sont complètement chaotiques et traduisent des "stop and go" irréels. Les paramètres de mouvement changent constamment au cours d'une simulation, produisant un déplacement aléatoire. Par contre, à l'échelle de masse ces incohérences peuvent être masquées. Ce type de modélisation pourrait révéler et prévenir des situations plausibles sur lesquelles un décideur devrait s'interroger.

Les modèles déterministes représentent les traces réelles de chaque mouvement de personnes dans une zone de simulation. Ces modèles génèrent des trajectoires réalistes mais ne sont valables que dans les zones étudiées et reflètent la réalité avec un certain degré de répétitivité. La traçabilité de ces trajectoires implique beaucoup de moyens sur le terrain, comme l'acquisition de données à partir d'enquêtes, de comptages de passages à des endroits clés, de questionnaires, etc.

Concrètement, le déplacement des personnes n'est ni totalement aléatoire, ni toujours le même. Nos déplacements suivent un certain schéma tout en conservant un degré de liberté. Ainsi un compromis doit être trouvé entre le totalement aléatoire, et la combinaison déterministe de traces. C'est ce qu'apportent les modèles hybrides qui visent à se rapprocher au mieux des déplacements réalistes.

1.2.1 Modèles de mobilité aléatoires

Ces modèles ont été introduits notamment pour réaliser des simulations de fonctionnement de réseaux mobiles. Ils ont servi à élaborer des scénarii de mobilité afin d'estimer les cellules radio visitées par un usager [Jeon 2000], [Akyildiz 2000], [Akyildiz 1996], [Ho 1995]. De tels modèles peuvent être appliqués sur un axe (1D) ou dans le plan (2D), selon les cas envisagés. En particulier, pour simuler les déplacements le long d'une route, la version 1D sera privilégiée. Dans ce cas chaque cellule est connectée à deux voisins : un précédent et un suivant. Dès qu'on s'intéresse aux mouvements de piétons par exemple, caractérisés par des changements de direction fréquents, un modèle 2D est à envisager. Dans ce cas, chaque cellule dispose de 4 voisins (cellules carrées) ou 6 voisins (cellules hexagonales).

1.2.1.1 Random Walk Mobility Model

Dans ce modèle de mobilité décrit dans [Camp 2002], un mobile se déplace vers une destination avec un angle et une vitesse prises de façon équiprobable entre des valeurs minimales et maximales. Chaque mouvement s'effectue sur un intervalle de temps constant. Les valeurs d'angle et de vitesse sur un intervalle de temps étant indépendantes de celles de l'intervalle précédent, ce mouvement traduit des changements brusques d'angle et de vitesse, ce qui est peu réaliste. A chaque instant,

le mobile se dirige donc vers une cellule voisine avec une probabilité p , ou reste dans la même cellule avec une probabilité $(1 - p)$.

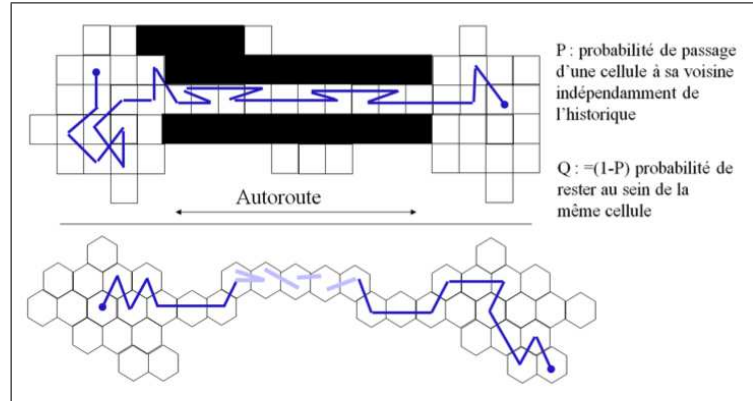


FIGURE 1.2 – Exemple de mouvements obtenus avec le modèle "Random Walk"

Les mouvements du mobile sont formalisés par des chaînes de Markov, décrites par leurs états et une matrice de probabilité de transition P . Chaque élément P_{ij} représente la probabilité de passer d'un état i à l'état j .

Ce modèle traduit des changements de direction fréquents. Sur une autoroute, il est inconcevable qu'un mobile marque fréquemment des allers-retours comme le montre la Figure 1.2. C'est tout le défaut de ce type de modèles à l'échelle de trajectoires individuelles.

1.2.1.2 Markovian Movement Model

Plusieurs auteurs utilisent ce modèle dans le cadre 1D [Burulitisz 2004], ou encore [Bar-Noy 1994]. Il est une extension du "Random Walk", décrit par le diagramme d'état Figure 1.3 et une matrice 3x3 de probabilité de transition :

$$P = \begin{bmatrix} q & 1 - q - v & v \\ p & 1 - 2p & p \\ v & 1 - q - v & q \end{bmatrix} \quad (1.1)$$

D'après le diagramme d'état, un usager peut être dans l'un des trois états suivants (Figure 1.3) :

- L : état de déplacement vers la gauche à partir de la cellule courante
- R : état de déplacement vers la droite à partir de la cellule courante
- S : état stationnaire (reste dans la cellule courante)

Wu & al. [Wu 2002] proposent une version 2D de ce modèle en définissant des distributions non uniformes des probabilités de transition entre cellules (HO). Ce

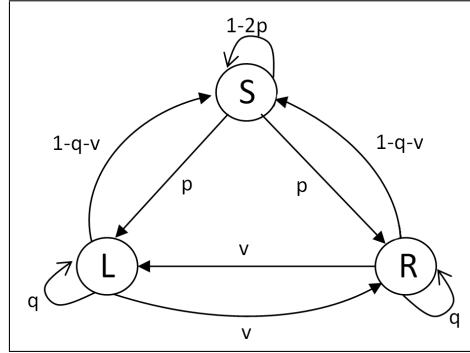


FIGURE 1.3 – Diagramme d'état du modèle de chaîne de Markov à une dimension (extrait de [Bar-Noy 1994])

type d'approches a notamment été utilisé par Perera & al. [Perera 2002] en 2002 pour simuler l'utilisation des réseaux UMTS, avec l'introduction de comportements de mobilité en fonction de la cellule et du type de mobile. Dans ce dernier article, les auteurs présentent les conditions de stationnarité et le calcul du nombre moyen de mobiles par maille en termes de taux d'arrivée, de probabilité de transition et de temps de séjour moyen.

En conclusion, ces modèles, appropriés pour modéliser le mouvement de piétons de cellules à cellules, ne sont pas réalistes. Aucun point de destination n'est défini, cela conduit à des déplacements assez chaotiques et sans objectif.

1.2.1.3 Normal Walk model

Voici une autre extension du "Random Walk", qui tient compte cette fois du mouvement précédent pour déterminer le mouvement courant. Ce modèle intègre le fait que la plupart des déplacements réels suivent soit le chemin le plus court pour rejoindre sa destination, soit une route pseudo-linéaire entre l'origine et la destination [Tsai 1999].

Un angle de rotation θ détermine une des six directions possibles d'une cellule hexagonale. La probabilité associée à cet angle est distribuée selon une loi normale avec une moyenne μ de 0° et un écart-type σ variant de $[5^\circ, 90^\circ]$: $N(0^\circ, \sigma^2)$. Ainsi cette loi favorisera les déplacements rectilignes tout en acceptant des courbures de trajectoire [Tsai 1999][Tuan 2003]. La Figure 1.4 montre trois exemples de trajectoires pour différentes valeurs de σ pouvant caractériser trois environnements différents (rural, suburbain, urbain). Il est ainsi constaté que plus σ est petit, plus les trajectoires seront douces, linéaires et surtout étendues sur de plus grandes zones.

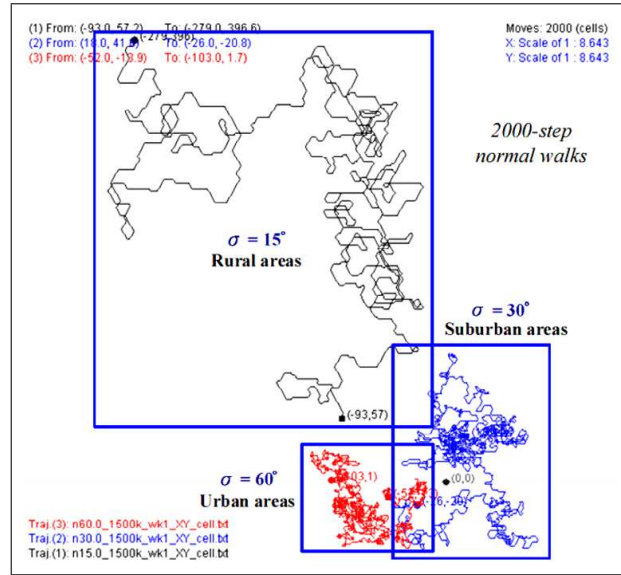


FIGURE 1.4 – Trois trajectoires de 2000 pas du modèle "Normal Walk" pour $\sigma = 15^\circ, 30^\circ$ et 60° (extrait de [Tuan 2003])

1.2.1.4 Random Waypoint Mobility model

Ce modèle est semblable au Random Walk Mobility, mais il intègre des temps de pause entre les instants de changements d'angle et de vitesse. Il est beaucoup utilisé dans la communauté des réseaux Ad hoc [Bettstetter 2003] et est implémenté dans une large majorité de simulateurs MANET (Mobile Ad hoc NETworks). A l'initialisation chaque mobile est positionné aléatoirement. Une destination de la zone de simulation ainsi qu'une vitesse de déplacement $V \in [V_{min}, V_{max}]$ sont également choisies aléatoirement. Au point de destination, le mobile y reste un certain temps $\Delta t \in [t_{min}, t_{max}]$, puis un nouveau couple (destination, vitesse) est choisi, sans aucune corrélation avec les valeurs précédentes. A tout moment un mobile est décrit par sa position $(x(t), y(t))$, sa vitesse courante $V(t)$ et son point de destination courant $(x_d(t), y_d(t))$.

Ce type de modèle est également étendu pour simuler des déplacements de groupes de noeuds formant un réseau PAN (Personal Area Networks) [Cano 2004a] [Cano 2004b] [Hong 1999]. Ceux-ci exploitent le concept de "piconets", qui sont de petits réseaux, normalement constitués d'au plus 8 noeuds, où un des noeuds est le maître dans la topologie.

1.2.1.5 D'autres modèles aléatoires

De nombreux autres travaux ont traité des modèles aléatoires. On notera en particulier le "Smooth Mobility Model" [Bettstetter 2001] qui permet de lisser les changements de vitesse et de direction. C'est le premier modèle aléatoire faisant apparaître des notions de classe de mobilité (par exemple piétons, voitures en centre ville, vélos) en définissant différents jeux de paramètres : vitesse max, probabilité des vitesses de préférences... De même Liang et Haas [Liang 2003] présentent le modèle de mobilité "Gauss-Markov". Ce modèle intègre la notion de vitesse pour estimer la position future. Il est notamment utilisé par Camp & al. [Camp 2002] et fournit la simulation illustrée Figure 1.5.

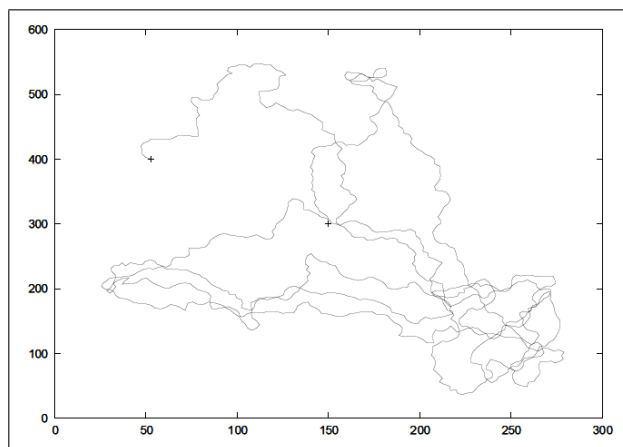


FIGURE 1.5 – Exemple de motif de déplacement d'un mobile utilisant le modèle de mobilité Gauss-Markov (extrait de [Camp 2002])

1.2.2 Modèles de mobilité déterministes

La seconde classe de modèles présentée ici, dite déterministe, vise à construire des règles à partir de données d'entrée telles que des enquêtes, des questionnaires et des comptages.

Le seul article [Scourias 1999] qui traite de cette approche présente une étude réalisée dans la région de Waterloo. L'objet de cette étude était de simuler des trajectoires de mobilité à partir des données collectées. Les informations utilisées pour construire ces trajectoires sont :

- la connaissance des motifs de déplacement au cours de la journée
- le chainage des activités réalisées

Afin d'alimenter le modèle présenté dans [Scourias 1999], une enquête a été réalisée sur les déplacements des personnes de 5 ans et plus. Cette enquête, du type Enquête Ménages/Déplacements (EMD), contient en particulier les heures de départ

et d'arrivée des déplacements, ainsi que le motif (travail, scolaire, accompagner quelqu'un, achats, loisirs, retour chez soi).

Des matrices sont réalisées pour refléter la durée des activités liées à chaque déplacement, ainsi que les transitions entre activités. Ces matrices définissent le comportement obtenu lors du processus de simulation.

Bien que l'auteur mette en avant l'intérêt d'un tel modèle, le travail demandé est particulièrement fastidieux, ce qui limite sa portée.

1.2.3 Modèles de mobilité hybrides

Bien que les modèles aléatoires et déterministes fournissent des trajectoires plausibles, ils restent éloignés des déplacements réels ou sont très coûteux à construire et ne peuvent pleinement être utilisés pour élaborer des politiques d'aménagement, en matière de transport par exemple. Les modèles de mobilité hybrides répondent à cette exigence, en combinant mouvements aléatoires et régularité de déplacements.

1.2.3.1 Modélisation des flux : approches issues des systèmes de transport

Les modèles présentés ici permettent de reproduire les principales caractéristiques des mouvements de personne. Ils sont largement utilisés dans les modèles de mobilité dits "intégrés" présentés section 1.2.3.2.

Fluid flow model : Il s'agit de modèles qui s'appuient sur des théories physiques issues de l'hydrodynamique [Lam 1997]. L'accent n'est pas mis sur le déplacement individuel mais au contraire sur une analyse du nombre moyen de personnes traversant la région étudiée.

Ces modèles sont particulièrement utilisés pour des études de transport, afin d'identifier le taux d'utilisation d'un axe routier [Markoulidakis 1997]. Dans ce modèle les mobiles appartenant à un même flux suivent la même direction à la même vitesse. Ce type de modèles a également été utilisé pour étudier l'impact du trafic autoroutier sur les réseaux mobiles [Lizee 2001], lors d'un ralentissement dû à un accident par exemple.

Gravity model : Là encore il s'agit de modèles directement inspirés de la physique. Cette famille de modèles est très souvent utilisée dans les modèles "intégrés" pour simuler les déplacements entre différentes régions géographiques. Un des avantages liés à ces modèles est qu'ils permettent de simuler des flux à différentes échelles,

du niveau local jusqu'au niveau national et international [Lam 1997]. Formalisons maintenant le fonctionnement de ce modèle. Soit $T_{i,j}$ la quantité de trafic se déplaçant d'une région i vers une région j . On a alors :

$$T_{i,j} = k_{i,j} P_i P_j \quad (1.2)$$

Avec : P_i la population habitant dans la région i

($k_{i,j}$) les coefficients caractérisant les déplacements pour aller de i à j .

La matrice K peut être construite de différentes manières pour définir les différentes variantes de ces modèles. En particulier, pour se rapprocher de la loi de la gravité de Newton, K peut être calculée proportionnellement à l'inverse du carré de la distance entre les zones i et j , d'où le nom de ces modèles. Dans des versions plus complexes, cette matrice peut également intégrer des données de trafic telles que des comptages routiers.

Bien que l'équation 1.2 décrive un trafic agrégé, il est possible de modéliser des déplacements par individu. En effet, en fixant l'attraction P_i de la région i , et la probabilité $T_{i,j}$ d'aller de i vers j , on peut agir sur des mobiles pour simuler leurs déplacements.

1.2.3.2 Modèles de mobilité intégrés

Ces modèles ont la particularité de regrouper en leur sein différents modèles adaptés à différents niveaux d'échelle. On distingue souvent dans ces modèles trois niveaux de précision, et donc trois stratégies de simulation.

Selon le type d'application, les modèles utilisés à chaque niveau d'échelle vont différer. Ainsi, dans [Lam 1997] est proposé un modèle de mobilité pour simuler l'impact des déplacements sur le fonctionnement des réseaux 3G. On distingue le niveau local, s'appuyant sur un modèle Markovien (cf 1.2.1.2), et les niveaux nationaux et internationaux qui s'appuient sur des approches gravitaires 1.2.3.1.

Dans [Markoulidakis 1997], le modèle de mobilité concerne l'échelle urbaine. Les trois niveaux de ce modèle sont donc beaucoup plus fins. Ainsi on distingue le niveau "rue", le niveau "quartier" et le niveau "agglomération".

Mais d'autres approches considèrent deux niveaux uniquement. C'est le cas par exemple dans [Liu 1998], où l'approche "Global-Local mobility model" (GLMM) est proposée. Alors qu'au niveau global un modèle déterministe est proposé, au niveau local c'est un modèle stochastique qui est choisi. Cette approche a permis de simuler le fonctionnement d'un réseau mobile, selon les déplacements des personnes.

Le modèle global a ainsi permis de simuler les transferts intercellulaires alors que le modèle local s'attachait à simuler les mobilités intracellulaires.

On notera que d'autres modèles de complexité variables ont été proposés, tels que [Erbas 2001, Erbas 2002] qui intègrent dans leurs modèles une prédiction des déplacements à venir. Ils s'appuient pour ce faire sur l'aspect régulier et répétitif de nos déplacements d'un jour à l'autre.

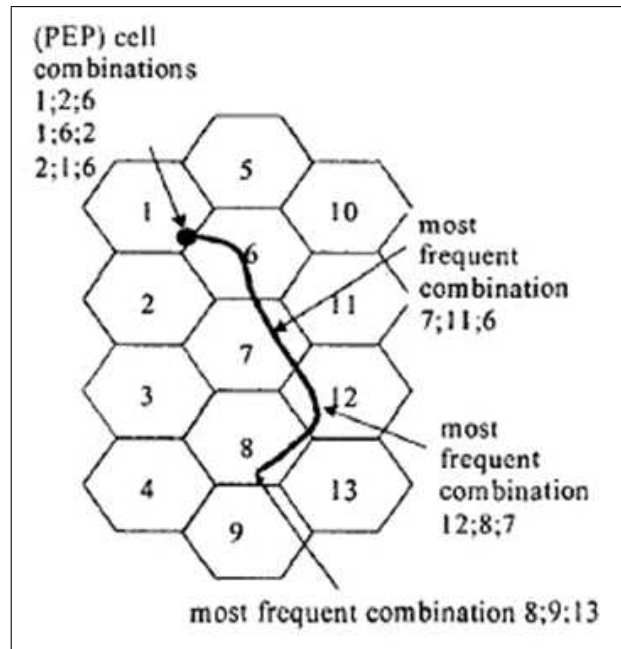


FIGURE 1.6 – Séquence de cellules d'un profil usager (extrait de [Erbas 2001])

Ces algorithmes laissent entrevoir les possibilités que peut apporter la téléphonie mobile pour caractériser les déplacements des personnes.

1.2.3.3 Modèles basés sur la géographie

On rencontre dans la littérature des modèles de mobilité s'appuyant sur une description fine du territoire.

Kammann & al. [Kammann 2003] proposent un modèle de mobilité basé sur des cartes géographiques. Ils ont en particulier proposé deux méthodes de déplacement, l'une utilisant l'algorithme de Lee [Lee 1961], développé à l'origine pour la constitution automatique de circuit imprimé, et l'autre reposant sur la diffusion des gaz dans l'espace vue en thermodynamique et utilisée pour trouver les chemins que peuvent prendre des robots [Azarm 1994].

Après avoir déterminé les chemins, un modèle de vitesse est appliqué. La vitesse

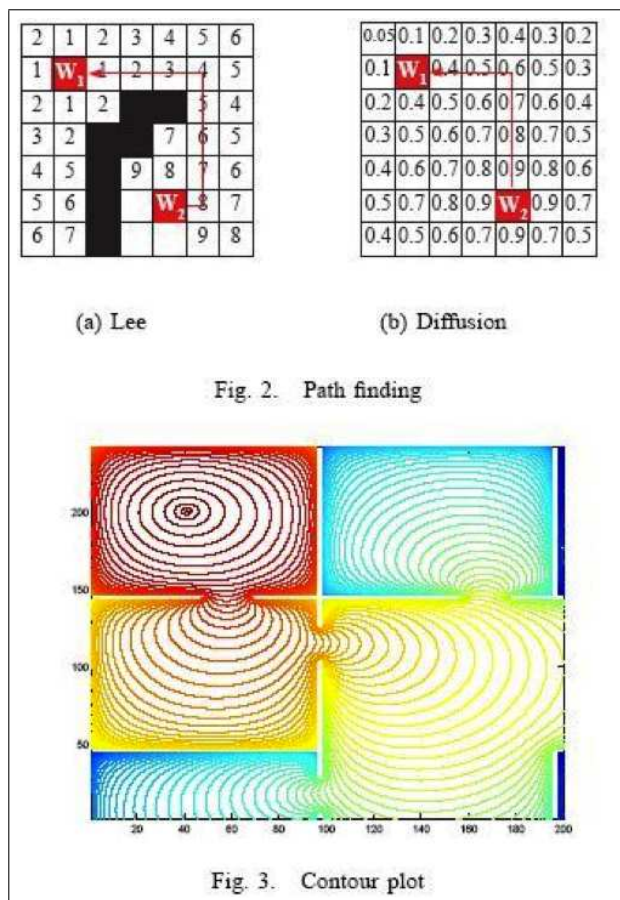


Fig. 2. Path finding

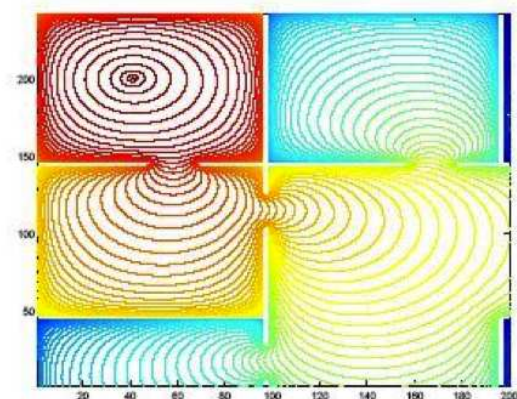


Fig. 3. Contour plot

FIGURE 1.7 – Méthodes de détermination des chemins (extrait de [Kammann 2003])

$v(t + T_s)$ est calculée suivant la valeur précédente $v(t)$ et une valeur d'accélération $a(t)$. Cette dernière est générée par un processus aléatoire. La vitesse est comprise ente $[v_{min}, v_{max}]$ et $0m.s^{-1} \leq v_{min} \leq v_{max}$:

$$v(t + T_s) = \min(\max(v(t) + a(t).T_s, v_{min}), v_{max}) \quad (1.3)$$

D'autres modèles également basés sur la géographie intègrent, en plus du réseau routier, une information sur la localisation des bâtiments. Dans ces modèles, la nature des bâtiments permet d'y associer une attractivité variant selon l'heure de la journée. Dans [Joumaa 2007] est présenté le modèle Mask-Based Mobility Model (MBMM).

Il s'agit d'une approche maillée où chaque maille (cellule carrée) a un poids d'attraction dépendant de la géographie et du moment de la journée. A chaque étape, le mobile se déplace vers une des 8 mailles voisines (cf. 1.8). La direction est choisie à chaque étape en fonction d'une chaîne de Markov de 9 états. La probabilité

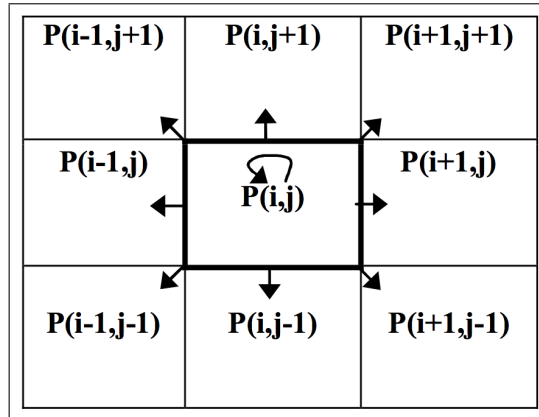


FIGURE 1.8 – MBMM : grille des probabilités de transition

de transition est calculée en fonction du poids de chaque maille voisine.

Le modèle V-MBMM [Ait Ali 2010b] est une adaptation de MBMM pour la simulation des déplacements de véhicules, toujours par une chaîne de Markov qui prend en compte la topologie et les poids d'attraction, mais cette fois des routes.

1.2.3.4 Le modèle à 4 étapes

Le modèle à quatre étapes est le modèle "classique" des études de trafic. Dans [Bonnell 2002] ce modèle est longuement développé et chacun de ses composants est analysé. Il s'agit d'un modèle qui est utilisé dans de nombreux projets de simulation de flux. Contrairement aux modèles habituels, il s'agit ici plutôt d'un schéma général à suivre. Son application suppose dans un premier temps de diviser le secteur d'étude en différentes zones géographiques ou logiques, généralement à partir de données socio-économiques, permettant d'agréger les différentes données. Ce modèle décrit les quatre étapes principales que nous rappelons succinctement.

La génération : Cette première étape a pour but de déterminer le nombre de déplacements, c'est-à-dire quantifier les flux entre les différentes zones. Il s'agit donc de savoir si les individus se déplacent ou non et si oui dans quelles zones. On cherche alors à modéliser ces choix en se plaçant du point de vue d'une zone. On obtient alors d'une part, l'ensemble des déplacements émis par une zone donnée et d'autre part l'ensemble des déplacements dont la destination se situe dans cette même zone. On parle alors d'émission et d'attraction d'une zone qui constitue la génération. Ce nombre de déplacements entrants ou sortants est fonction de nombreux paramètres, notamment les caractéristiques d'occupation du sol des zones : les services, la population, l'emploi, etc. Il y a deux moyens de calculer les émissions et attractions par zone : une méthode statistique (régression linéaire) ou une utilisation de taux de mobilité

obtenus par enquête-ménages [Nguyen-Luong 2000a]. En sortie de cette étape, on obtient par zone deux vecteurs, un pour les émissions E , l'autre pour les attractions A . La génération attribue donc à chaque zone ses vecteurs E/A .

La distribution : Dans cette seconde étape, on s'intéresse cette fois à la situation où les individus se déplacent, d'une zone O origine à une zone D destination. Il s'agit en fait de lier les émissions et les attractions précédemment obtenues par couple de zone O/D tous modes confondus. Cette étape consiste donc à reconstituer les matrices représentatives des flux pour tous les couples de zones O/D en s'ajustant au mieux aux observations de l'aire d'étude. A cette étape les volumes de déplacement entre les différentes zones du découpage initial sont connus. Plusieurs modèles sont utilisés à ce niveau comme le maximum d'entropie ou le modèle gravitaire.

Le choix modal ou la répartition modale : On s'intéresse dans cette étape au choix du mode de transport des individus c'est-à-dire par quel moyen ils se déplacent. La matrice "tous modes" produite précédemment est éclatée en matrices par modes. Il en ressort une estimation des volumes de déplacement pour chaque moyen de transport à prendre en compte. Les matrices O/D pour les différents modes à modéliser sont générées à l'issue de cette étape. Les méthodes mathématiques généralement utilisées sont la régression multiple, les tables de parts modales par segment de population, ou encore des courbes sigmoïdales pour des modèles agrégés.

L'affectation : Cette dernière étape consiste à modéliser le choix de l'itinéraire des individus pour se rendre d'une zone O à une zone D . Concrètement, il s'agit ici d'affecter toutes les matrices précédemment générées sur le réseau à l'aide de techniques opérationnelles (algorithme du plus court chemin, optimisation, etc.). A la fin de cette étape, on obtient des réseaux chargés, c'est-à-dire que les prévisions (telles que par exemple le nombre de véhicules, de voyageurs ou encore la vitesse et le temps de parcours) sont distribuées et connues sur chaque arc. Pour les réseaux routiers, une technique d'affectation statique avec contrainte de capacité est généralement employée. Pour les réseaux de transport en commun on voit plus communément une affectation "tout ou rien". Dans les recherches effectuées, plusieurs projets se basent sur le logiciel DAVISUM pour effectuer cette étape.

En résumé, ce modèle se base sur l'offre et la demande. Les trois premières étapes définissent le calcul de la demande, la quatrième étape calcule l'affectation sur l'offre. L'ensemble de ces quatre étapes permet ainsi une affectation modale des flux de déplacements. Par ailleurs, ce modèle localise toutes les activités à une période de temps donnée (heure de pointe par exemple).

Les premiers modèles à quatre étapes sont développés aux États-Unis dans les années 1950. Ils sont importés en France à partir des années 1960. Historiquement,

les premiers modèles ont une approche agrégée et plutôt déterministe. Orientés automobile aux États-Unis, ils subissent de légères adaptations pour répondre au besoin grandissant du transport en commun dans les grandes villes de France, qui continueront cependant à l'exploiter dans sa dimension agrégée. Le modèle GLOBAL [RATP 2010] de la RATP est ainsi conservé (modèle classique à 4 étapes agrégé). Il correspond à la zone d'étude Région Ile-de-France et sud de l'Oise.

Par la suite les modèles désagrégés ont permis de tenir compte de l'hétérogénéité des comportements individuels. En effet, il existe plusieurs manières de modéliser les individus à des échelles différentes en fonction des besoins. Pour comprendre, à l'échelle microscopique, qui correspond au niveau de désagrégation maximum, on s'intéresse à l'individu et au détail qui le caractérise. On parle alors de données désagrégées. Une échelle intermédiaire consiste en une agrégation des données précédentes suivant des caractéristiques individuelles communes (classe sociale, modalités identiques, agrégation géographique, etc.). A l'échelle macroscopique, on s'intéresse alors à une population entière (nombre total de déplacements, moyenne d'âge, etc.).

Il est donc intéressant de constater que les niveaux d'agrégation peuvent s'opérer à différents niveaux ou modules et peuvent être de plusieurs types (agrégation temporelle, agrégation géographique, etc.). D'après les recherches et les exemples de modèles trouvés il n'est pas démontré qu'un niveau d'agrégation soit meilleur que les autres, par contre il est établi qu'à chaque niveau d'agrégation, le modèle perd en consistance mais gagne en stabilité. On peut alors envisager pour toutes les phases d'un modèle à quatre étapes d'utiliser au choix des données agrégées ou désagrégées en fonction de l'utilité finale du modèle.

Ainsi par rapport au modèle classique agrégé, quelques améliorations ont été apportées par une approche plus désagrégée dans les différentes étapes décrites précédemment [Nguyen-Luong 2000a]. On parlera alors de modèle à quatre étapes hybrides (ex : logiciels EMME2, MINUTP ou TRIPS utilisés en France) par opposition au modèle à quatre étapes classique. De plus en plus, les éditeurs présentent souvent une suite de logiciels adaptables et modulaires (un module pour la génération, un pour l'affectation, etc.). De même une suite de logiciels basée sur un modèle statique agrégé peut comporter des modules dynamiques désagrégés.

En ce qui concerne les améliorations dans la génération, il faut trouver le bon compromis, entre un système entièrement désagrégé dans lequel chaque demande de déplacement est étudiée individu par individu pour être additionnée sur une population à la fin du processus, et un système agrégé dans lequel tous les individus dans une zone sont supposés être moyens. D'après le rapport de l'IAURIF [Nguyen-Luong 2000b], il semble que la segmentation de marché soit un compromis très intéressant lorsqu'on ne dispose pas de moyens de développer un modèle désagrégé. C'est une approche par exemple exploitée dans le logiciel VISEM-DAVISUM (PTV). Cependant elle se trouve vite limitée en termes de prévisions dans le temps notamment dans l'étape

de choix modal.

Concernant l'amélioration dans le choix modal, beaucoup de modèles mathématiques ont été développés. Avec la désagrégation, plusieurs modèles de choix discrets ont été proposés (LOGIT, PROBIT, DOGIT). Probabiliste et basé sur la théorie d'utilité d'un certain mode pour un déplacement donné, c'est le modèle LOGIT qui est utilisé dans la plupart des cas de modèles désagrégés au niveau de la modalité (mise en oeuvre par le logiciel HIELLOW).

Au final, étant basé sur la distinction des déplacements par motif de déplacement, cette approche à quatre étapes suppose l'indépendance de ces motifs dans un découpage géographique donné. Les modèles d'affectation conventionnels sont statiques (affectation des déplacements sur le réseau en utilisant une matrice agrégée sur une période moyenne) et il n'y a pas de chargement graduel dans le temps.

Pourtant d'après les recherches effectuées, le modèle à quatre étapes classique ou hybride est un modèle reconnu et satisfaisant dans la simulation des transports notamment à l'échelle macroscopique (encore utilisé aujourd'hui). En effet, les modèles existant (ex : GLOBAL RATP [RATP 2010]) semblent bien adaptés à l'étude des infrastructures routières ou de transport en commun à l'échelle macroscopique et remplissent parfaitement leur rôle en matière de dimensionnement des infrastructures et de calcul de leur rentabilité socio-économique.

Cependant, il faut ajouter un bémol car la modélisation des transports aujourd'hui impose de s'intéresser également à l'urbanisation notamment avec les modèles d'occupation des sols tels que les modèles "LUTI" (Land-Use Transport Integrated models). En termes d'urbanisation, la modélisation de trafic classique (statique) à quatre étapes est dépassée même si elle est encore largement utilisée par défaut en France dans la pratique. C'est à partir de ce constat qu'est né le projet SIMAURIF [Nguyen-Luong 2008] liant un modèle d'urbanisation URBANSIM et un modèle dynamique de prévision de trafic DAVISUM-METROPOLIS, véritable projet pionnier en France dans le domaine de la modélisation de l'interaction urbanisation-transport (cf. p11 [PREDIT 2007]).

La réalité de la mobilité urbaine à une échelle plus microscopique introduit une problématique différente. La focalisation au niveau de l'individu devient nécessaire. Cette complexité nouvelle réside dans la prise en considération de facteurs du comportement des individus. Dans ce domaine, les modèles à quatre étapes classiques sont dépassés. Pour cause, avec la plupart de ces modèles, l'affectation en pratique est statique. Conséquence, chaque mobile apparaît simultanément sur chaque arc de l'itinéraire, défiant la notion réelle d'espace-temps. Face à cette approche déterministe, apparaît alors une notion plus stochastique d'espace, de temps et d'activités qui insiste sur les ressorts de la mobilité plutôt que sur les volumes de trafic et fait apparaître le concept de "Time-geography" (cf. p12 [PREDIT 2007]). On parle alors d'approche "activité-centrée" mettant en avant le comportement des individus dilués plus finement dans le temps.

1.2.3.5 Modèles basés sur l'activité

D'après les recherches, les comportements de mobilité reposent sur des composantes sociales et des contraintes spatio-temporelles, d'où la notion d'espace-temps-activité. Si l'approche classique tendait à réduire la mobilité à un simple déplacement, l'approche activité-centrée quant à elle, basée sur l'activité individuelle, introduit la complexité réelle du déplacement au niveau de l'individu. Elle impose une segmentation plus détaillée de l'activité et un traitement plus élaboré des décisions de déplacement et des lieux associés.

Les comportements humains sont individuellement modélisés en représentant la façon dont chaque individu choisit entre plusieurs options suivant ses perceptions propres (préférence, habitude, incertitude, alternative, limite de temps ou d'argent, manque d'information, etc.). Les données d'entrées sont obligatoirement de plus en plus fines.

Les modèles de transport basés sur l'activité sont en développement dans plusieurs pays et sont généralement considérés comme les modèles de transport du futur (URBANSIM).

Par ailleurs, ce nouveau modèle de demande permet de simuler une chaîne de déplacements ou plus exactement une suite d'activités (changement d'état des individus à un autre) et s'avère donc intéressant dans la modélisation microscopique de la mobilité urbaine pour notre cas d'étude. On distingue 2 types de représentations pour les suites d'activités, les suites temporelles et les suites événementielles.

Les disparités avec le modèle classique sont nombreuses et peuvent être résumées comme suit (cf. p8 [PREDIT 2007]) :

	Modèle à 4 étapes	Modèle activité-centrée
Objectif du modèle	Prévisions de trafic pour planifier les investissements	Comprendre les comportements de déplacement
Unité d'observation géographique	Approche agrégée : groupement d'ilots ou communes	Approche désagrégée à l'échelle de l'individu ou du ménage
Unité d'observation du déplacement	L'origine et la destination sont prises en compte	Décomposition en séquences d'activités
Sources de données	Les données socio-économiques proviennent d'enquêtes variées : recensement (RGP), enquête transport (ETC), Enquêtes Ménages/Déplacements (EMD), enquêtes cordon.	Nécessite des données sur les activités individuelles à partir d'enquêtes réalisées localement (EMD)
Interaction données socio-économiques et déplacements	Relation indirecte entre les données de population et de trafic	Relation directement issue de l'enquête
Résultats des analyses	Le nombre de déplacements est consigné dans une matrice zone à zone	Les déplacements sont analysés dans l'espace et le temps

Cette confrontation avec l'approche classique met non seulement en évidence le niveau de précision des données d'entrée mais également le niveau de complexité méthodologique en termes de modélisation informatique. C'est pour répondre à ce dernier point que s'opère le vrai changement avec l'apparition des systèmes multi-agents dans les années 1980. Dans un contexte de puissance informatique grandissant, il constituera un véritable bon en avant dans la modélisation de la mobilité.

1.2.3.6 Autres modèles hybrides

Kim & al. proposent un modèle de mouvement 3D dans des immeubles avec cage d'escalier [Kim 1998] et avec ascenseurs [Kim 2000].

Les usagers se déplacent de façon rectiligne sur un étage jusqu'à ce qu'ils changent de direction (gauche, droite, arrière suivant une distribution de Poisson) puis continuent tout droit, ainsi de suite. Si le point où ils tournent est dans la zone de cage d'escalier ou d'ascenseur, les abonnés bougent horizontalement ou verticalement selon certaines probabilités. La vitesse du déplacement horizontal ou vertical est uniformément réparti entre $[0, V_{min}]$ et $[0, V_{max}]$ respectivement pour les piétons.

Tian & al. [Tian 2002] proposent un modèle de mobilité basé sur les graphes. Un centre ville est approximé par un graphe, dont les sommets représentent les lieux fréquentés par les usagers et les arcs les connexions entre ces lieux. Les mouvements des usagers sont donc limités aux arcs entre les sommets. Chaque sommet a le même poids, c'est-à-dire qu'aucun point n'est plus attractif qu'un autre pour l'utilisateur mobile.

Les sommets origine et destination de chaque mobile sont choisis aléatoirement. La vitesse de déplacement est comprise entre $[V_{min}, V_{max}]$. Le mouvement se fait toujours par le chemin le plus court. A l'arrivée au point de destination, le mobile reste un certain temps sur place $[tP_{min}, tP_{max}]$, avant de partir vers un nouveau sommet aléatoire.

Ce modèle est proche du "Random Waypoint" (cf. 1.2.1.4). La seule différence pour le rendre plus réaliste est de contraindre les mouvements sur un graphe construit à partir de la connaissance de la géographie.

Une amélioration de ce type de modèle serait d'allouer une certaine attraction des sommets.

A nouveau pour les réseaux Ad Hoc, dans [Bittner 2005] est présentée une modélisation basée sur des zones connectées de différentes densités par un graphe. Les auteurs affirment que ce modèle n'est pas aussi spécifique que la plupart des modèles les plus réalistes qui permettent de simuler certains scénarii de mobilité.

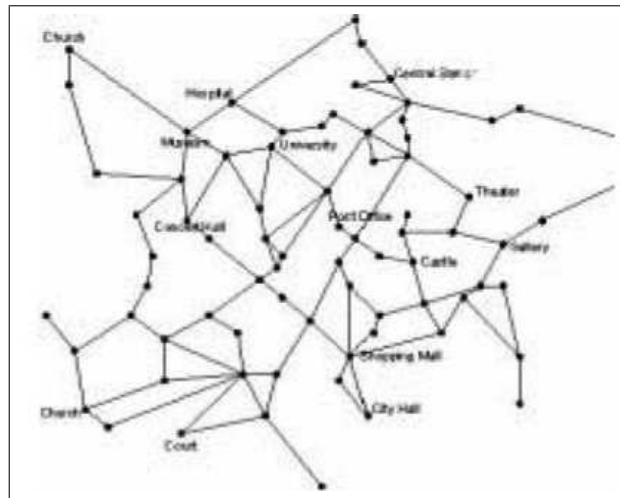


FIGURE 1.9 – Graphe du centre ville de Stuttgart (extrait de [Tian 2002])

Cette modélisation permet de préserver cette structure hétérogène, c'est-à-dire la répartition du trafic hétérogène de forte densité sur certaines zones ou cluster associé et de faible densité sur d'autres zones ou des axes. Le poids des arcs du graphe définit la probabilité qu'un mobile choisisse ce chemin lorsqu'il quitte une zone constituant le sommet. Pour chaque sommet, un intervalle de temps de séjour dans la zone est donné. Ce dernier est ensuite distribué uniformément dans cet intervalle. Le déplacement est modélisé de deux manières différentes selon où se situe le mobile :

- A l'intérieur des sommets par un "Random Waypoint Mobility model"
- Sur les arêtes, le déplacement se fait sur des choix aléatoires respectant les poids de chaque arête et la vitesse.

La difficulté de cette modélisation est la définition des zones et du graphe qui peut être très complexe pour les systèmes radio-mobiles couvrant des zones géographiques très élargies. Des niveaux de hiérarchisations pourraient être envisagés suivant le type et la taille des zones à étudier.

Pour information, Jardosh & al. [Jardosh 2005] proposent l'utilisation des mêmes modèles de mobilité aléatoire (RWP...) pour les réseaux Ad Hoc, basés cette fois sur un graphe de Voronoï afin de définir les chemins réels entre obstacles physiques.

1.3 Cas de la téléphonie comme marqueur de mobilité

Les téléphones mobiles sont à l'origine de nouvelles formes de connexion, tant entre les personnes qu'avec les lieux et les infrastructures urbaines. Le lien étroit existant entre l'utilisateur de mobile et son téléphone a dès 1994 incité des chercheurs

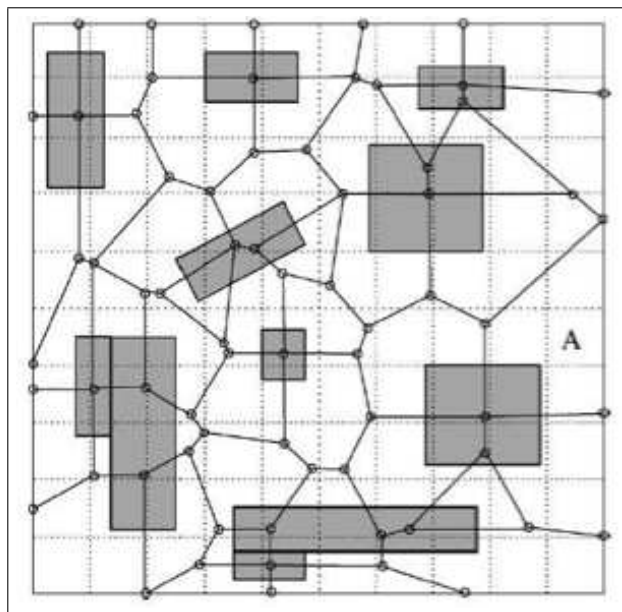


FIGURE 1.10 – Terrain simulé par un graphe de Voronoï (extrait de [Jardosh 2005])

à étudier les déplacements de ces appareils : si l'on est capable de comprendre les mouvements des téléphones on en déduira une information pertinente sur la mobilité des personnes. Les premières expérimentations se sont pourtant révélées peu concluantes, en particulier en raison de la faible précision des données et du faible nombre d'abonnés jusqu'au début des années 2000.

Alors que jusqu'en 2006 quelques rares expérimentations seulement ont été réalisées [Steenbruggen 2011], les premiers travaux significatifs ont été menés par une équipe du Massachusetts Institute of Technology (MIT), au laboratoire SENSEable City. Ils ont ainsi réalisé plusieurs expérimentations utilisant ces nouvelles technologies pour étudier ce qu'ils appellent le Paysage Mobile [Ratti 2006], comportement dans l'espace et dans le temps des sociétés urbaines. Contrairement au principe des applications de type Location Based Services (LBS) [Ahas 2005], ici on ne s'adresse pas directement à l'utilisateur du mobile. Il s'agit donc d'exploiter les données téléphoniques sans intervention de l'utilisateur, en observant uniquement le fonctionnement du réseau lui-même.

Depuis 2006, plusieurs travaux de recherches ont ainsi été menés pour exploiter les informations qui transitent en permanence par les réseaux de télécommunications afin d'observer les mouvements de la ville à différents moments de la journée. Alors que pendant longtemps un des principaux freins à l'utilisation de cette source de données était la difficulté d'accès aux données elles-mêmes, la libéralisation actuelle des données notamment au travers de l'Open Data, incite (bien que très timidement encore!) les opérateurs à mettre à disposition certaines informations sur la façon dont leur réseau est utilisé. La localisation obtenue est bien moins précise qu'avec

un GPS, cependant la masse de données liée à la présence généralisée des relais téléphoniques et au grand nombre d'utilisateurs permet d'envisager une exploitation très pertinente.

1.3.1 Le projet CAPITAL [Systems 1997]

Ce premier projet réalisé dans la région de Washington DC (USA) fut initié en 1994 pour une durée de 27 mois. Il s'agissait d'utiliser la téléphonie mobile pour apporter une meilleure compréhension de la mobilité urbaine. Le projet CAPITAL (Cellular Applied to ITS Tracking And Location) avait ainsi pour objectif de fournir une information temps réel sur l'état du trafic routier, à travers l'observation de données provenant directement d'un opérateur téléphonique. En raison de la taille trop importante des cellules radios et certainement du manque d'expérience de ce type, les résultats se sont avérés peu concluants. Le niveau de précision de l'ordre d'une centaine de mètres ne permettait pas de construire une information fiable sur l'état du trafic.

1.3.2 Les travaux du MIT : une source de référence

Le laboratoire SENSEable City du MIT publie les premiers résultats de visualisation dynamique de la mobilité à partir des téléphones portables en 2006. Les techniques d'étude de la mobilité ont ainsi été développées et expérimentées dans des contextes variés de différentes villes à travers le monde. A l'origine de ces travaux un architecte italien, Carlo Ratti, un passionné des interactions pouvant exister entre la ville et les nouvelles technologies. Nous citons ci-après une suite de travaux menés au MIT sur ce sujet.

1.3.2.1 City out of Chaos : Social Patterns and Organization in Urban Systems [Pulselli 2006]

Ce travail est mené avec l'University of Siena et en partenariat avec un opérateur téléphonique dans la région de Milan sur un territoire de 20km x 20km. Les résultats obtenus sur 16 jours d'observation ont permis de faire ressortir des cartes de chaleur qui reflètent le niveau d'utilisation des téléphones portables en fonction des lieux et des moments de la journée.

1.3.2.2 Real Time Rome [Francesco 2006]

A la même période que le projet mené à Milan, le MIT présente le projet Real Time Rome à la Biennale de Venise. Il s'agit de l'un des premiers exemples d'un système de surveillance urbaine temps réel qui permet de cartographier les déplacements des téléphones portables en y ajoutant d'autres données de mobilité issues des réseaux de transport (bus et taxis). Les projections issues de ce travail fournissent un point de vue sur le fonctionnement de la ville de Rome, d'un point de vue géographique et socio-économique. Les cartes projetées lors cet événement permettent de montrer :

- la répartition de différents groupes sociaux dans la ville tels que les touristes et les résidents
- l'occupation des zones de la ville et le déplacement des personnes dans les zones pendant les événements spéciaux
- les endroits les plus attractifs à Rome
- le niveau de corrélation entre la distribution du transport public (taxi, bus) et la densité de la population

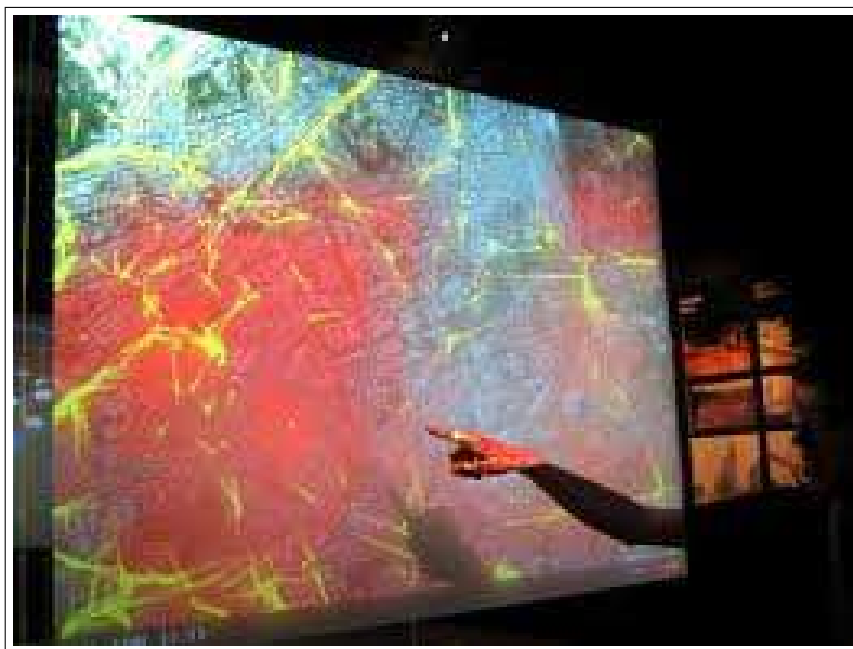


FIGURE 1.11 – Vue de Rome en temps réel

La Figure 1.11 propose une illustration de ce qui est observable. La couleur jaune représente les lignes autobus et la couleur rouge correspond à la densité de population en temps réel.

Quelques années après, Calabrese et al. [Calabrese 2009] ont enrichi les résultats

de 2006 à travers le projet WikiCity. Les données utilisées dans WikiCity sont :

- localisation des bus et taxis par GPS
- bruit de la circulation (traffic noise) à partir de capteurs sans fil installés à Rome
- intensité du trafic (traffic intensity)
- densité des touristes obtenue à partir de l'analyse de l'IMSI des utilisateurs localisés
- densité des piétons obtenue à partir des données de localisation des utilisateurs supposés marcher à faible vitesse
- vitesse des véhicules

1.3.2.3 Mobile Landscapes : Graz in Real Time [Ratti 2007]

A Graz, en Autriche, Ratti et al. [Ratti 2007] ont cartographié le remplissage des zones de la ville en temps réel, ainsi que les déplacements des utilisateurs. Afin d'éviter des problèmes de confidentialité, la collecte des trajets a eu lieu lors de l'exposition M-City auprès de volontaires. Les personnes souhaitant participer à l'événement se sont enregistrées par SMS et ont été suivies durant 24h. Lors de l'exposition, le public a ainsi pu voir s'animer les cartes d'intensité du trafic téléphonique en direct. Le projet réalisé à Graz a permis aux auteurs d'aborder les rapports entre ville, technologie de communication et individu. Il s'agit d'un travail de fond qui aborde les questions fondamentales que l'on se pose lorsqu'on souhaite aborder ce sujet. En particulier sont abordés les apports de la téléphonie en tant que nouvel outil : l'influence de la téléphonie sur notre perception de la ville, l'utilisation de la téléphonie en tant qu'outil d'enquête et d'analyse (comment les citoyens peuvent à travers ces outils participer activement à la construction de leur environnement), etc. Mais les auteurs vont plus loin en s'interrogeant sur les modes de représentation cartographique à inventer pour retranscrire les informations construites ou encore sur l'impact en matière de vie privée : comment peut-on concilier information de mobilité et liberté individuelle, vie privée ?

1.3.2.4 Live Singapore [Kang 2013]

Il s'agit d'un projet réalisé en collaboration avec SMART (Singapore MIT Alliance for Research and Technology). Les résultats ont été présentés lors d'une exposition organisée au "Musée d'Art de Singapour" en 2011. L'objectif de ce travail est le développement d'une plate-forme qui permette de visualiser l'activité de la ville de Singapour, dans un contexte dynamique et fonctionnant en temps réel. La Figure 1.12 montre les différentes vues proposées dans ce projet. A gauche se trouve une cartographie de la température. La température dans les villes est généralement plus élevée qu'à la campagne, et plus il fait chaud, plus le recours à la climatisation est

important, et donc plus il y a de production de chaleur. Cette image permet de se rendre compte des lieux les plus chauds et les plus énergivores. On aperçoit ensuite une carte isochrone qui permet, en fonction du moment de la journée, d'identifier les lieux où l'on peut se rendre par durée de trajet (en 5min, en 10min, etc.) La troisième vue rend compte du taux d'utilisation de la téléphonie mobile. L'intensité de chaque maille traduit le niveau d'utilisation des téléphones. La vue suivante traduit le taux d'utilisation des taxis, selon les lieux mais aussi selon la pluviométrie. L'avant-dernière illustration montre les échanges de fret, en bateau et en avion, entre Singapour et l'ensemble du monde. Enfin, la dernière vue permet de se rendre compte des répercussions que peut avoir un événement comme le Grand prix de Singapour sur la façon dont la population se déplace et utilise son téléphone portable.

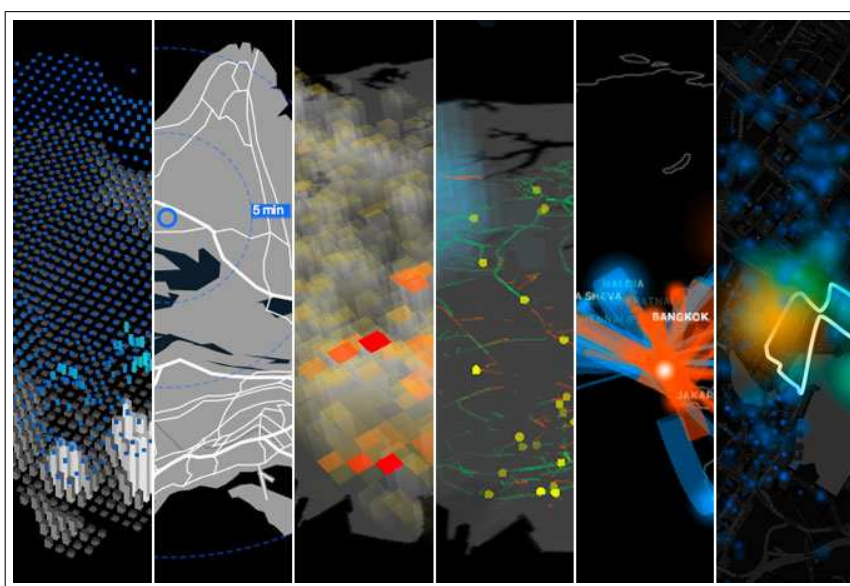


FIGURE 1.12 – Illustration des vues de la mobilité à Singapour

L'ensemble des cartes proposées est intégré à un système d'information géographique qui fournit une nouvelle façon de voir la ville. L'accent est ici mis sur la représentation des résultats, afin de rendre compte de la façon la plus juste des tendances sociales, économiques et de mobilité dans une ville.

1.3.2.5 D'autres contributions...

Afin d'être complet on citera encore quelques travaux menés par le MIT, et ouvrant des pistes intéressantes dans le domaine :

- Etude des choix modaux en matière de transport à partir des données téléphoniques croisées avec des données de transport [Holleczek 2013]
- Estimation de l'erreur faite lors de l'utilisation des mobiles pour identifier les

trajectoires des utilisateurs [Hoteit 2013]

- Proposition d'une méthodologie permettant d'utiliser les données téléphoniques appliquées aux problèmes de transport [Calabrese 2013]. Cette méthodologie comprend les filtres à appliquer aux données issues de millions d'utilisateurs pour donner un maximum de pertinence aux résultats. La validation proposée repose sur une comparaison entre la mobilité déduite du modèle et des mesures odométriques réalisées par l'inspection de la sécurité de la région concernée.

1.3.3 Autres travaux utilisant la téléphonie mobile

1.3.3.1 Understanding individual human mobility patterns

Une étude à grande échelle est réalisée par Gonzalez et al. [Gonzalez 2008], publiée en 2008 dans *Nature*. Elle portait sur l'étude des déplacements de 100.000 personnes anonymes à partir de leurs téléphones mobiles durant 6 mois. Chaque fois qu'un utilisateur a émis ou reçu un appel ou un SMS l'antenne de rattachement a été enregistrée. Ceci a permis de reconstruire les déplacements quotidiens de ces personnes. Par ailleurs 206 téléphones ont été enregistrés toutes les 2 heures durant une semaine. Cette étude a permis de constater et de quantifier le haut degré de régularité de la façon dont les gens se déplacent, tant dans l'espace que dans le temps.

1.3.3.2 De la téléphonie mobile pour de l'infotrafic

Suite à la première étude de 1994 avec le projet CAPITAL, quelques travaux ont été faits pour tenter de construire de l'information sur le trafic routier en temps réel [Pan 2006][Steenbruggen 2011][Caceres 2012]. Cette nouvelle approche appliquée à l'infotrafic a notamment été utilisée pour quantifier le flux qu'engendre un événement tel qu'un concert [Reades 2007][Calabrese 2010].

1.3.3.3 Elaboration de matrices Origine/Destination

Les travaux les plus récents dans ce domaine visent à générer des matrices Origine/Destination des déplacements (matrices O/D) [Caceres 2007] [Calabrese 2011] [Frias-Martinez 2012] [Duan 2011]. Il s'agit là d'un élément important dans les domaines de l'ingénierie des transports. En effet ces données sont habituellement collectées grâce à des enquêtes très chères et donc renouvelées que très rarement (au mieux tous les dix ans). Se pose donc également le problème de la durée de validité de ces enquêtes, qui ont pour effet de figer le trafic à un instant donné. Le recours

à la téléphonie mobile pour élaborer les matrices O/D a montré son efficacité. Les résultats ont en effet été comparés aux résultats obtenus par les méthodes conventionnelles [Caceres 2007] [Calabrese 2011]. Il s'agit certainement d'une exploitation majeure à l'avenir, en raison du faible coût de collecte de la donnée.

1.3.3.4 Peut-on identifier des trajectoires de déplacements ?

En 2012, des chercheurs d'IBM (International Business Machines) ont élaboré un modèle pour optimiser les transports en commun en Côte-d'Ivoire. L'étude menée avec Orange dans le cadre du projet D4D (Data for Development) s'est appuyée sur les relevés de cinq millions d'abonnés, soit 2,5 milliards d'enregistrements collectés entre décembre 2011 et avril 2012. Ceci a donné lieu au projet nommé African Bus Routes Redrawn Using Cell-Phone Data dont une capture de la fenêtre de contrôle est donnée en Figure 1.13.



FIGURE 1.13 – Illustration du réseau de bus actuel (en jaune) et des préconisations de nouvelles lignes

Cette étude a permis d'apporter des préconisations pour faire évoluer le réseau

de transport en commun d'Abidjan, la plus grande ville du pays, comptant 539 bus, 5000 minibus et 11000 taxis. L'impact estimé de ces modifications est une réduction du temps de déplacement de 10%.

1.3.3.5 Que nous disent nos mobiles sur notre façon de vivre ?

Les opportunités qu'apportent ces nouvelles sources de données ont ces 5 dernières années intéressé également les sociologues. Ainsi des travaux s'attachent à utiliser les données téléphoniques pour identifier des comportements sociaux, tels que les mouvements selon les moments de la journée [Ahas 2010][Pulselli 2008] ainsi que la régularité des déplacements d'un jour à l'autre [Sevtsuk 2010][Song 2010]. Ici les téléphones ne sont pas utilisés pour construire une donnée directement exploitable d'un service précis mais plus pour comprendre comment nos vies ont évolué au cours des dernières années.

On notera particulièrement le travail de [Phithakkitnukoon 2012] qui portait sur plus d'un million d'utilisateurs de téléphones mobiles au Portugal. Les auteurs ont mesuré le rayon géo-social des utilisateurs, rayon dans lequel vit et se déplace l'utilisateur. Ils ont ainsi pu établir les liens existant entre le rayon géo-social et la densité d'habitation de la zone concernée, ainsi qu'avec le lien social (niveau d'activité téléphonique).

Alors que ces travaux se développent, il reste la question de la fiabilité des données utilisées et de leur représentativité. Couronné et al. [Couronne 2011] ont défini une nouvelle métrique permettant de caractériser l'intensité de la mobilité et sa fiabilité.

Nous proposons dans nos travaux exposés dans le chapitre suivant d'aller plus loin en intégrant la connaissance de la nature socio-économique du territoire. Ceci permettra de déterminer la représentativité des événements téléphoniques, en fonction notamment des caractéristiques des lieux et du temps.

1.4 Synthèse et perspectives

Les modélisations de la mobilité présentées dans ce chapitre possèdent chacune des caractéristiques propres liées à la prise en compte de la spécificité des individus, de l'environnement, de la vitesse de déplacement, etc. Selon ces caractéristiques, ils sont plus ou moins disposés à modéliser complètement et correctement la mobilité. Cette aptitude concerne aussi bien les trajectoires que les déplacements de masses de population sous la forme de nuages. Les inspirations de ces modèles sont mathématiques avec une grande part pour les probabilités. Ces modèles décrivent des phénomènes

aléatoires, communément appelés processus stochastiques.

Parmi les modèles présentés, trois approches ont été particulièrement développées dans ce chapitre : les approches **agrégées**, les approches **désagrégées** et les approches **activité-centrée**. Il s'agit dans les trois cas d'approches orientées vers les domaines du transport, domaines qui nécessitent de traduire le plus fidèlement possible les déplacements sur un territoire. Les **modèles agrégés**, tel que le modèle classique à quatre étapes, sont les modèles les plus anciens et les plus simples à appréhender. Leur principale caractéristique est de reposer sur une agrégation des flux de déplacement : l'individu est masqué derrière un volume de trafic zone à zone. Aussi, leur faible degré de finesse les rend essentiellement adaptés aux modélisations des déplacements interurbains. Parce qu'ils sont particulièrement simples à mettre en oeuvre, ce sont encore aujourd'hui des modèles largement utilisés pour définir des politiques de transport. Les **modèles désagrégés** permettent un niveau de finesse supplémentaire, en considérant l'individu lui-même lors de la modélisation des déplacements. Ceci permet une meilleure prise en compte des éléments qui caractérisent chaque individu, dans la modélisation de ses déplacements. Ainsi, des facteurs tels que le niveau de salaire peuvent être pris en compte pour améliorer le niveau de précision des déplacements modélisés. On notera qu'un niveau d'agrégation intermédiaire, entre agrégé et désagrégé, est parfois utilisé. Une agrégation par classe est alors effectuée, telle qu'un regroupement par catégorie socio-professionnelle par exemple. Enfin, toujours dans un souci de fournir des modèles les plus proches possible de la réalité, les **modèles activité-centrée** ont été introduits récemment. Alors que les deux premières classes de modèle reflètent la mobilité de façon figée dans le temps, en heure de pointe par exemple, cette dernière classe considère le triplet espace-temps-activité. Il s'agit incontestablement de la classe la plus complexe à mettre en oeuvre, mais également de la classe la plus précise. Les modèles de cette nature permettent de considérer le chainage des activités tout au long de la journée, et par là même, les déplacements impliqués au cours du temps. Ils s'appuient en général sur des enquêtes détaillées telles que les Enquêtes Ménages/Déplacements. Les travaux actuels tendent à montrer que les modèles de transport du futur seront de cette nature.

Qu'en est-il de l'utilisation de la téléphonie mobile comme marqueur de mobilité ? Nous avons vu que des travaux récents utilisent la téléphonie mobile pour construire une meilleure connaissance de la mobilité. Les premières expérimentations datant de 1994 se sont révélées peu concluantes. En effet, le faible nombre de données disponibles et la localisation au niveau de la cellule n'ont pas permis d'obtenir un niveau de précision suffisant pour construire une information pertinente. C'est à partir de 2006, en même temps que nous avons initié le projet Territoire Mobile, que les premiers résultats significatifs sont apparus avec le projet Real Time Rome mené par le MIT. Les résultats présentés dans ce projet, et les suivants menés à Graz (2007) et Singapour (2013), semblent prometteurs. En particulier, ces projets montrent l'avantage considérable qu'apporte la téléphonie en proposant un échantillon

de la population jamais étudié dans le cadre d'enquêtes, et à un coût significativement moindre. Cependant, ces études s'attachent essentiellement à présenter des résultats sous forme de calques visuels, de façon macroscopique. Il est possible ainsi de superposer différentes sources de données, telles que la position des bus ou des taxis en même temps que l'utilisation des réseaux cellulaires. Mais le niveau de précision ne permet pas de déterminer le nombre de personnes se trouvant dans un secteur géographique fin. Seule une tendance de répartition peut être perçue. Les travaux les plus aboutis, présentés cette année, explorent les possibilités existantes pour traduire quantitativement les données téléphoniques en nombre de déplacements. Ainsi, pour des applications de transport, certains travaux visent à reconstituer des matrices Origine/Destination à partir des données téléphoniques. Pour l'instant, les résultats obtenus n'ont permis que d'identifier les déplacements domicile/travail, en fonction du lieu de rattachement du téléphone le jour et la nuit. Mais ces travaux laissent entrevoir le très grand potentiel que recèlent les données téléphoniques.

Les modèles de mobilité classiques permettent donc d'estimer de façon quantitative le nombre de déplacements effectués entre différentes zones d'un territoire. Ils présentent toutefois l'inconvénient d'être approximatif en raison du faible nombre de données et de la rareté des collections car elles sont extrêmement chères. Les modèles de mobilité s'appuyant sur la téléphonie permettent quant à eux de bénéficier d'un grand nombre de données, dont l'actualisation peut se faire facilement. Ils présentent toutefois l'inconvénient d'être géographiquement très peu précis et de ne pas être facilement généralisables à l'ensemble des territoires.

Dans le chapitre suivant nous proposons d'affiner la façon dont se répartissent géographiquement les utilisateurs de mobiles avec une dynamique spatio-temporelle. Le modèle proposé intègre la description du territoire pour construire un schéma complet de la mobilité associant géographie, activités et déplacements dans un réseau cellulaire. Ce modèle ouvre la voie à une nouvelle façon d'appréhender la mobilité, afin de bénéficier des avantages proposés par les modèles classiques qui considèrent le triplet espace-temps-activité, et des avantages qu'apporte la téléphonie mobile à travers la quantité de données disponibles notamment.

La téléphonie mobile : un marqueur de présence

Dans ce chapitre est présenté un nouveau modèle de répartition dynamique de la population à partir de données téléphoniques. Nous tentons de répondre à un besoin toujours plus marqué de modéliser finement les déplacements de la population. Cette modélisation constitue l'élément essentiel pour permettre d'analyser et comprendre les besoins en services de mobilité. L'utilisation de la téléphonie comme marqueur de présence apparaît alors comme une nouvelle approche, permettant d'imager autrement la distribution spatio-temporelle de population. La première partie de ce chapitre est destinée à présenter de façon détaillée le modèle proposé. L'accent sera mis sur la façon dont s'articulent dans ce modèle les données téléphoniques d'une part, et la structure du territoire d'autre part, à travers les données géographiques et socio-économiques. Dans la deuxième partie du chapitre, une phase de validation étayée sur un cas d'étude réel est présentée associant une collectivité et un opérateur téléphonique, une force pour cette étude.

Sommaire

2.1	Introduction	40
2.2	La téléphonie mobile : une nouvelle source de données	41
2.2.1	Structure et fonctionnement d'un réseau cellulaire	41
2.2.2	Quelles données sont disponibles?	45
2.3	Un nouveau modèle de distribution de la population	47
2.3.1	Mobilité et caractérisation socio-économique du territoire	48
2.3.2	Modèle de distribution spatio-temporelle des mobiles	55
2.4	Résultats et analyse	58
2.4.1	Environnement d'analyse	58
2.4.2	Constitution des données de référence	60
2.4.3	Expérimentation et résultats	64
2.4.4	Construction d'une <i>Signature Territoriale</i>	70
2.4.5	Généralisation de la <i>Signature Territoriale</i>	73
2.5	Synthèse	73

2.1 Introduction

Dans le souci de fournir des informations de plus en plus précises et à un coût maîtrisé, de nouvelles sources d'information sont à inventer pour localiser des individus et identifier les flux de mobilité au fil du temps. Une meilleure prise en compte des déplacements des habitants d'un territoire doit ainsi permettre de déployer des offres de transports en commun de meilleure qualité, ou encore d'aménager des voiries dédiées aux cyclistes, de construire des stations de vélo ou de voitures partagées. A l'heure où les voitures électriques constituent une nouvelle façon de se déplacer, de grands investissements en matière d'infrastructures de recharges sont à prévoir. Pour accompagner toutes ces évolutions, il apparaît indispensable d'améliorer notre connaissance de la mobilité des personnes.

Les téléphones portables constituent sans aucun doute une source de données avérée pour un retour informatif sur les déplacements de personnes. Les données de réseau de téléphonie mobile : appels entrants et sortants pour chaque cellule, les transferts intercellulaires (handover), etc. fournissent de nouveaux marqueurs de la mobilité. Leur pertinence est d'autant plus forte que le téléphone portable est devenu en quelques années omniprésent [Rheingold 2002]. L'International Telecommunication Union (ITU), agence des nations unies en charge des Technologies de l'Information et de la Communication (TIC) indique que le nombre d'abonnés dans le monde est passé de moins d'1 milliard à plus de 5 milliards au cours de ces dernières années. La figure 2.1 représente cette tendance au niveau mondial. Par ailleurs, toujours selon l'ITU, en 2010 90% de la population mondiale était couverte par un signal cellulaire contre 61% en 2003. Nous assistons bien à une évolution rapide qui a des répercussions fortes sur notre vie quotidienne : travail, loisirs, déplacements, etc. [Campbell 2008]

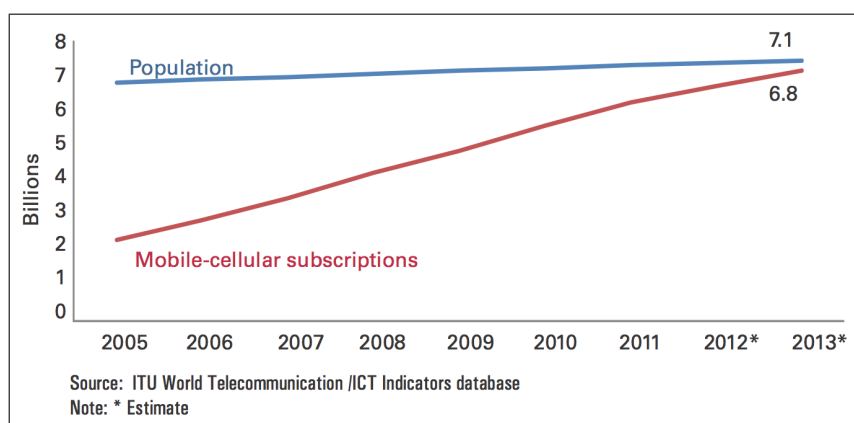


FIGURE 2.1 – Evolution du nombre d'abonnements téléphoniques dans le monde

Alors que jusqu'à présent le déploiement des réseaux mobiles s'appuyait sur une

connaissance des déplacements de la population, les téléphones portables peuvent maintenant également améliorer la compréhension sur la façon dont les personnes se déplacent. La téléphonie sans fil devient alors un traceur de mobilité, disponible tout au long de la journée, tous les jours de la semaine, pour des jours travaillés ou de vacances.

Dans ce chapitre nous proposons d'élaborer une caractérisation spatio-temporelle du terrain basée sur une analyse statistique des données de déplacements et l'usage de la téléphonie mobile. Un modèle de distribution de la population "multicritère" est ainsi présenté. Il tient compte à la fois de critères géographiques, sociologiques, événementiels, et temporels. Ce modèle est complété par un tableau de bord permettant d'imager la répartition dynamique des individus à travers une carte géographique. Nous visons à fournir un nouveau moyen d'appréhender la mobilité tout au long de la journée. Un cas d'étude a permis de valider la modélisation sur des données réelles appartenant au département du Territoire de Belfort, en France, et à la société Orange.

Enfin, on notera que le modèle que nous proposons a servi à alimenter divers travaux de recherche, pour simuler le déplacement d'individus afin de valider des protocoles de communication inter-véhiculaires ou véhicule à infrastructure [Ait Ali 2010a] [Ait Ali 2010b], pour élaborer un nouveau modèle de simulation de la mobilité [Joumaa 2009a] [Joumaa 2009b] ou encore pour proposer une nouvelle approche de configuration des réseaux cellulaires eux-mêmes, pour le développement de la 4G notamment [Chariete 2013].

La section suivante nous permet de présenter les éléments clés du fonctionnement d'un réseau de téléphonie mobile et d'identifier les éléments disponibles pour une modélisation de la mobilité. Ensuite, la section 2.3 est consacrée à la présentation de notre modèle, basé en premier lieu sur la téléphonie mobile mais intégrant également une connaissance fine de la caractérisation spatio-temporelle du terrain. Enfin, la troisième section fournit une analyse des résultats dans un contexte réel, résultats qui ont été validés dans le cadre du projet *Territoire Mobile*, labellisé par le pôle de compétitivité Véhicule du Futur, de 2006 à 2009.

2.2 La téléphonie mobile : une nouvelle source de données

2.2.1 Structure et fonctionnement d'un réseau cellulaire

La principale caractéristique d'un système radio mobile est de permettre aux personnes de se déplacer librement tout en assurant une continuité de service sur l'ensemble du territoire couvert. Pour répondre à ce besoin l'opérateur répartit des stations de base (émetteurs radio) sur le territoire à couvrir par lesquelles les appels

ou les données pourront être transférés vers les mobiles. Lorsqu'un mobile établit une communication, un événement de localisation est généré sur la station de base qui gère la communication.

Définition 1 (Station) *Une station, également appelée station de base, est un relai radio permettant d'acheminer toute forme de communication avec un mobile. A une station est toujours associée une zone de couverture qui dépend des caractéristiques de la station (paramétrage de l'antenne) et de l'environnement (relief, construction, etc.).*

Définition 2 (Site) *Un site est le lieu géographique repéré par ses coordonnées où est placée une station. Un même site peut regrouper plusieurs stations permettant de couvrir des orientations différentes, de proposer diverses technologies (UMTS, GSM) et éventuellement plusieurs opérateurs.*

Définition 3 (Cellule) *On appelle cellule la zone de couverture radio d'une station. Il s'agit de la portion de territoire sur laquelle les communications peuvent être acheminées par la station associée.*

Pour permettre aux personnes d'être continuellement connectés lorsqu'elles se déplacent, la majeure partie du territoire est couverte par plusieurs cellules. Afin de calculer, visualiser et gérer les recouvrements entre cellules, les opérateurs utilisent des méthodes et des logiciels de planification de réseaux. En particulier, ils ont recours à des modèles de propagation des ondes radio leur permettant de déterminer en chaque lieu quelles sont les antennes susceptibles d'acheminer une communication, et avec quelle probabilité. C'est ce que nous appellerons la probabilité de prise de communications (PPC)¹. La Figure 2.2 montre sur un schéma théorique comment ces superpositions cellulaires opèrent, avec sous forme maillée les PPC de trois antennes.

La Figure 2.3 est une illustration de l'enchevêtrement réel de cellules où chaque couleur représente une cellule et l'intensité de la couleur est le reflet de l'intensité de la probabilité de prise de communication (PPC) de la station.

Lorsqu'un mobile se déplace, il détecte successivement les stations de base en fonction de la couverture de chacune d'elles, de leur niveau de saturation et de paramètres sur les antennes qui gèrent l'admission des terminaux. Ce déplacement a un effet réduit sur le réseau si le téléphone est en veille (pas d'appel, pas de transfert de données) car celui-ci fait une mise à jour de la localisation par groupe de cellules (Location Area - LA). Par contre dans le cas où le téléphone est en cours de communication, tout changement de cellule implique un dialogue avec le

1. Ce concept est la propriété d'Orange et ne sera pas détaillé dans cette thèse.

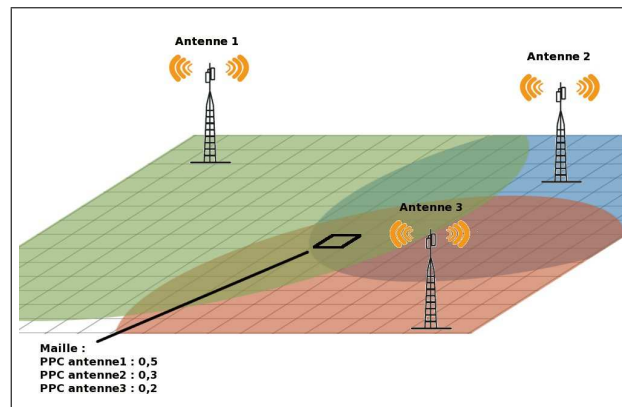


FIGURE 2.2 – Maillage et PPC

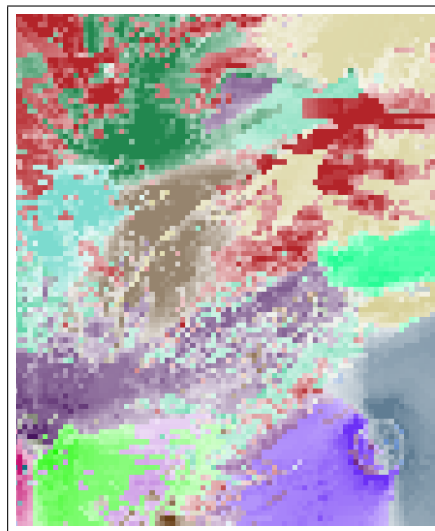


FIGURE 2.3 – Couverture radio

réseau pour changer la connexion du terminal d'une station à une autre (transfert intercellulaire ou Handover).

Définition 4 (Handover) *On appelle Handover ou Handoff (HO) le transfert intercellulaire d'un mobile.*

Plusieurs événements sont générés sur le réseau lorsqu'un mobile en cours de communication se déplace de façon à le raccrocher à la station la plus adaptée au contexte. Ainsi un terminal peut être tracé au sein de plusieurs cellules en fonction de sa mobilité en cours d'appel et générer un ensemble d'événements sur le réseau. Les paragraphes suivants illustrent ce phénomène au travers d'un exemple en deux étapes :

1. La personne 1 initie, depuis la cellule A, un appel à la personne 2 présente dans la cellule B (Figure 2.4). Un appel sortant est comptabilisé pour la station ou antenne A, et un appel entrant est comptabilisé pour la station B.

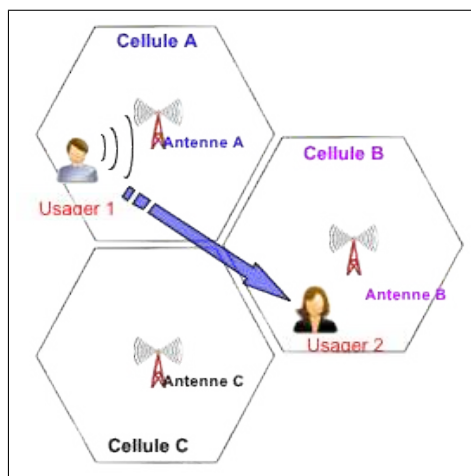


FIGURE 2.4 – Génération d'un événement de communication

2. La personne 1 se déplace alors, pendant la communication, dans la cellule C. Le mobile de la personne 1 doit donc opérer un transfert intercellulaire (Figure 2.5). Ce type d'événement fait comptabiliser un HO sortant pour la cellule A et un HO entrant pour la cellule C.

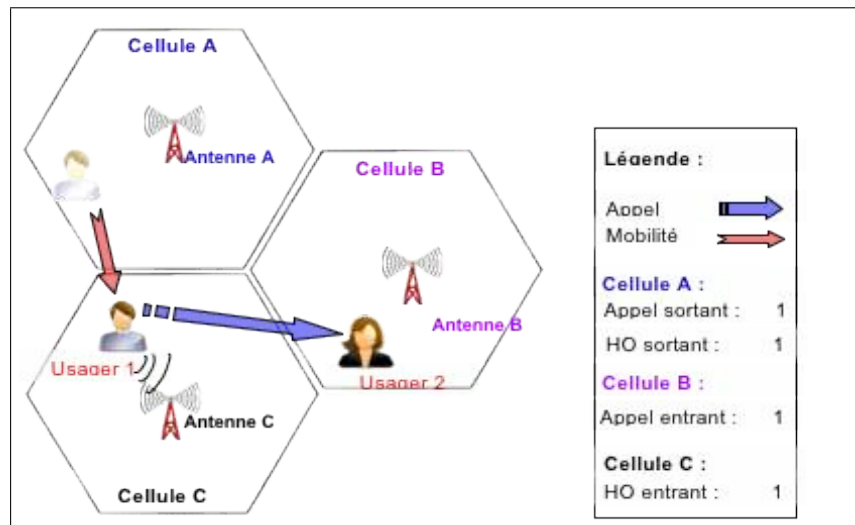


FIGURE 2.5 – Génération d'un événement de transfert intercellulaire

Le HO permet de localiser la présence de personnes, par contre ce n'est pas un indicateur de la direction des flux. Lorsqu'un HO entrant est comptabilisé pour une cellule, nous pouvons affirmer qu'une personne est entrée dans celle-ci, mais nous ignorons sa provenance. De même, si nous comptabilisons un HO sortant pour une cellule, nous pouvons affirmer qu'une personne a quitté cette cellule, mais nous ignorons sa destination. Bien entendu, le réseau identifie le besoin de transfert intercellulaire mais cette information n'est pas stockée, seuls les événements anonymes d'entrée ou de sortie de cellules sont archivés.

Nous ne détaillerons pas ici les mécanismes impliqués dans le choix de la station de rattachement, ni dans le déclenchement d'un transfert intercellulaire. Le lecteur souhaitant plus de détails pourra se référer au livre [Tabbane 2002].

2.2.2 Quelles données sont disponibles ?

Tout ce qui se passe sur un réseau cellulaire n'est pas enregistré pour des questions de technologie, de volume de données, de stratégie opérateur, de protection des libertés (CNIL), etc. Pour des raisons de limitation de ressources sur le lien radio, il n'est pas possible de connaître à tout moment la station de rattachement de chaque téléphone. Toutefois les opérateurs captent et conservent les événements suivants :

- les communications entrantes (information de communications) : ce sont les appels/SMS/transfert de données reçus par la personne se trouvant dans la cellule, c'est un marqueur de présence.
- les communications sortantes (information de communications) : ce sont les

appels/SMS/transfert de données émis par la personne se trouvant dans la cellule, c'est un marqueur de présence.

- les "handover" (information de signalisation) : c'est un marqueur de mobilité étant donné qu'il s'agit de communications ayant débuté dans une cellule et qui se poursuivent dans une autre cellule.

Les données de trafic et de signalisation se présentent en général sous la forme d'une matrice individus/variables où les individus sont les stations, caractérisées par leur identifiant, assurant la couverture de la zone étudiée, et les variables sont les événements. Dans notre travail sur l'utilisation de ces données pour modéliser la mobilité, nous avons regroupé les événements par 1/4 d'heure de 6h00 à 23h45, soit 71 plages horaires. Ce pas est tout à fait ajustable pour augmenter ou diminuer la précision temporelle du modèle. Voici la liste des données utilisées par le modèle que nous proposons :

- Le nombre d'appels sortants. Celui-ci indique le nombre de personnes se trouvant dans la cellule c durant le quart d'heure t ayant initié une communication depuis leur téléphone mobile.
- Le nombre d'appels entrants. Celui-ci indique le nombre de personnes se trouvant dans la cellule c durant le quart d'heure t ayant reçu une communication sur leur téléphone mobile.
- Le nombre d'appels total = le nombre d'appels entrants + le nombre d'appels sortants.
- Les handover sortants. Ils indiquent le nombre de personnes quittant la cellule c pour une cellule voisine durant le quart d'heure t .
- Les handover entrants. Ils indiquent le nombre de personnes entrant dans la cellule c depuis une cellule voisine durant le quart d'heure t .
- La somme et la différence des handover sortants et entrants. Ces calculs seront utiles pour la détection de pôles plus ou moins attracteurs sur la zone géographique étudiée.

De plus, pour faire le lien entre les communications et le territoire, nous utiliserons une matrice de couverture radio pour l'ensemble des antennes du territoire étudié. Les données de couverture radio utilisées ici sont des données maillées à un pas de 25m. Cette matrice contient pour chaque maille la valeur de puissance reçue et provenant de l'antenne considérée. Elles permettent de localiser les zones couvertes par chaque antenne et de connaître la PPC associée. Enfin cette matrice permet par superposition de faire le lien entre les cellules et les données géographiques (parcelles, communes, bâtiments, routes) du territoire.

2.3 Un nouveau modèle de distribution de la population

Dans cette section nous présentons notre approche permettant de modéliser les déplacements, qui s'appuie sur trois éléments essentiels :

- l'utilisation des réseaux téléphoniques sans fil
- les caractéristiques du sursol
- les objectifs de déplacements des personnes

Ce modèle a été mis au point et expérimenté dans le cadre d'un projet labellisé par le Pôle de Compétitivité "Véhicule du Futur" [Lamrous 2008] avec comme partenaire le SMTC Belfort, la société Orange, l'Agence d'Urbanisme et les collectivités territoriales du Territoire de Belfort. La vocation du projet est de proposer une modélisation et une vision via un logiciel informatique des déplacements de l'ensemble de la population d'un territoire afin de planifier des services de transport en commun. Cette approche constitue ainsi une nouvelle façon d'appréhender l'environnement urbain. Ce modèle se place dans la catégorie des modèles basés sur l'activité couplé avec des données issues des réseaux de téléphonie cellulaire.

Afin de permettre de visualiser les résultats, nous avons développé la plateforme logicielle "geoLogic" proposant un tableau de bord ayant pour fond une carte géographique (cf. Figure 2.6) à laquelle se superpose plusieurs calques de couleur pour scruter les déplacements des individus. Ce tableau de bord permet de visualiser les déplacements sur le territoire couvert sur une journée complète au quart d'heure près et à un pas de 25 mètres. Tous les moyens de zoom et de sélection sont disponibles à l'utilisateur. Une présentation plus détaillée de geoLogic est proposée dans le chapitre 4. Nous utiliserons des sorties graphiques de geoLogic pour illustrer des composants du modèle de mobilité de ce chapitre.

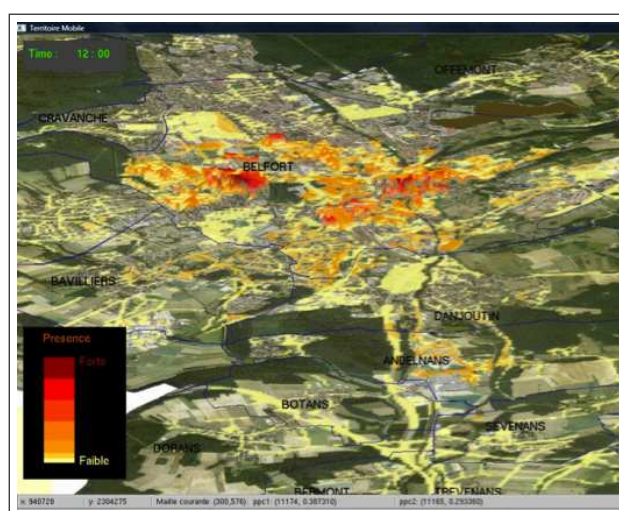


FIGURE 2.6 – Cartographie temporelle de la distribution de la population belfortaine

Dans les projets issus de la littérature, le niveau de précision est celui de la couverture radio, or ces tailles peuvent varier de façon significative. Alors que dans un contexte urbain dense, elles peuvent couvrir un rayon de 300m, elles peuvent hors agglomération s'étendre à des rayons de plusieurs kilomètres. Selon les besoins, ce niveau de découpage du territoire peut être suffisant, mais dans beaucoup de situations un niveau plus fin est nécessaire. L'introduction ici d'une connaissance de la topographie permettra d'estimer de façon probabiliste comment les mobiles se répartissent à l'intérieur des zones de couverture des antennes considérées, grâce aux données socio-économiques et géographiques.

2.3.1 Mobilité et caractérisation socio-économique du territoire

Les informations sur le trafic téléphonique reportent le volume global de trafic écoulé (communications et signalisations) sur chacune des cellules. A ce stade la connaissance du trafic est imprécise et macroscopique puisque les cellules de l'opérateur ne permettent pas de savoir où le mobile se trouve à l'intérieur de celle-ci. Les procédés de localisation que nous proposons permettent de calculer la distribution non uniforme des appels à l'intérieur de chacune des cellules en exploitant toutes les informations spatiales et temporelles disponibles. En général, de telles données décrivant la géographie d'un territoire sont disponibles dans les services des collectivités traitant des Systèmes d'Information Géographique (SIG). En France, ces données sont souvent issues de l'IGN qui géo-référence de nombreuses caractéristiques ; elles sont complétées par les collectivités elles-mêmes qui disposent d'une connaissance très précise de leur territoire. Notons également qu'à l'heure du développement de l'Open Data, des outils tels que OpenStreetMap constituent une nouvelle source de données en plein essor.

A partir des données géographiques et socio-économiques, l'objectif est d'établir un profil d'attraction temporelle propre à chaque maille du territoire. Ce profil ne peut être envisagé que dans un contexte dynamique. A chaque maille est ainsi attribué un vecteur décrivant son attraction au cours du temps (poids d'attraction spatio-temporelle). De cette manière, il est possible d'associer différents profils pour différents types de jours (semaine, weekend, vacances, etc.). La Figure 2.7 décrit le processus complet permettant d'aboutir à cette pondération. Ainsi, à partir d'une description vectorielle du territoire, une première étape consiste à appliquer une grille de maillage à ces données. Chaque maille contient alors la proportion que représente chaque classe de terrain qui la compose. Dans la figure chaque carré de couleur en haut à droite indique la classe dominante. Un profil horaire à trois états est alors associé aux classes de sursol, pour décrire en fonction du temps si une classe est attractive de façon forte (Heure Pleine - HP), faible (Heure Creuse - HC) ou non attractive ($\emptyset$). Ce profil est alors couplé à des données statistiques issues de l'INSEE pour attribuer un poids d'attraction spatio-temporelle final à chaque maille. Ce poids est donc variable dans l'espace et dans le temps.

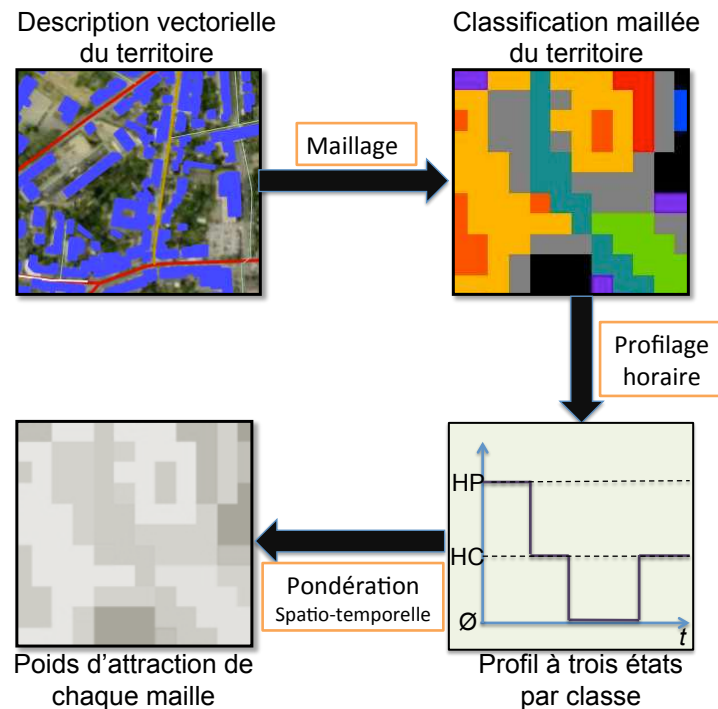


FIGURE 2.7 – Process complet de la pondération spatio-temporelle du territoire

Les sections suivantes décrivent de façon détaillée ce processus qui mène à la description du territoire. Nous présenterons ensuite (cf. section 2.3.2) comment cette description est couplée au réseau de communications (PPC) pour distribuer les mobiles dans les mailles du territoire.

2.3.1.1 Données spatiales socio-économiques et géographiques

Le territoire est caractérisé par l'information dite de sursol qui indique l'occupation du sol (route, bâti, eau, etc.). Une approche hiérarchique nous permet de structurer cette information en différentes classes de terrain.

Le terrain est dans un premier temps divisé en grandes catégories (cf. Figure 2.8) : commerce, activité, habitat, etc. Seuls les lieux pouvant significativement définir une concentration de population seront importants dans la suite du travail mais tous les types sont pris en compte. Ces catégories sont couramment utilisées dans les études traitant de la mobilité. C'est le cas des enquêtes ménages/déplacements par exemple, qui recensent les motifs de déplacement directement liés à ces catégories. La catégorie *commerce* contient ainsi toutes activités non professionnelles de type achat, loisir, etc. La catégorie *activité* contient toutes activités professionnelles ou liées à la scolarité. Et la catégorie *voirie* concerne toutes activités liées aux déplacements.

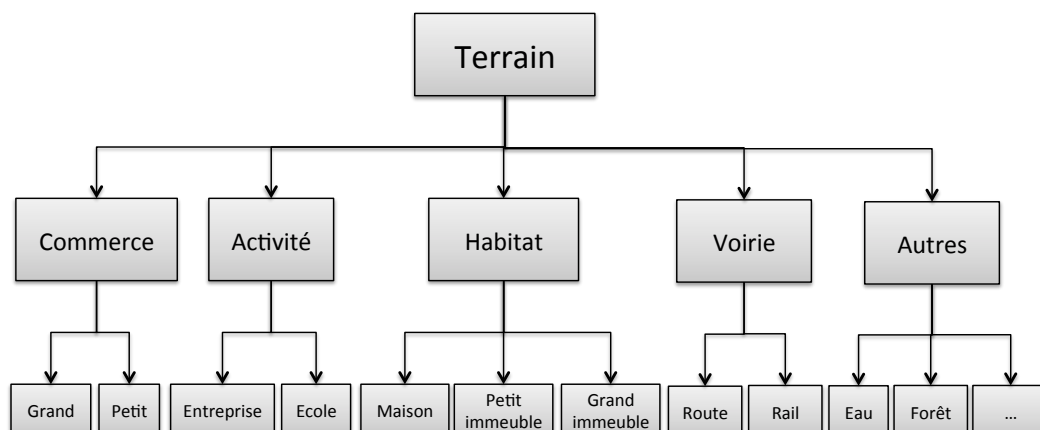


FIGURE 2.8 – Organisation du sursol

Chaque catégorie est subdivisée en plusieurs classes afin d’apporter plus de finesse lors de la distribution des mobiles. L’habitat sera par exemple décomposé en maison, petit immeuble et grand immeuble, ce qui permettra lors d’affectation de personnes au domicile d’affecter plus ou moins de personnes selon la taille du bâtiment. De la même façon, les commerces sont distingués selon leur taille : un petit commerce de proximité n’aura pas le même pouvoir d’attraction qu’un centre commercial.

Ainsi, la description du territoire selon cette arborescence est tout d’abord disponible sous forme vectorielle, puis la première étape du traitement est le maillage et la classification de ces données. Cette opération relève d’un géo-traitement à base d’intersection de polygones. Il ne sera pas détaillé ici mais est publié dans [Lamrous 2008]. Voyons à présent comment est attribué un profil horaire à chaque classe de terrain, soit comment inclure une dimension temporelle aux données spatiales.

2.3.1.2 Un profil temporel des classes de terrain à trois états

Selon ses caractéristiques socio-économiques, un lieu attire des personnes différemment à chaque moment de la journée et d’un jour à l’autre. Pour rendre compte de ce phénomène nous procédons à un découpage du temps en plages horaires qui sera associé au découpage géographique du territoire en mailles de l’étape précédente. Une bonne connaissance de l’occupation du sursol du territoire nous permettra ainsi de distribuer les mobiles rattachés à un relai téléphonique différemment selon l’heure de la détection. Nous pouvons citer l’exemple d’un plan d’eau urbain interdit à la baignade et à la navigation, ou encore d’un établissement scolaire pouvant recevoir des personnes en journée mais pas la nuit. Nous élaborons ainsi un schéma temporel de caractérisation du sursol.

Chaque classe élémentaire de terrain représentée Figure 2.8 a un pouvoir d'attraction à trois états qui dépend du jour et du moment de la journée. Nous définissons un tel profil horaire pour une classe de la manière suivante : chaque plage horaire Δt (généralement de 15 minutes) peut prendre la valeur 0 lorsqu'un nombre non significatif de personnes peut se trouver dans cette classe, 1 lorsque la classe est considérée en heure creuse (faiblement fréquentée) et 2 lorsqu'un nombre important de personnes peut se trouver dans cette classe. Cette approche à trois niveaux (absence, faible présence, forte présence) apporte une grande souplesse dans la construction et dans l'évolution de la matrice d'appartenance horaire. Il s'agit donc d'un compromis entre une approche binaire du tout ou rien et une approche continue qui complexifierait la construction de ces profils.

	6h	6h15	6h30	6h45	7h	...	
Petits commerces	0	0	0	0	0	0	Commerce
Grands commerces	0	0	0	0	0	0	
Entreprise	0	0	0	0	0	0	Activité
Ecole	0	0	0	0	0	0	
Maisons <10m	2	2	2	2	2	2	Habitat
Petits immeubles < 20m	2	2	2	2	2	2	
Grands immeubles > 20m	2	2	2	2	2	2	
Routes	1	1	1	1	1	1	Voie
Rail	0	0	0	0	1	1	
Jardins / Parcs	0	0	0	0	0	0	Autres
Forêt	0	0	0	0	0	0	
Eau	0	0	0	0	0	0	

FIGURE 2.9 – Table de profilage temporel du sursol

Un extrait de la table de profilage temporel du sursol à trois états est donné Figure 2.9 de 6h à 7h (en colonnes). Chaque ligne correspond à une classe de terrain pour un lundi d'une période travaillée. Le profil change d'un jour à l'autre de la semaine et éventuellement d'une semaine à l'autre pour refléter au plus près la réalité.

Plus formellement, à la table de profilage temporel est associée la matrice A , qui est définie de la façon suivante :

$$A = (a_{ij})_{1 \leq i \leq p, 1 \leq j \leq n}$$

avec :

p : le nombre de classes élémentaires de sursol

n : le nombre de plages horaires dans une journée

$$a_{ij} = \begin{cases} 2 & \text{si la classe } i \text{ est considérée en} \\ & \text{HP à la période } j \\ 1 & \text{si la classe } i \text{ est considérée en} \\ & \text{HC à la période } j \\ 0 & \text{sinon (pas de personne affectée)} \end{cases} \quad (2.1)$$

Cette étape nous a permis de définir une matrice décrivant le profil temporel à trois états de chaque classe de terrain. Il s'agit maintenant de convertir cette matrice en poids d'attraction spatio-temporelle variables dans le temps pour chaque classe, avant d'en déduire un poids d'attraction pour chaque maille.

2.3.1.3 Un poids d'attraction spatio-temporelle pour le territoire

Nous abordons maintenant la phase qui permet de définir la façon dont la nature du sursol influence la répartition des mobiles sur le territoire. Cette répartition se déroule en deux temps. Tout d'abord, il s'agit de déterminer un poids d'attraction global pour chaque classe définie dans le tableau 2.8; puis une deuxième étape consiste à distribuer les poids globaux sur les plages de la journée en tenant compte des trois niveaux possibles de présence de personnes (absence, faible présence, forte présence).

Répartition globale des poids de chaque classe de terrain Pour procéder à cette étape nous avons considéré les données statistiques fournies par l'INSEE pour la France sur la période 2007-2009, afin de définir une estimation du temps moyen passé dans chaque classe. Remarquons que ces poids constituent une donnée d'entrée de l'algorithme, et peuvent être modifiés, pour tenir compte de caractéristiques propres à chaque territoire.

La période de nuit n'est pas prise en compte, toutes les personnes sont considérées appartenant à la classe habitat. Pour la période de la journée (6h - 0h), nous considérons le pourcentage de temps moyen passé par catégorie (ramené à l'ensemble de la population) tel que décrit dans le tableau 2.1. La première ligne du tableau issue de l'INSEE donne la répartition entre les catégories commerce, activité, habitat, voirie et autre. On remarque que 18% du temps (catégorie autre) ne peut être attribuée à une classe identifiée. On choisit alors de répartir cette proportion de temps entre les classe identifiées au prorata de celles-ci, ce qui donne la nouvelle répartition de la deuxième ligne du tableau. La répartition au prorata est arbitraire mais neutre; elle pourrait être faite autrement selon le territoire. Cette répartition est définie pour un jour travaillé (hors vacances scolaires).

Pour définir la proportion de temps passé dans chaque classe de sursol du

	Commerce	Activité	Habitat	Voirie	Autres
Enquête brute	25%	25%	22%	10%	18%
Enquête modifiée	30,49%	30,49%	26,83%	12,2%	0%

TABLE 2.1 – Temps moyen passé par catégorie

Catégorie	% catégorie	Classe de sursol g pour <i>ground</i>	% classe	Répartition globale ω_g
Commerce	30,49%	petit commerce	30%	9,15%
		grand commerce	70%	21,34%
Activité	30,49%	entreprise	67%	20,43%
		école	33%	10,06%
Habitat	26,83%	maison	6,25%	1,68%
		pt immeuble <20m	31,25%	8,38%
		gd immeuble >20m	62,5%	16,77%
Voirie	12,2%	route	75%	9,15%
		rail	25%	3,05%

TABLE 2.2 – Répartition de la population par thème

tableau 2.8, les catégories sont elles-mêmes réparties selon des proportions guidées par des statistiques de l'INSEE. Le tableau 2.2 synthétise cette répartition. Les deux premières colonnes qui décrivent les catégories sont directement issues du tableau 2.1. Les colonnes suivantes concernent les classes de sursol et donnent, d'abord la répartition interne entre les classes de sursol dans chaque catégorie, puis la répartition globale par rapport à l'ensemble du temps passé qui croise les catégories et les classes de sursol. Prenons l'exemple de la catégorie *activité* qui représente 30,49% du temps passé. 33% de ce temps est attribué à la classe école, et 67% à la classe entreprise, soit 33% de 30,49% = 10,06% de temps pour la classe école et de même 20,43% de temps pour la classe entreprise. La répartition attribuée aux classes de sursol provient d'une estimation de la proportion de population concernée ; elle doit être affinée par territoire d'étude.

La répartition du temps global par classe de sursol est ainsi donnée dans la dernière colonne du tableau 2.2. Il s'agit des poids ω_g de chaque classe de sursol g qu'il faut à présent répartir dans le temps.

Répartition temporelle de chaque classe de terrain La dernière étape de la caractérisation temporelle des classes de terrain consiste à déterminer le poids d'attraction de chacune d'elles tout au long de la journée pour un type de jour donné. C'est ce poids final qui permettra de distribuer statistiquement les mobiles à

l'intérieur de la cellule radio. Cette étape consiste donc à répartir les poids calculés globalement pour chaque classe sur l'ensemble des plages horaires de la journée, en s'appuyant sur la matrice A qui décrit le profil des classes en trois états.

Soient les définitions suivantes :

G l'ensemble des classes de sursol (G pour *ground*)

g une classe de sursol, $g \in G$

ω_g^{hp} le poids de la classe g en heure pleine

ω_g^{hc} le poids de la classe g en heure creuse

ω_g le poids global de la classe g calculé précédemment

$\omega_g^0 = 0$ le poids de la classe g en heure d'absence (ce poids peut être mis à une valeur différente de 0 si on veut traduire une très faible présence).

n_g^{hp} le nombre de plages horaires en heure pleine pour la classe g

n_g^{hc} le nombre de plages horaires en heure creuse pour la classe g

n_g^0 le nombre de plages horaires d'absence pour la classe g

Pour une classe donnée g , la matrice A contient les valeurs 0, 1 ou 2 selon les plages horaires. Le poids ω_g est alors réparti linéairement sur l'ensemble des plages absence, HC et HP de la journée, avec une proportion x à définir entre HC et HP selon le territoire.

Calculons à présent les poids ω_g^{hp} et ω_g^{hc} de la classe g en heure pleine et en heure creuse respectivement. Ils vérifient :

$$\begin{cases} \omega_g^{hc} n_g^{hc} + \omega_g^{hp} n_g^{hp} + \omega_g^0 n_g^0 = \omega_g \\ \omega_g^{hp} = x * \omega_g^{hc} \end{cases} \quad (2.2)$$

Avec $\omega_g^0 = 0$, il s'en suit que les poids d'attraction des classe de sursol sont donnés par :

$$\begin{cases} \omega_g^{hc} = \frac{\omega_g}{x n_g^{hp} + n_g^{hc}} \\ \omega_g^{hp} = x * \frac{\omega_g}{x n_g^{hp} + n_g^{hc}} \end{cases} \quad (2.3)$$

Chaque maille m est à présent caractérisée par un vecteur d'attraction $\omega_m = [\omega_m^1, \dots, \omega_m^n]$, construit à partir des classes de sursol couvrant la maille en respectant la proportion de couverture de chacune. Chaque composante ω_m^t , avec $t \in [1..n]$, est définie par :

$$\omega_m^t = \sum_{g \in G} \omega_g^t \cdot p_{m,g} \quad (2.4)$$

avec :

ω_m^t le poids d'attraction spatio-temporelle de la maille m à l'instant t

ω_g^t le poids de la classe de sursol g à l'instant t

$$\omega_g^t = \begin{cases} \omega_g^{HP} & \text{si } a_{g,t} = 2 \\ \omega_g^{HC} & \text{si } a_{g,t} = 1 \\ \omega_g^{\emptyset} & \text{si } a_{g,t} = 0 \end{cases} \quad (2.5)$$

$p_{m,g}$ la proportion de classe de sursol g dans la maille m

L'étape suivante de notre modèle consiste à distribuer les mobiles sur un territoire.

2.3.2 Modèle de distribution spatio-temporelle des mobiles

Dans cette partie nous rassemblons les éléments présentés jusqu'ici, les données téléphoniques et les poids d'attraction spatio-temporelle, pour aboutir à la répartition des mobiles sur l'ensemble du territoire. La Figure 2.10 illustre ce qui est présenté dans cette section.

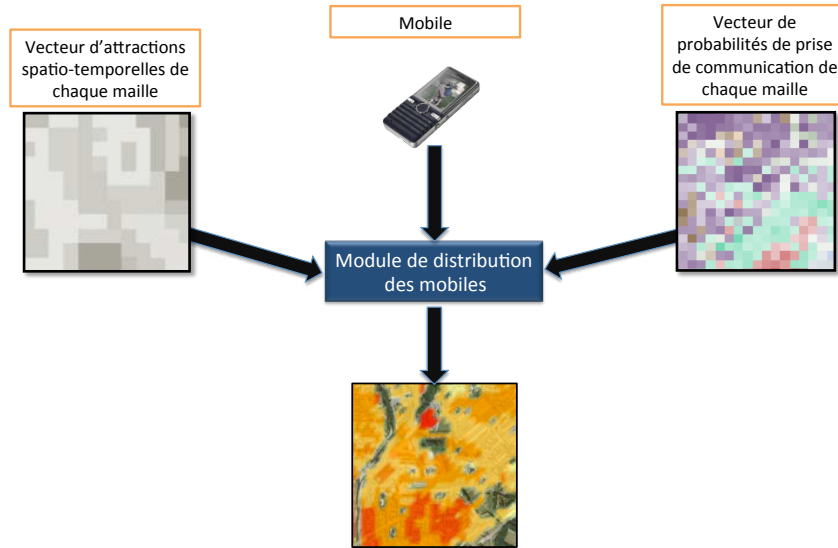


FIGURE 2.10 – Module de distribution des mobiles

Nous procéderons à une répartition en deux temps. Dans un premier temps les poids d'attraction spatio-temporelle ne sont pas intégrés. Ceci permettra de vérifier qu'une répartition des mobiles sans prise en compte du territoire a une forte tendance à être homogène. Puis dans un deuxième temps les poids d'attraction seront intégrés pour réaliser une distribution des mobiles plus réaliste.

2.3.2.1 Modèle de distribution de mobiles selon la cartographie des Probabilités de Prise de Communication (PPC)

Comme nous l'avons vu, les cellules radio sont découpées en mailles carrées qui définissent le niveau de précision que l'on souhaite atteindre. Ainsi, pour chaque maille est générée une structure multicouche qui intègre toutes les cellules qui la recouvre ainsi que les PPC associées. De la même façon, la journée est découpée en plages horaires de Δt . Pour l'application numérique nous fixons le pas de maillage pdm à 25m de côté et le pas de temps Δt à 15 minutes. La répartition des mobiles sur les mailles d'une cellule radio pendant l'intervalle Δt est faite proportionnellement aux PPC de chaque maille. Ainsi la densité des mobiles sur chaque maille m est déterminée en fonction de la représentation de la PPC de la maille relativement à l'ensemble des PPC de toutes les mailles m' de la cellule. Pour une station s sur une maille m à l'instant t (pour simplifier la notation, t sera assimilée à la période de temps Δt) nous avons la répartition suivante :

$$n_{m,s}^t = \frac{ppc_{m,s}}{\sum_{m' \in C_s} ppc_{m',s}} \cdot n_s^t \quad (2.6)$$

avec :

$n_{m,s}^t$ le nombre de mobiles rattachés à la station s réparti sur la maille m à l'instant t

C_s l'ensemble des mailles appartenant à la zone de couverture de la station s i.e. l'ensemble des mailles du territoire pour lesquelles la PPC associée à s est non nulle

$ppc_{m,s}$ la probabilité lorsqu'un mobile est sur la maille m que son appel soit acheminé par la station s

$ppc_{m',s}$ la probabilité lorsqu'un mobile est sur la maille m' que sa communication soit acheminée par la station s

n_s^t le nombre de mobiles rattachés à la station s à l'instant t

Remarquons dans le modèle (2.6) que les PPC sont invariantes dans le temps, seule l'apparition des mobiles est variable.

Nous avons ainsi défini la distribution des mobiles à l'intérieur d'une cellule radio indépendamment des propriétés spatio-temporelles du sursol. Ce même traitement est appliqué à l'ensemble des cellules, puis nous définissons pour chaque maille un cumul de présences issues de l'ensemble des cellules qui contiennent cette maille i.e. pour lesquelles la PPC est non nulle. Ainsi pour chaque maille m à l'instant t nous obtenons la répartition globale de l'ensemble des mobiles de la façon suivante :

$$n_m^t = \sum_{s \in S_m} n_{m,s}^t \quad (2.7)$$

avec : S_m l'ensemble des stations couvrant, PPC non nulle, la maille m .

La Figure 2.11 illustre la répartition des mobiles sur un territoire sans prise en compte des propriétés du sursol. Nous apercevons les limites d'une telle approche qui est basée uniquement sur les caractéristiques du réseau à travers la PPC. En effet, même si en ingénierie de réseaux on cherche à avoir une bonne qualité de signal dans les lieux de forte fréquentation, la réciproque n'est toutefois pas vraie. Ce n'est pas parce qu'il y a un bon signal (une bonne PPC) en un lieu donné que nécessairement il y a beaucoup de personnes en ce lieu.

La Figure 2.11 illustre bien ce phénomène. On aperçoit notamment que, sans prise en compte du sursol, le modèle distribue autant de mobiles dans un plan d'eau que dans les habitations voisines, de façon très uniforme.

2.3.2.2 Modèle de distribution spatio-temporelle intégrant les propriétés du sursol

Le modèle vu jusqu'ici ne tenait compte que de la couverture radio. L'étape suivante consiste à intégrer dans le calcul les vecteurs de poids d'attraction spatio-temporelle ω_m de chaque maille tels que définis dans la section 2.3.1.3.

En appliquant cette pondération au modèle de distribution des mobiles précédent, nous obtenons pour la maille m à l'instant t l'affectation suivante :

$$n_m^t = \sum_{s \in S_m} \left(\frac{ppc_{m,s} \cdot \omega_m^t}{\sum_{m' \in C_s} ppc_{m',s} \cdot \omega_{m'}^t} \cdot n_s^t \right) \quad (2.8)$$

L'intégration des poids d'attraction spatio-temporelle avec le modèle de distribution de mobiles a pour effet de réorganiser la distribution des mobiles de façon plus réaliste. Pour s'en rendre compte la Figure 2.12 présente le résultat de l'application

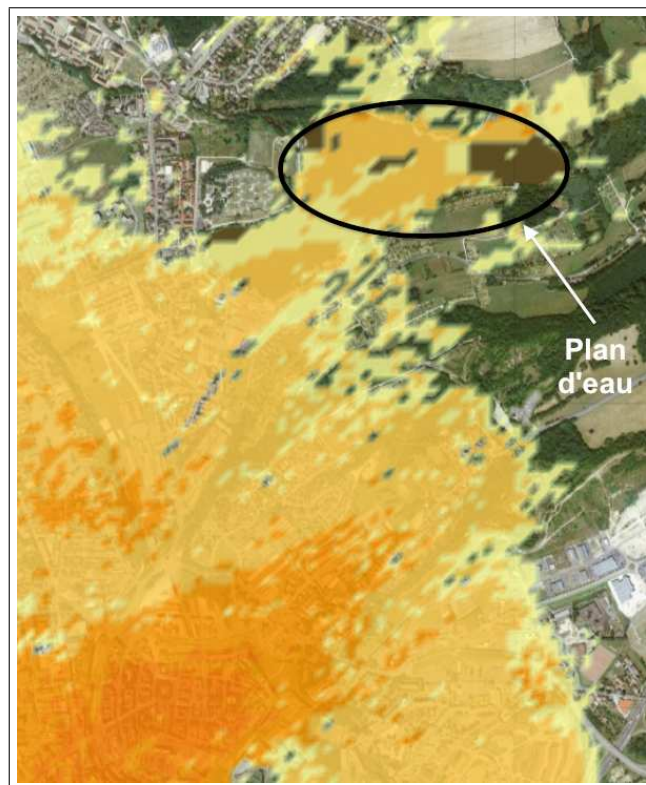


FIGURE 2.11 – Distribution des mobiles à un instant t sans prise en compte des propriétés du sursol

de ce modèle, à gauche sans prise en compte de ces poids et à droite en les intégrant. En reprenant notre exemple précédent, nous pouvons ainsi remarquer que la partie de la population qui était affectée à un plan d'eau lorsque seules les PPC étaient considérées se retrouve réaffectée dans des mailles plus appropriées.

La visualisation de la répartition des mobiles avec et sans prise en compte des poids d'attraction spatio-temporelle donne une nouvelle visibilité de la présence des personnes. Toutefois, une analyse numérique des résultats doit permettre de vérifier cette approche. Ce point est abordé dans la section suivante.

2.4 Résultats et analyse

2.4.1 Environnement d'analyse

Une bonne analyse de la mobilité nécessite de disposer d'un outil efficace pour représenter les résultats des calculs sur les déplacements. Une plate-forme logicielle capable de gérer à la fois des données spatiales telles que des données géographiques

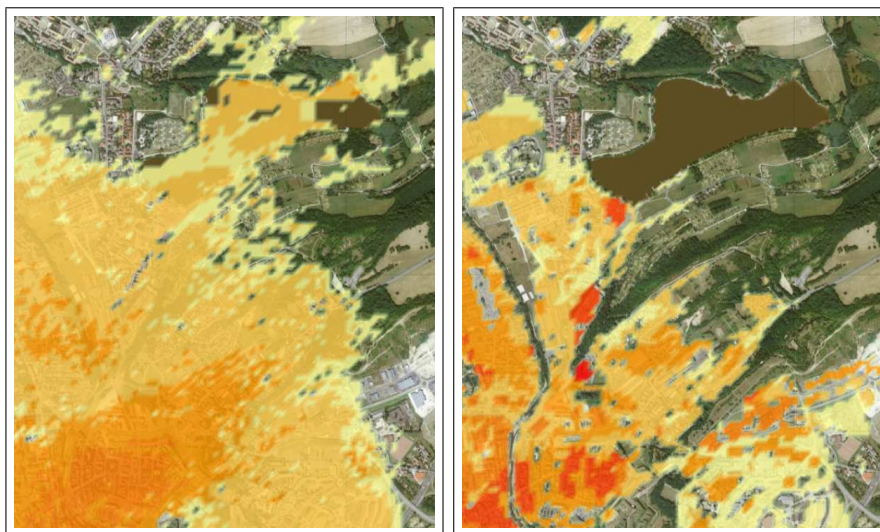


FIGURE 2.12 – Distribution des mobiles à un instant t sans et avec prise en compte des poids d'attraction spatio-temporelle

ou de localisation, ainsi que des données temporelles reflétant la mobilité apparaît alors comme un élément essentiel.

De nombreux outils existent dans le domaine des SIG (Systèmes d'Information Géographiques) pour représenter des données géo-localisées. Toutefois la prise en compte simultanément de l'espace et du temps est rare.

Afin de permettre une grande souplesse d'utilisation nous avons donc développé notre propre plate-forme de représentation des résultats. Elle dispose des fonctionnalités habituellement présentes dans les outils SIG classiques, telles que l'affichage de fonds de carte ainsi que de données vectorielles ou ponctuelles de type shapefile (cf. Figure 2.14). Il est par ailleurs possible de se déplacer dans cet environnement, de zoomer sur un lieu précis, etc. Chaque objet chargé, tel que les limites de communes ou les contours des bâtiments, constitue un calque qu'il est possible d'afficher ou masquer selon les besoins.

Les données peuvent être statiques ou dynamiques. Alors que les éléments topographiques tels que les axes routiers ou les bâtiments ne varient pas dans le temps, mais leurs propriétés quant à elles peuvent évoluer au cours de la journée. En particulier un supermarché sera toujours au même endroit mais aura un pouvoir d'attraction différent à 9h, 16h ou 23h et cela suivant le jour considéré. De même la localisation de la population, tant au niveau macroscopique, sous forme de masse, qu'au niveau microscopique, sous forme ponctuelle ou de trajectoire, varie dans le temps. Pour rendre compte de cette dynamique l'application que nous proposons permet de définir la journée de l'étude et le moment de la journée à observer. Enfin il est possible de faire défiler le temps entre 2 plages horaires, de faire varier la vitesse de

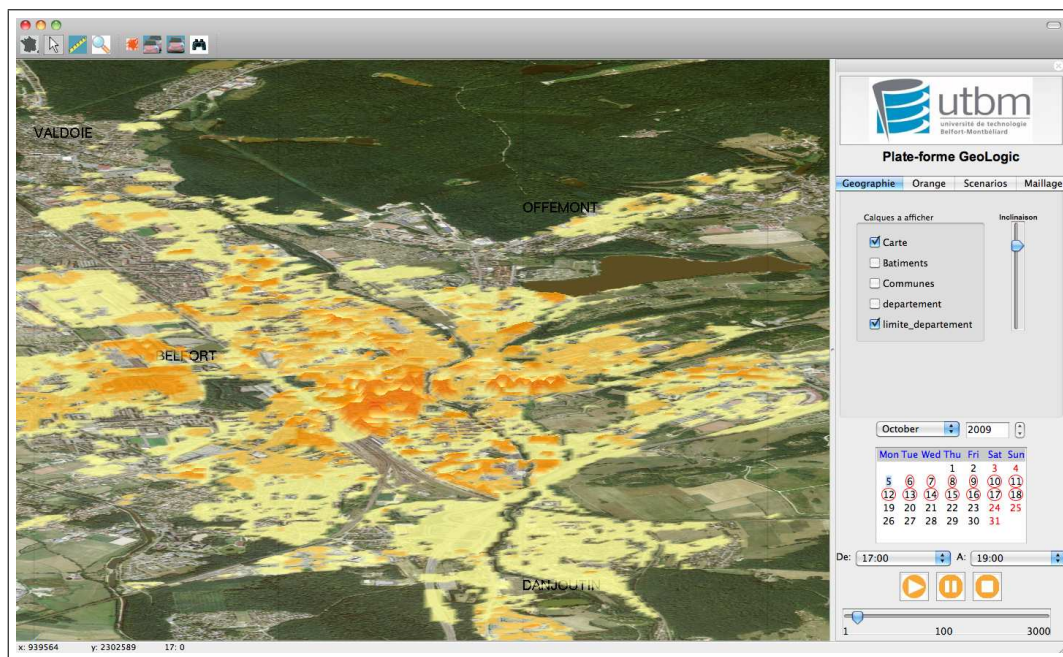


FIGURE 2.13 – Plate-forme geoLogic



FIGURE 2.14 – Informations vectorielles : SIG shapefile

défilement du temps, ou encore de mettre l'application en pause.

2.4.2 Constitution des données de référence

Les données utilisées pour constituer la base de référence sont l'Enquête Ménages/Déplacements (EMD) ainsi que les données INSEE du territoire. Les résultats de l'enquête effectuée en 2004/2005 (enquête redressée) fournissent une base de données contenant notamment les 435 360 enregistrements correspondant aux réponses apportées par chaque personne enquêtée sur les déplacements de la veille, redressées à l'ensemble de la population. Les déplacements contenus dans la base sont ceux ayant pour origine ou destination une zone prédéfinie de la CAB et ayant été effectués en semaine. Ainsi l'EMD nous permet de connaître les déplacements des habitants du territoire étudié tout au long de la journée.

L'EMD est donc riche d'informations sur la manière dont la population se déplace sur le territoire. Les traitements que nous avons effectués sur ces données ont deux objectifs : d'une part fournir des résultats visualisables dans le logiciel développé afin de se rendre compte de la façon dont la population évolue (avec une information variable dans l'espace et dans le temps à différente échelle spatiale i.e. pour une zone particulière ou pour tout le territoire de l'EMD) et d'autre part de constituer la base de référence pour comprendre le lien pouvant exister entre déplacements et activité téléphonique.

2.4.2.1 Agrégation temporelle

L'enquête réalisée contient l'horodatage précis des déplacements de la population sondée. Dans le cadre du projet nous avons échantillonné les journées en périodes de 15 minutes qui correspondent à une précision temporelle exploitable pour l'analyse des déplacements dans le cadre d'une étude sur le transport urbain. Une première étape consiste donc à affecter à chacun des 435 360 enregistrements la plage horaire (le 1/4 d'heure) lui correspondant pour pouvoir agglomérer ces affectations et disposer d'une information globale par période.

2.4.2.2 Agrégation géographique

Une classe de déplacements est un ensemble zone d'origine, zone de destination, plage horaire, motif de déplacement au sens de l'EMD. Le motif est facultatif car les déplacements pourront être distingués ou non selon le motif. Il faut donc calculer pour chaque classe ainsi définie le nombre d'enregistrements correspondant. Ce calcul répond donc à la question : Combien de personnes sont parties de la zone A pour aller à la zone B au cours du 1/4 d'heure t (éventuellement pour le motif m) ?

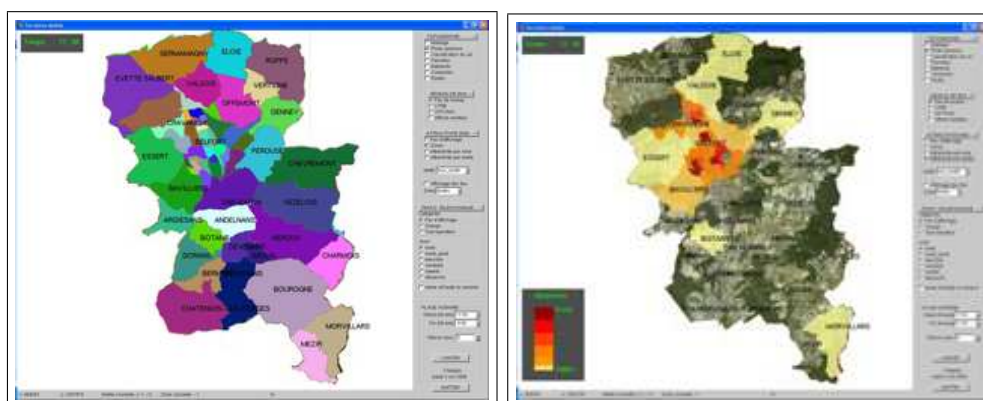


FIGURE 2.15 – Zonage et attractivité des zones à 17h

La Figure 2.15 illustre à gauche le découpage du territoire en différentes zones

(chaque couleur représente une zone). A droite est représentée l'attractivité selon l'EMD de chaque zone. Un niveau de couleur allant du blanc au rouge traduit le niveau d'attractivité de la zone. Ainsi les zones rouges foncées indiquent des lieux à forte attractivité.

2.4.2.3 Calcul de la différence cumulée au cours de la journée

L'évolution de la population dans les zones au cours de la journée est une information utile que l'on peut calculer à partir de l'évolution de l'attractivité de chaque zone au cours du temps. Pour cela, il faut prendre en considération non seulement les flux entrants et sortants à l'instant considéré, mais y ajouter également tous les flux entrants opérés depuis le début de la journée (cumul des arrivées) et retrancher de la même manière la somme des flux sortants depuis le début de la journée (cumul des départs). Le point de référence temporel considéré est à 6h du matin.

Cette information s'exprime de la façon suivante : à un moment donné, est-ce que la zone s'est remplie (teintes rouges) ou vidée (teintes bleues) depuis 6h00 du matin par rapport à la population initiale ? Et quel est le solde en nombre de personnes arrivées ou parties ?

Les teintes claires expriment un faible solde dans un sens ou dans un autre, soit un équilibre des arrivées et des départs. Les teintes foncées présentent toutes un solde positif dans les rouges et négatif dans les bleues.

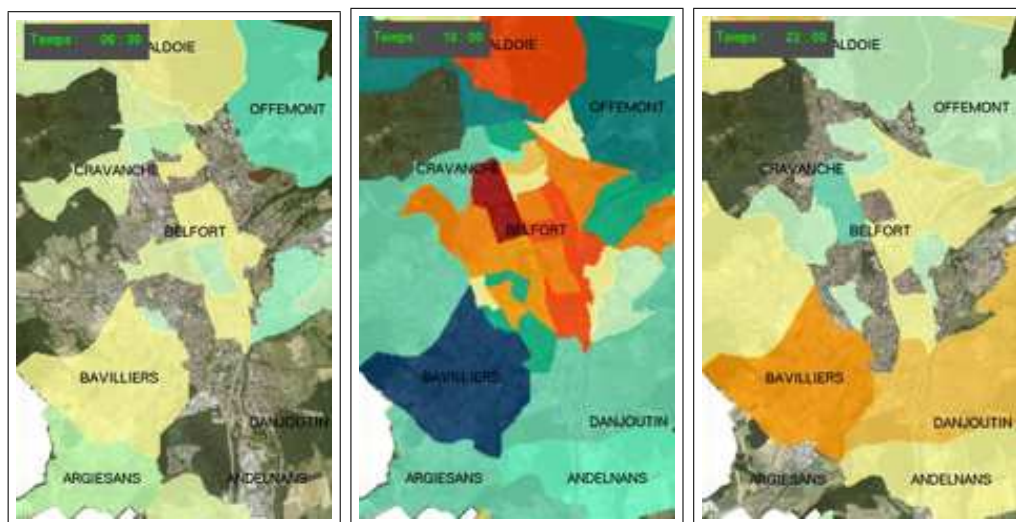


FIGURE 2.16 – Différence cumulée à 6h30, 10h et 22h

La Figure 2.16 présente la différence cumulée à 6h30, 10h et 22h. Alors que le matin à 6h30 il n'y a encore eu que très peu de mouvements, à 10h nous observons

que les zones industrielles et de centre ville se sont particulièrement remplies au détriment des zones périurbaines. A 22h la tendance se rapproche de la situation de 6h30, marquant un retour au domicile.

2.4.2.4 De la différence cumulée à la présence

Dans le cadre de cette section nous nous intéressons à la mise au point d'un procédé pour fournir une information pertinente sur la localisation de la population au cours de la journée.

Afin de faire un parallèle entre les données téléphoniques et l'enquête ménages déplacements, nous devons donc traduire les données de l'EMD en termes de population présente à chaque instant et en chaque endroit du territoire représenté par les zones. Ce calcul qui fournit une information utile en soi sera également la base de comparaison utilisée pour valider la localisation identifiée par les téléphones portables, ces résultats étant construits de manière totalement indépendante.

Pour ce traitement nous considérons la différence cumulée présentée en 2.4.2.3. Pour chaque zone nous ajoutons la population résidente issue de l'INSEE et nous obtenons ainsi une approximation du nombre de personnes présentes chaque 1/4 d'heure de la journée. Il n'y a en effet que les déplacements effectués entre minuit et 6h00 qui ne sont pas pris en considération. L'information produite revient donc à reprendre les résultats de la différence cumulée translatés du nombre d'habitants de chaque zone à tout moment de la journée.



FIGURE 2.17 – Présence à 3 périodes de temps

Les résultats de présence sont représentés Figure 2.17 pour des calculs effectués à 6h, 10h et 22h. La carte de gauche représentant le résultat à 6h00 correspond au recensement INSEE ramené aux zones de l'EMD. A 10h00 nous pouvons observer la nouvelle répartition de la population en fonction des activités de chacun. Enfin la carte de droite montre la répartition à 22h00 qui est semblable à celle observée le matin à 6h00 : la population a presque repris sa distribution initiale.

2.4.3 Expérimentation et résultats

La comparaison des résultats téléphoniques avec l'EMD a pour but d'établir le lien pouvant exister entre l'activité téléphonique, la mobilité des personnes et le territoire. A cette fin nous proposons ici de comparer la répartition spatiale des mobiles d'après notre modèle exposé en section 2.3.1 à l'EMD, ce qui doit nous permettre d'identifier dans quelle mesure les personnes utilisant des téléphones portables sont représentatives de l'ensemble de la population selon le moment de la journée et le lieu. Le principe de validation emploie les zones de l'enquête ménage : pour une zone donnée, l'expérimentation doit montrer comment la présence des personnes localisées par les téléphones portables reflète la présence des personnes issues de l'enquête sur cette même zone. Il faut donc dans un premier temps agglomérer dans les zones les présences détectées via les téléphones sur l'ensemble des mailles du territoire par l'application de notre modèle de répartition spatio-temporelle du trafic téléphonique. Nous pouvons alors comparer les tendances des courbes issues de l'EMD et de la localisation.

Les résultats de localisation doivent être observés quasi obligatoirement dans des contextes de variation spatiale et temporelle. Par exemple, comment évolue au cours d'une journée travaillée la présence de population sur les zones industrielles ? Ou sur les zones commerciales les samedis ? Etc. Il faut donc à la fois regarder le lieu et le moment.

Nous présentons ci-après les comparaisons entre les courbes de présences de l'enquête EMD et les courbes de présences issues du modèle de répartition du trafic téléphonique (les identifiants des zones sont ceux issus de l'enquête). Nous montrons ces comparaisons sur 4 cas de figure représentatifs des phénomènes que nous avons observé sur l'ensemble des 47 zones.

- 2 zones urbaines : zone d'activité industrielle (Techn'hom) et zone du Centre Ville avec de nombreux petits commerces et habitat de type immeuble.
- 2 zone périurbaines : Andelnans et Bavilliers qui sont deux villages avec population résidente (500 habitants à Andelnans et 5400 à Bavilliers) et quelques commerces.

Les courbes sur les 4 graphiques (Figures 2.18 à 2.21) indiquent :

- En rouge, les résultats de l'enquête EMD, c'est-à-dire les données de référence. L'ordonnée tient compte de la population qui habite sur les lieux observés donc le point départ dépend de la zone ; par exemple 0 pour la zone d'activité et 3000 pour la zone centre ville. La zone d'activité fournit une très bonne référence de traçage de la mobilité puisqu'elle n'inclut pas de résidents.
- En bleu, les résultats de notre modèle de distribution spatio-temporelle, c'est-

à-dire les données à valider. Les données sont multipliées par 100 pour que l'amplitude de ces courbes soient dans le même champ de lecture que les courbes de référence et faciliter la comparaison (sauf les données d'Andelnans qui sont multipliées par 10). Le point initial est toujours à 0 c'est-à-dire que l'on n'inclut pas la population résidente dans la mesure.

Une moyenne mobile sur 8 valeurs est calculée sur chaque courbe pour donner les tendances et faciliter l'analyse.

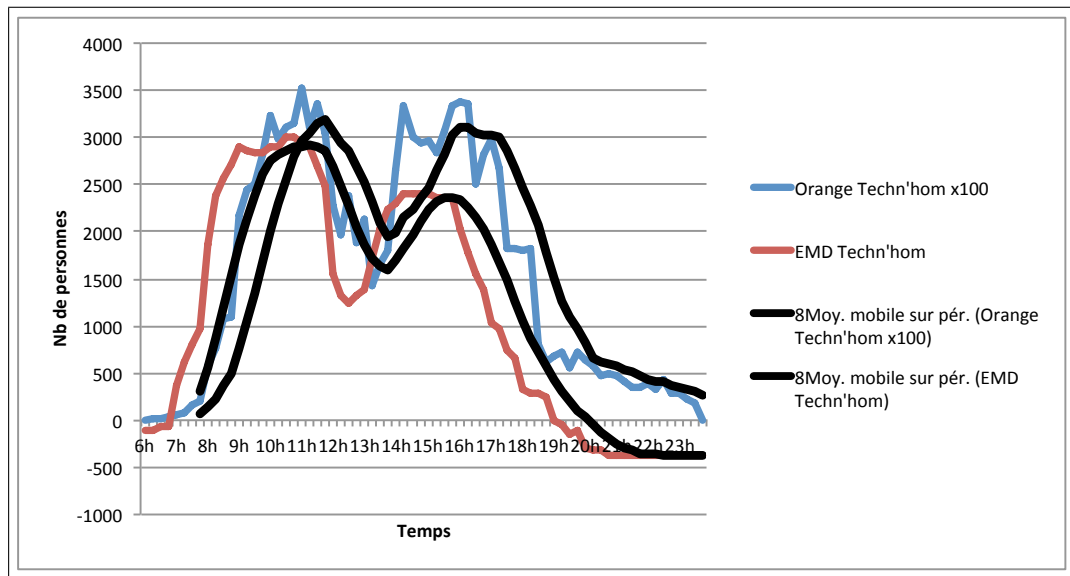


FIGURE 2.18 – Comparaison EMD / modèle en zone d'activité industrielle sans habitat

La zone d'activité sans résident initial montre Figure 2.18 deux courbes dont les tendances sont bien corrélées. Il en est de même pour la zone de centre ville (cf. Figure 2.19). Les autres zones qui sont essentiellement des zones d'habitation ont des courbes qui semblent suivre une corrélation en opposition de phase tel qu'illustré en Figures 2.20 et 2.21. Pour mieux observer l'inversion de phase des zones d'habitation, nous avons retourné les courbes correspondantes. Les Figures 2.22 et 2.23 montrent cette symétrie des données issues du modèle par rapport à l'axe $f(x) = 150$ et $f(x) = 2500$ respectivement.

Des types de zones apparaissent donc selon l'impact du sursol, résidentiel ou pas, et avec activités économiques ou pas. On peut ainsi identifier deux catégories de zones, caractérisées par les points suivants :

- Sur les 2 zones urbaines, les courbes sont très corrélées : la présence des personnes distribuées par le modèle suit étroitement celle des personnes vues par l'EMD. Il s'agit principalement de zones non habitées (Techn'hom) ou

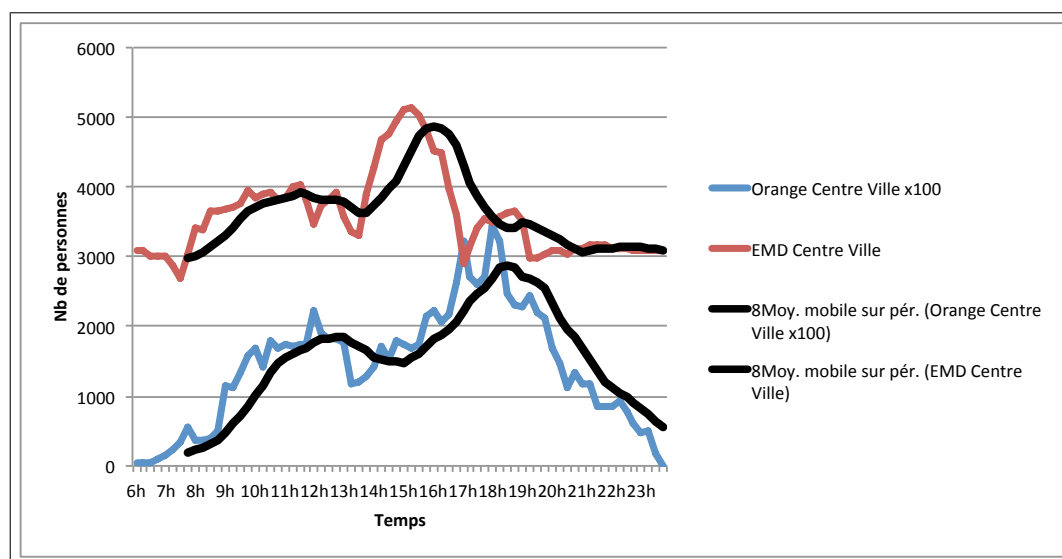


FIGURE 2.19 – Comparaison EMD / modèle en centre ville avec habitats et commerces

faiblement habitées (centre ville) où les personnes se rendent pour travailler ou faire leurs courses du fait d'une typologie d'activités économiques. L'emploi des téléphones mobiles est bien corrélé à la présence des personnes en déplacement sur ces zones.

- Sur les 2 zones péri-urbaines, les courbes sont inversées ou corrélées en inversion de phase. Quand les personnes sont chez elles à des heures de faible mobilité (dans la matinée, à l'heure du déjeuner et le soir après 21h30), elles utilisent peu leur téléphone mobile ; dans les périodes de transitions où des déplacements sont effectués (arrivées pour midi, arrivées en fin de journée entre 17h00 et 21h30), les personnes utilisent leur téléphone mobile. Les deux courbes sont harmonieusement inversées. Là aussi le périmètre horaire propice aux déplacements est détectable par l'activité téléphonique.

Notons que les 4 zones présentent un glissement sur les pics issus de notre modèle par rapport à ceux de l'EMD. Ce glissement est homogène : il se produit vers 12h-12h30 et vers 17h30-19h30 et toujours en décalage d'une heure plus tard avec notre modèle (2h pour la zone centre ville). Ce qui est intéressant dans ce phénomène est qu'il est visible quel que soit le type de zones, on peut donc supposer que le modèle pourrait être calibré. Cependant, il s'agit de périodes à mobilité forte. Côté trafic téléphonique cela signifie un fort usage du téléphone sur ces périodes. Côté EMD cela signifie plus d'incertitude sur les résultats enregistrés pendant l'enquête. En effet, lors de l'enquête, les personnes enquêtées s'expriment sur ce qu'elles ont fait comme déplacement la veille. Ainsi, de façon assez conventionnelle, pour prendre un exemple, une personne aura déclaré qu'elle part du travail à 17h-17h30 et rentre à son domicile. L'enquête traduit cela par le fait de marquer une zone d'origine à

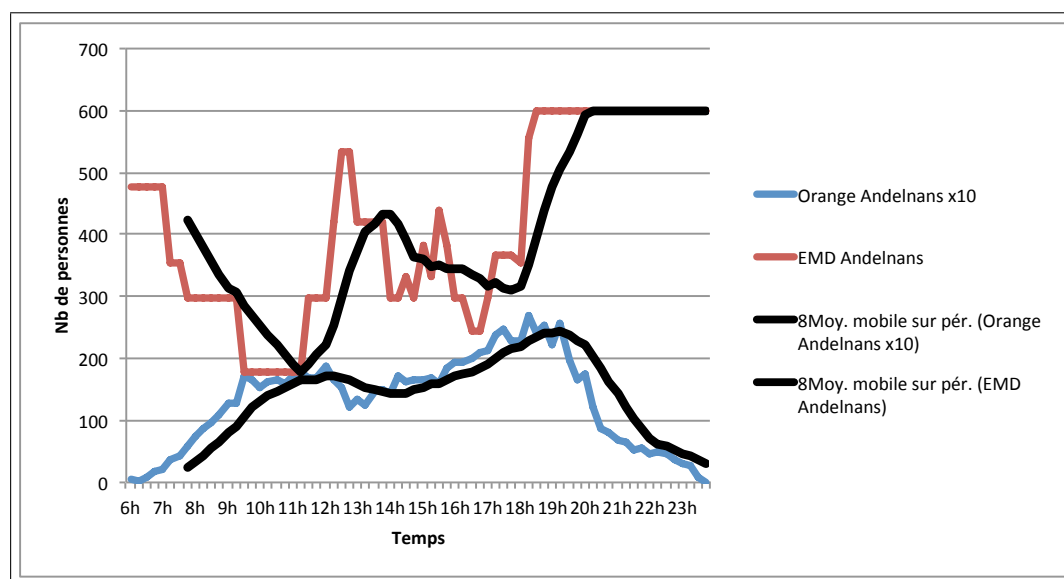


FIGURE 2.20 – Comparaison EMD / modèle en secteur périurbain 1 avec habitat et très peu de commerce

17h (zone de travail) et une zone de destination à 17h30 (zone du domicile). En pratique, l'enquête exclut donc très facilement une période de temps, entre le travail et le domicile, pendant laquelle il y a une forte variation d'un jour à l'autre, la plus forte en réalité. Un jour, la personne va rentrer à cette heure-là ; le lendemain, 1h30 plus tard car elle aura fait des courses ; le surlendemain, 30 min plus tard car elle aura travaillé un peu plus, etc. Ces phénomènes sont visibles lorsque les contraintes de mobilité sont souples, en général quand on sort du travail plutôt que quand on y va, soit plutôt aux heures où nous avons observé un glissement des résultats du modèle par rapport à l'EMD. On ne peut donc pas en conclure définitivement qu'il faut calibrer le modèle mais plutôt que c'est la partie la plus stochastique du comportement humain et qu'une enquête sur une journée, comme c'est fait pour l'EMD, ne traduit pas forcément correctement cette dimension aléatoire des comportements. Ces hypothèses demanderaient à être vérifiées sur d'autres villes, pour ce qui est des données issues de l'EMD, et sur plusieurs semaines pour ce qui est de données issues de la téléphonie.

Par ailleurs, nous voyons qu'une simple lecture naïve de l'activité téléphonique ne permet pas en général de conclure sur le nombre de personnes présentes dans une zone. Et pour cause, les téléphones ne sont pas utilisés de façon linéaire tout au long de la journée. Il est vraisemblable que nous appelons plus facilement à 12h qu'à 23h. De même le lieu semble avoir un effet direct sur l'utilisation du téléphone portable. En zone très résidentielle, on favorisera sans doute le téléphone fixe étant à la maison, alors que dans les commerces ou au travail, le portable sera privilégié.

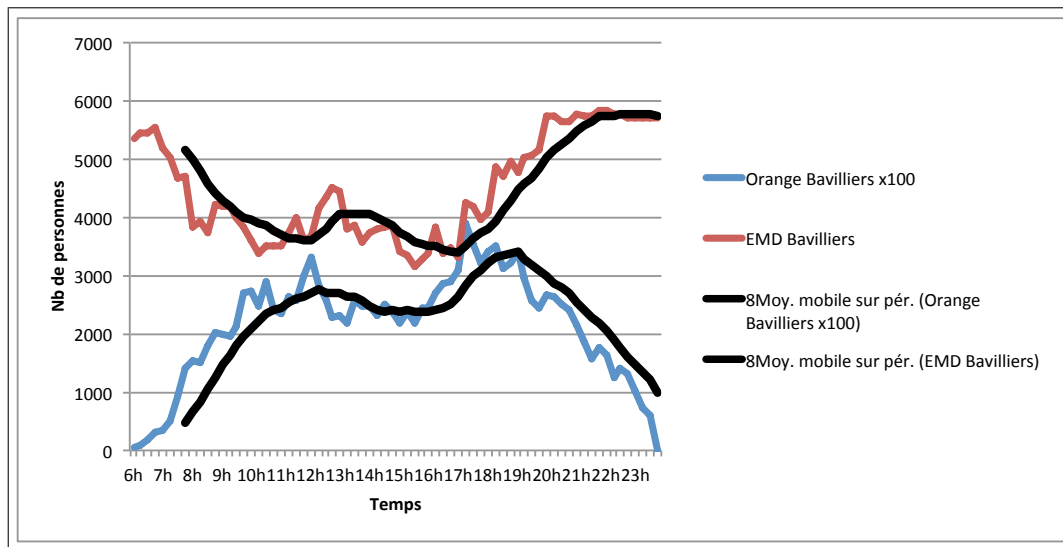


FIGURE 2.21 – Comparaison EMD / modèle en secteur périurbain 2 avec habitat et très peu de commerce

Il apparaît alors nécessaire de ne pas se contenter de traduire l'activité téléphonique directement en niveau de présence mais de déterminer la représentativité de l'activité téléphonique dans un contexte spatio-temporel.

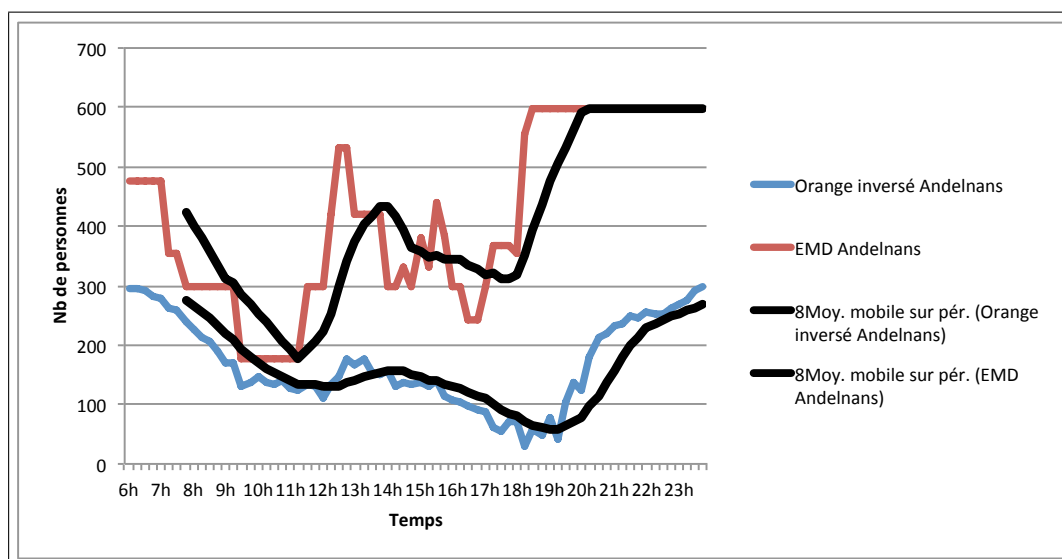


FIGURE 2.22 – Symétrie du modèle en secteur périurbain 1 avec habitat et très peu de commerce

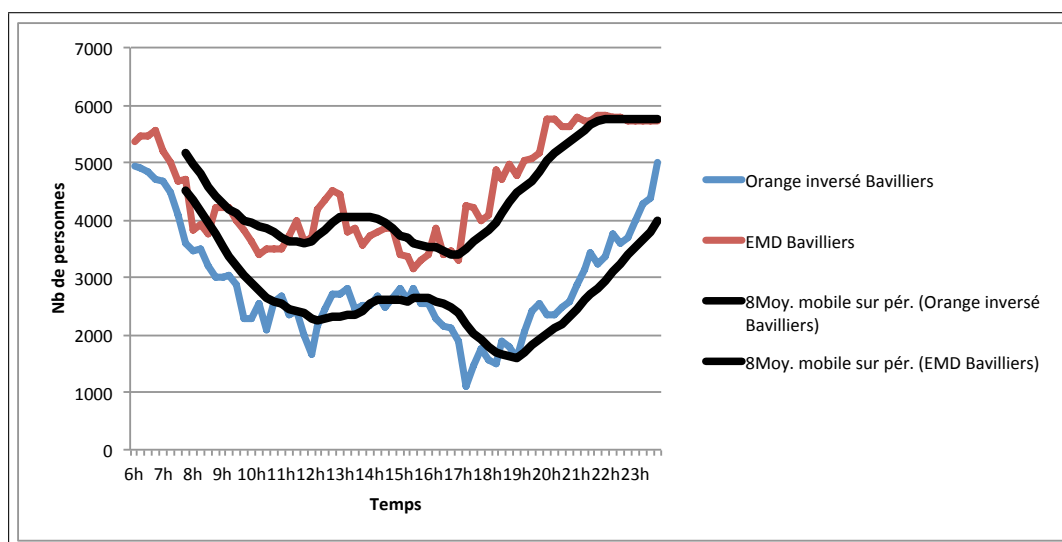


FIGURE 2.23 – Symétrie du modèle en secteur périurbain 2 avec habitat et très peu de commerce

2.4.4 Construction d'une *Signature Territoriale*

Les résultats précédents ont illustré le fait que la façon de téléphoner dépend directement du type de lieu (l'espace) et du moment (la journée, l'heure). Il apparaît donc fondamental d'étudier le niveau de stabilité des communications, non pas au cours de la journée mais d'une semaine à l'autre. En effet, si l'utilisation du téléphone varie de façon significative pour un même type de jour, dans la même zone et au même moment de la journée, aucune conclusion sur la représentativité des communications ne peut être tirée. A l'échelle de l'ensemble des personnes, l'activité téléphonique serait chaotique et ne laisserait aucun espoir d'en tirer une information sur la répartition de la population. Au contraire, si pour une même zone, le même jour et au même moment, nous pouvons observer un comportement stable d'une semaine à l'autre, cela montre que la façon de téléphoner est directement et étroitement liée au lieu. A ce titre les communications constitueraient un bon marqueur de la présence dans une zone.

Pour vérifier ceci, nous introduisons la notion de *Signature Territoriale*. Nous la définissons comme étant le coefficient multiplicateur qui, pour chaque jour type et à chaque moment, permet de passer du nombre d'événements téléphoniques au nombre de présents issus de l'EMD. Ce coefficient devrait certainement être important le soir par exemple, étant donné qu'un appel à 22h représente certainement beaucoup plus de monde présent qu'un appel à 12h. Il s'obtient naturellement par la division euclidienne du nombre de personnes comptées par l'EMD par le nombre d'événements téléphoniques.

Définition 5 (Signature Territoriale) *On appelle Signature Territoriale d'une zone la matrice caractéristique permettant d'identifier pour chaque type de jour et chaque horaire le nombre de personnes que représente un événement téléphonique.*

Afin d'étudier la stabilité et donc la pertinence de cette signature, reprenons les zones utilisées précédemment. Nous allons nous focaliser sur les journées connues pour être à la fois semblables et de plus forte mobilité, le mardi et le jeudi [Jeannic 2011]. Par ailleurs, pour renforcer cette validation, nous comparons ces journées au cours de deux semaines différentes. La figure 2.24 illustre les résultats obtenus pour les zones que nous avons identifiées comme étant en phase (centre ville et zone d'activité industrielle). Les courbes hautes concernent le centre ville et les courbes basses la zone d'activité industrielle.

En abscisse est représentée l'heure de la journée, et en ordonnée le coefficient permettant de passer du nombre de personnes issu du modèle proposé au nombre de personnes issu de l'EMD. Concernant les courbes de la zone d'activité industrielle (courbes basses), on constate qu'elles sont à 0 passé 19h. Etant donné qu'il s'agit d'une zone sans habitat, cela signifie que, d'après l'EMD, il y a plus de personnes

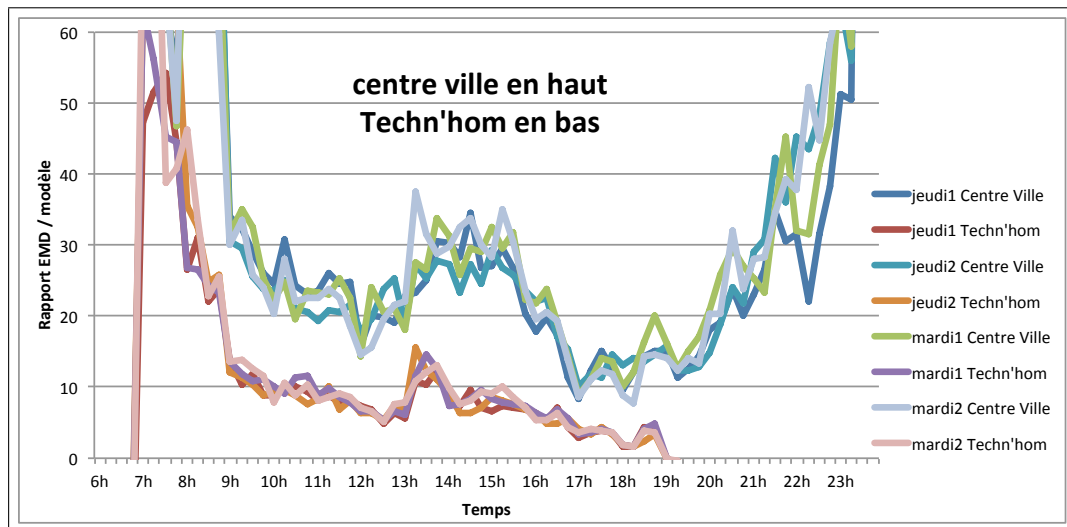


FIGURE 2.24 – Signature Territoriale pour le centre ville (courbes hautes) et la zone d'activité industrielle (courbes basses) - pour deux mardis et jeudis différents

qui sont parties depuis le début de la journée (6h du matin) que de personnes qui sont arrivées. L'EMD indique donc un nombre négatif de personnes passé 19h. Cela montre les limites de ces enquêtes, qui sont la meilleure référence que l'on ait aujourd'hui mais qui comportent de nombreuses incertitudes quant à leurs résultats.

Les 4 courbes de chaque zone se révèlent être très fortement semblables. De plus, d'une zone à l'autre les courbes sont bien distinctes. La signature se révèle être un identifiant propre à chaque zone et stable d'un jour à l'autre.

Nous avons de la même façon construit les signatures correspondant aux zones périurbaines. Le résultat observable sur la Figure 2.25 révèle une tendance très semblable à ce que nous venons de commenter. La signature permet d'identifier clairement chaque zone, avec en haut les courbes de la signature de Bavillers, et en bas celles d'Andelnans.

Nous identifions donc très nettement une tendance stable propre à chaque zone pour les 4 jours observés, quel que soit le type de zone. A ce titre nous montrons ici que la façon dont nous avons utilisé les événements du réseau cellulaire conduit à un marqueur fiable de la présence des personnes. De plus, l'observation généralement faite dans les études de transport que le mardi et le jeudi ont des profils de mobilité très semblables est ici également vérifiée.

Enfin, la lecture détaillée des graphiques fournit des indicateurs sur la représentativité des communications mobiles. Un pic bas à 12h précise au centre ville (Figure 2.24) montre que, proportionnellement au nombre de personnes présentes, on appelle beaucoup plus. Le ratio descend alors à 15 ce qui signifie qu'un événement généré

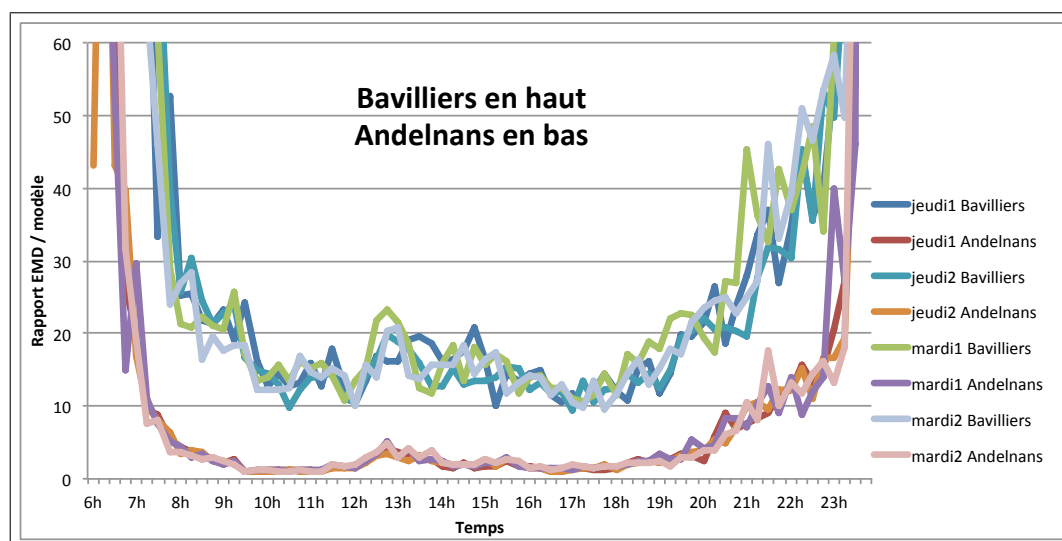


FIGURE 2.25 – Signature Territoriale pour les zones périurbaines - pour deux mardis et jeudis différents

au cours du 1/4 d'heure correspond à 15 personnes présentes. Au contraire, un pic haut vers 15h avec un ratio de 35 signifie qu'un événement téléphonique représente 35 personnes présentes. On constate des pics bas également à 17h et 18h mais pas entre les deux. Sur la zone d'activité le pic de 12h est moins marqué : on téléphone de manière plus étalée dans le temps avec un léger pic à 12h30. On peut également constater un pic haut entre 13h et 14h sur la zone d'activité. Il s'agit d'un moment où le nombre de communications par personne diminue, c'est l'heure du déjeuner. L'ensemble de ces observations demanderait bien entendu à être confronté à des relevés sur plusieurs semaines mais nous ne disposons pas d'autres relevés actuellement.

La comparaison des deux profils de zone permet d'identifier deux tendances distinctes. Alors que les zones de centre ville et d'activité industrielle ont une évolution forte entre 12h et 14h (pic bas suivi de pic haut), ce phénomène est nettement moindre dans les zones périurbaines. De plus, alors que dans les zones périurbaines le niveau de la signature est très semblable le matin et l'après-midi, dans les zones d'activité ce niveau est différencié, nettement plus élevé le matin qu'en fin d'après-midi. Cela montre que dans les zones d'activité l'utilisation du mobile est très dépendante de l'heure et du couplage entre l'heure et l'activité territoriale. Dans les zones d'habitation, nous observons une tendance beaucoup plus linéaire en journée. Entre 9h et 20h, l'utilisation du téléphone se fait de façon assez uniforme, quelle que soit l'heure.

L'étude approfondie de ces *signatures territoriales* se révèle être riche de sens, et fournit de nombreuses informations sur le fonctionnement du territoire.

2.4.5 Généralisation de la *Signature Territoriale*

Nous avons montré que la façon d'utiliser le téléphone mobile est globalement très stable d'une semaine à l'autre pour une même zone, au même jour et au même moment. Par contre, elle peut différer significativement d'une zone à l'autre. Nous avons ainsi pu quantifier pour chaque zone le nombre de personnes représentées par l'activité téléphonique. La question qui se pose est de savoir dans quelle mesure cette signature peut être déduite de la constitution de la zone, en termes de données de sursol et socio-économiques. Est-il possible d'identifier des caractéristiques territoriales qui permettent d'estimer la signature correspondante ?

Malgré nos efforts dans ce sens, il apparaît que cette signature territoriale suit des règles particulièrement subtiles que nous n'avons pas encore clairement identifiées. Nous avons eu recours à une approche à base de régression linéaire et d'un estimateur des moindres carrés pour tenter d'identifier la manière dont contribue chaque classe de terrain à la pondération globale des zones. Or nous avons constaté que la façon de communiquer via les mobiles n'est pas toujours directement corrélée à la constitution territoriale des zones (habitat, voirie, commerces, etc.) mais qu'elle intègre d'autres facteurs, peut-être sociaux. Il apparaît ainsi que l'usage des mobiles dans les habitations diffère d'un quartier à un autre. De la même façon les téléphones sont utilisés différemment dans les entreprises selon les zones de localisation des entreprises concernées ; là encore il faudrait regarder d'autres facteurs tels que la nature de l'activité de l'entreprise.

Une orientation qui nous semble prometteuse est le recours aux techniques de classification (souvent appelées clustering) telles que l'analyse en composante principale (ACP) afin de regrouper les zones ayant des profils semblables. Nous laissons cette partie en perspective pour des travaux futurs.

2.5 Synthèse

Le modèle proposé dans ce chapitre s'inscrit dans un triptyque téléphonie - mobilité - territoire. Les travaux présentés sont novateurs à plusieurs titres, tout d'abord, de part les données utilisées. Contrairement à ce que font les outils classiques de modélisation de la mobilité à partir d'enquêtes et d'estimations, notre approche vise à fournir des résultats qui sont des mesures de déplacements issus du terrain. Ensuite, ces travaux sont novateurs à travers l'unification de données de déplacement et de données du territoire. Contrairement à ce qui commence à être fait à travers l'utilisation des données issues de la téléphonie mobile, le modèle que nous proposons s'appuie sur la localisation des pôles attracteurs/générateurs de mobilité constitués des habitations, des commerces, des entreprises, des écoles etc., tout ce qui constitue une activité humaine, et à ce titre est un modèle de la catégorie activité-centrée.

Dans une première partie a été présenté le fonctionnement général d'un réseau de téléphonie mobile. L'objet de cette partie était de fournir un éclairage sur les données issues de cette technologie qui peuvent être utiles à une modélisation de la mobilité. Deux types de données ont été identifiés, les données de couverture radio et les données de trafic. Alors que les données de couverture radio s'attachent à définir la portion de territoire couverte par chaque antenne, avec une probabilité de prise de communication pour chaque maille (PPC), les données de trafic rendent compte du nombre d'événements enregistrés par chaque station, événements liés à des communications. Les données de couverture sont des données statiques, fixées par la topographie et le paramétrage des stations. Les données de trafic sont quant à elles dynamiques, variant au cours de la journée et d'un jour à l'autre.

La deuxième partie de ce chapitre contient la description du modèle proposé pour une répartition spatiale et temporelle des personnes. Ce modèle qui vise à affiner les connaissances actuelles sur la mobilité regroupe des données téléphoniques et des données de description du territoire. Il s'appuie sur la notion de classe d'attraction qui, à chaque classe de sursol, attribut un vecteur dynamique de pondérations dépendant du type de jour et du moment de la journée. La mise en correspondance de la cartographie des PPC, de la cartographie des poids d'attraction spatio-temporelle, et du trafic téléphonique dans notre modèle a donné lieu à une distribution spatio-temporelle des personnes. La mise en oeuvre du modèle proposé dans la plateforme geoLogic a donné lieu au développement d'un logiciel d'analyse visuelle de la mobilité. Il offre la possibilité d'étudier des moments et des lieux précis, en continu ou en arrêt sur image sur la carte. Il est ainsi possible d'en tirer des informations utiles pour l'aménagement d'un service. Cela permet par exemple de fixer les fréquences de passage des bus, ou encore de déterminer les meilleurs emplacements pour des arrêts des bus. Sur le territoire considéré, l'étude est basée sur la présence de personnes détectées par leurs communications téléphoniques au cours de la journée.

Enfin, afin d'expérimenter notre modèle, une étude comparative est proposée. Ainsi la distribution des personnes issue de notre modèle est mise en regard de données de localisation construites à partir d'une enquête ménages/déplacements, ainsi que de données statistiques issues du recensement de l'INSEE. Deux types de zones territoriales ont ainsi été identifiés, avec deux types d'utilisation des téléphones portables bien distincts. D'une part il y a les zones de forte attractivité, tel que le centre ville, ou les zones d'activité. Il s'agit des zones particulièrement intéressantes à étudier car elles engendrent une forte mobilité. D'autre part il y a les zones de type résidentiel. On y retrouve les quartiers d'habitation ou les communes périphériques. Ces zones faiblement attractives sont souvent plus difficiles à étudier. Elles suivent un schéma de présence et de mobilité plus variable. L'approche proposée a permis de montrer que dans les zones de forte attractivité la présence de personnes à partir d'enquêtes est étroitement corrélée avec l'activité téléphonique. Dans ce cas, l'activité téléphonique est un bon marqueur direct de la population présente. Dans les zones d'habitation par contre, l'activité téléphonique s'est révélée être en opposition

de phase. Paradoxalement, plus il y a de personnes dans ces zones, moins il y a d'activité téléphonique via les mobiles. Le phénomène mis en évidence est qu'une forte fréquentation de ces zones signifie qu'il s'agit de résidents, qui utilisent de façon préférée le téléphone fixe, et ne génère ainsi que peu d'activité téléphonique sur le réseau mobile. Une étude complémentaire s'avère alors nécessaire pour permettre d'effectuer la transcription de l'activité téléphonique en nombre de personnes présentes. Pour étudier de façon plus fine le lien existant entre l'usage du téléphone et la présence de personnes, le concept de signature territoriale a été proposé. Il a permis de mettre en évidence ce lien, de façon différenciée pour chaque zone étudiée. Cette signature a également permis de vérifier le très bon niveau de stabilité de l'indicateur téléphonique d'une semaine à l'autre, ce qui renforce l'intérêt du modèle que nous avons proposé.

Une modélisation précise des déplacements est un élément essentiel pour guider une décision en matière d'aménagement territorial. Mais souvent, savoir comment une population se déplace n'est pas suffisant. En s'appuyant sur des données de mobilité, la seconde partie de ce mémoire est consacrée à la définition et l'utilisation d'algorithmes permettant d'optimiser les aménagements au sein d'un territoire. Le cas d'étude réel présenté concerne le déploiement d'un service d'auto-partage pour une collectivité. Mais pour introduire cette seconde partie, le chapitre suivant s'attache tout d'abord à présenter les éléments permettant de comprendre les métaheuristiques combinatoires en contexte multiobjectif.

Deuxième partie

Optimisation combinatoire
multiobjectif de l'auto-partage

Optimisation combinatoire multiobjectif : état de l'art

Dans ce chapitre est introduit le problème de l'optimisation combinatoire multiobjectif. Dans un premier temps les principaux concepts liés au sujet sont présentés, ainsi que les principales définitions. En particulier, le problème de non-comparabilité des solutions multiobjectif constitue un fil conducteur de ce chapitre. Dans une deuxième partie un inventaire des heuristiques couramment utilisées est proposé. Elles sont présentées regroupées en trois grandes catégories : les approches par transformation en problème mono-objectif, les approches non-Pareto et enfin les approches Pareto. Une présentation est ensuite faite des heuristiques hybrides, qui combinent des heuristiques provenant des différentes catégories précédentes. Enfin, pour appréhender l'évaluation d'un ensemble de solutions Pareto optimales, un inventaire des principaux indicateurs d'évaluation est proposé.

Sommaire

3.1	Introduction	80
3.2	Notions de base et définitions	81
3.2.1	Optimisation mono-objectif	81
3.2.2	Formulation d'un problème d'optimisation multiobjectif	82
3.2.3	Espace de décision vs espace objectif	82
3.2.4	Optimisation combinatoire et métaheuristiques	83
3.2.5	Dominance et Pareto optimalité (le problème de la relation d'ordre)	85
3.3	Métaheuristiques pour les problèmes multiobjectifs	87
3.3.1	Classification des principales approches	87
3.3.2	Les approches par transformation mono-objectif	87
3.3.3	Les approches non-Pareto	92
3.3.4	Les approches Pareto	93
3.3.5	Les approches hybrides	101
3.4	Evaluation des algorithmes multiobjectifs : une approche à base d'indicateurs	103
3.4.1	Indicateur de contribution	104
3.4.2	Indicateur ε -additif	105
3.4.3	Indicateur hypervolume	105
3.4.4	Indicateur de distance générationnelle	108
3.5	Synthèse	108

3.1 Introduction

De nombreux domaines tant scientifiques, économiques qu'industriels ou encore issus de la vie de tous les jours nécessitent de prendre les bonnes, voire les meilleures décisions. Alors que nous effectuons quotidiennement des choix à partir de simples raisonnements, tel que choisir un itinéraire en voiture, de nombreux problèmes, industriels notamment, sont bien trop complexes pour être abordés sans outils efficaces. L'optimisation intervient comme branche issue des mathématiques pour résoudre de tels problèmes. La quantité de travaux publiés ces dernières années atteste d'une discipline en pleine évolution, dans laquelle de nombreux axes restent encore à explorer.

Alors que dans certaines situations il est possible d'identifier un objectif unique à optimiser, un grand nombre de problèmes réels imposent de considérer plusieurs objectifs à la fois, souvent contradictoires entre eux. L'optimisation multiobjectif vise ainsi à optimiser ces problèmes complexes. Une des grandes difficultés vient du fait que la meilleure solution pour un objectif est rarement la meilleure sur l'ensemble des objectifs. De ce fait, contrairement à l'optimisation mono-objectif qui admet une solution optimale unique, la solution d'un problème d'optimisation multiobjectif (Multiobjective Optimization Problem - MOP) est en réalité un ensemble de solutions de compromis connu sous le nom de Pareto optimal. Toute solution de cet ensemble est telle qu'il n'existe aucune solution réalisable améliorant simultanément tous les critères. On parle alors d'une solution non-dominée.

L'ensemble Pareto optimal peut dans certains cas être de très grande taille. De plus pour la plupart des problèmes réels un tel ensemble n'est pas connu. Ainsi l'un des enjeux principaux en optimisation combinatoire est de fournir un algorithme qui permette d'approcher au mieux cet ensemble Pareto optimal.

Ce chapitre présente les différents concepts nécessaires à la bonne compréhension de ce domaine. Ils fournissent en particulier les éléments permettant d'appréhender le chapitre 4. En plus des concepts principaux, nous faisons ici un historique des principales métaheuristiques permettant de résoudre un MOP. Ces méthodes sont organisées autour de trois grandes familles : les approches par transformation en mono-objectif, les algorithmes non-Pareto, et enfin les approches de type Pareto. Dans cette dernière catégorie nous présenterons, en plus des approches classiques, les techniques visant à hybrider des algorithmes issus de différentes familles, et les approches à base d'indicateurs.

Une partie des définitions, concepts et algorithmes présentés dans ce chapitre est notamment inspirée des livres [Talbi 2009] et [Coello 2007]. La thèse d'Arnaud Liefooghe est également une référence en la matière.

La première section de ce chapitre nous permet de revenir sur les principales

définitions et notions de base rencontrées dans le domaine de la résolution d'un MOP. Dans la deuxième section nous présentons les principales métaheuristiques rencontrées dans ce domaine. La section 3.4 présente les approches permettant d'évaluer le niveau de performance d'un algorithme. L'accent est mis sur les indicateurs rencontrés pour refléter le niveau de qualité d'un ensemble de solutions. Enfin la dernière section traite de l'intervention du décideur face à un problème multiobjectif.

3.2 Notions de base et définitions

3.2.1 Optimisation mono-objectif

La recherche opérationnelle s'est dans un premier temps adressée aux problèmes d'optimisation mono-objectif (Optimization Problem - OP). Aujourd'hui encore de nombreux travaux de recherche concerne ce type de problèmes, notamment en proposant de nouvelles théories évolutionnaires ou encore des approches coopératives. De plus, dans le cas de l'optimisation multiobjectif certaines approches visent à réduire le nombre d'objectifs pour appliquer les algorithmes, souvent plus simples, disponibles pour résoudre les problèmes d'optimisation mono-objectif.

Définition 6 (Problème d'optimisation mono-objectif) *Dans le cas général, un problème d'optimisation mono-objectif consiste à maximiser (ou minimiser) une fonction $f(x)$:*

$$OP \begin{cases} \max & z = f(x) \\ \text{tel que :} & x \in X \end{cases} \quad (3.1)$$

x est ici le vecteur de n variables de décision $x = (x_1, x_2, \dots, x_n)^t$, avec t l'opérateur de transposition. X désigne l'espace de recherche regroupant toutes les solutions réalisables.

Définition 7 (Optimum global) *Etant donné un problème d'optimisation, la solution x^* est dite maximum global ssi :*

$$\forall x \in X : f(x^*) \geq f(x) \quad (3.2)$$

Souvent, lors de la résolution de problèmes réels, il n'est pas possible de garantir qu'une solution est l'optimum global sur l'ensemble de l'espace de recherche. On explore alors un sous-espace $X_1 \subset X$ de l'espace de recherche pour trouver l'optimum global x_1^* de X_1 , vérifiant $f(x^*) \geq f(x_1^*)$ toujours pour un problème de maximisation.

Nous le voyons ici, le contexte mono-objectif permet toujours d'identifier la solution optimale, au moins sur X_1 . Lorsque le problème d'optimisation devient multiobjectif, nous verrons qu'il en est tout autre chose.

3.2.2 Formulation d'un problème d'optimisation multiobjectif

La formulation générale d'un problème d'optimisation multiobjectif (Multiobjective Optimization Problem - MOP) est de la forme :

$$MOP \begin{cases} \max & z = f(x) \\ \text{où} & f(x) = (f_1(x), f_2(x), \dots, f_n(x)) \\ \text{tel que :} & x \in X \end{cases} \quad (3.3)$$

avec $n \geq 2$ le nombre de fonctions objectif, X l'ensemble des solutions réalisables dans l'espace de décision, $x = (x_1, \dots, x_k) \in X$ est le vecteur regroupant les variables de décision. Dans le cadre de problèmes combinatoires X est un ensemble discret. De plus soit $Z \subseteq \mathbb{R}^n$ tel que $Z = f(X)$, Z est l'ensemble des solutions réalisables dans l'espace objectif.

L'unique différence avec les problèmes d'optimisation mono-objectif réside dans le fait que $f(x)$ ne prend plus ses valeurs dans $\mathbb{R}$ mais dans $\mathbb{R}^n$.

3.2.3 Espace de décision vs espace objectif

L'espace de décision regroupant les variables $x = (x_1, \dots, x_k) \in X$ est de dimension k . Il s'agit de l'espace dans lequel sont représentées les solutions.

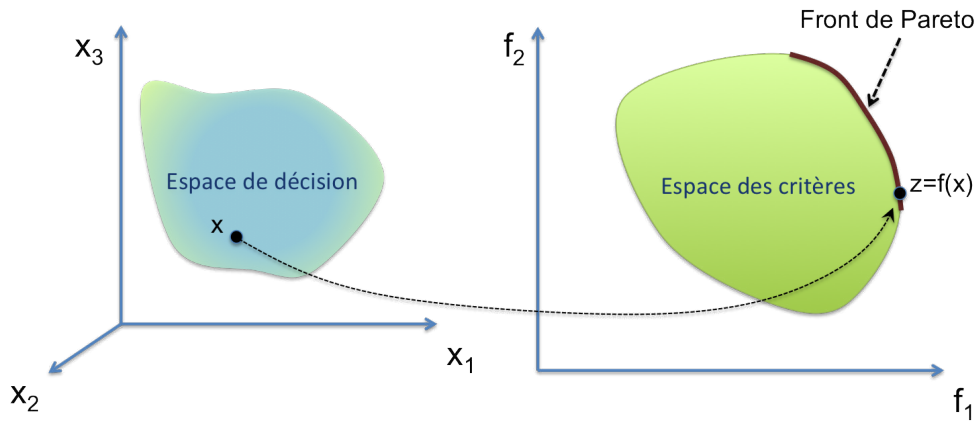


FIGURE 3.1 – Espace de décision / espace objectif

L'espace objectif, également appelé espace des critères, de dimension n permet de projeter l'évaluation des solutions. La Figure 3.1 illustre ces deux espaces, avec à gauche l'espace de décision de taille $k = 3$ dans lequel est représentée la solution x . À droite se trouve l'évaluation correspondante pour $n = 2$ fonctions à optimiser.

Sur cette figure représentant un problème de maximisation $z = f(x)$ appartient au front de Pareto : x est une solution Pareto optimale.

Dans la plupart des situations en optimisation le choix d'une solution se fait sur la base de l'étude de l'espace objectif. Il apparaît cependant parfois intéressant d'observer également les solutions dans l'espace de décision.

3.2.4 Optimisation combinatoire et métaheuristiques

Le cadre dans lequel s'inscrit les travaux de cette thèse concerne l'optimisation combinatoire. La fonction à optimiser, souvent appelée **fonction coût** ou **fonction objectif**, est définie sur un ensemble discret. Il s'agit donc de trouver l'optimum (maximum ou minimum) parmi un ensemble fini. Bien que dans la théorie il suffirait d'énumérer toutes les solutions possibles et de retenir la meilleure, il n'est bien souvent matériellement pas possible d'effectuer cet inventaire exhaustif dans un temps "raisonnable". Le temps de calcul constitue alors une contrainte forte à prendre en compte. Selon le problème considéré, une durée acceptable pour que le programme propose une solution sera de quelques millisecondes - dans une situation d'urgence à bord d'un véhicule par exemple - à quelques semaines lorsqu'il s'agit de planifier une nouvelle ligne de production dans un site industriel par exemple.

3.2.4.1 Complexité d'un problème et NP-complétude

Le but de cette section n'est pas de détailler la notion de complexité, qui est largement traitée par ailleurs. Pour une introduction plus complète à ce domaine le lecteur est renvoyé au célèbre ouvrage de Garey et Johnson [Garey 1979]. L'objet de cette partie est de rappeler quelques fondamentaux de la complexité qui sont liés aux problèmes de décision.

La résolution d'un problème passe par la mise au point d'un algorithme. Bien que la notion d'algorithme remonte à l'antiquité dans la démarche, elle n'a été formalisée réellement qu'avec l'apparition de la machine de Turing. Nous ne reviendrons pas ici sur ces fondements de l'informatique, qui sont remarquablement remis en perspective dans l'ouvrage de Gilles Dowek *Les métamorphoses du calcul* [Dowek 2007]. Un algorithme est ainsi généralement considéré comme efficace si sa complexité est bornée par une fonction polynomiale en la taille de l'instance à résoudre.

Définition 8 (Problème de classe P) *La classe P regroupe l'ensemble des problèmes de décision dont la solution est atteignable en un temps polynomial par une machine de Turing **déterministe**.*

Il existe toutefois un grand nombre de problèmes de décision pour lesquels il n'y a pas d'algorithme polynomial connu.

Définition 9 (Problème de classe NP) *La classe NP regroupe l'ensemble des problèmes de décision pouvant être résolus en un temps polynomial par une machine de Turing **non déterministe**.*

Remarque : Un problème appartenant à la classe P ne signifie pas nécessairement qu'il sera résoluble par un programme déterministe en temps "raisonnable". Inversement, certains problèmes de la classe NP admettront des instances pouvant être résolues par un algorithme déterministe.

On parle également souvent de problèmes NP-complets. Il s'agit des problèmes les plus difficiles de la classe NP, qui sont au moins aussi difficiles que tous les problèmes NP. De fait tout problème de la classe NP doit pouvoir se ramener à un problème NP-complet par réduction polynomiale.

3.2.4.2 Résolution exacte : approche déterministe

Lorsque les problèmes le permettent, le choix de l'optimum peut se faire suite à une exploration complète de l'espace de recherche. Ces approches, souvent basées sur le parcours d'arbres, ont l'avantage de fournir de façon certaine la solution optimale du problème. Sans faire l'exploration complète, d'autres algorithmes s'appuient sur les propriétés du problème pour garantir la solution optimale.

Ainsi, l'algorithme du simplexe par exemple permet de résoudre un problème souvent de façon efficace, en parcourant la fermeture convexe de l'ensemble de recherche. Cet algorithme s'applique toutefois aux problèmes garantissant cette propriété de convexité : problèmes en variable continue, ou problèmes en variables entières avec une matrice des contraintes unimodulaire.

De même, l'algorithme de Branch and Bound repose sur une énumération implicite où le problème est découpé en sous-problèmes (séparation) évalués à l'aide d'une relaxation, généralement continue ou lagrangienne.

On citera également les approches de type programmation dynamique qui permettent de résoudre un problème de la même famille puis de trouver une relation de récurrence liant les solutions de ces deux problèmes, ou encore les méthodes polyédrales dans lesquelles des contraintes sont ajoutées pour rendre convexe le domaine de solutions admissibles.

3.2.4.3 Résolution approchée : (méta)heuristiques

Lorsque le recours à un algorithme exact n'est pas possible pour résoudre un problème en raison d'un temps de calcul trop élevé, le recours à un algorithme approché permet toutefois de proposer une solution. Ces algorithmes approchés sont également appelés **heuristiques**. Le principe est ici de proposer une solution approchée de la meilleure qualité possible dans un temps limité, en exploitant notamment la structure du problème considéré. Les heuristiques présentent l'avantage de pouvoir résoudre un problème *NP* en un temps polynomial. Parmi les heuristiques les plus simples figurent les algorithmes gloutons. Le principe est de construire la solution itérativement en effectuant des choix optimaux localement. Cette démarche générique est à adapter selon la structure du problème posé. C'est là la caractérisation d'une heuristique : elle est définie par les spécificités du problème.

Bien que les heuristiques soient souvent des approches stochastiques, il est important de noter qu'elles peuvent être déterministes.

Les **métaheuristiques** se distinguent des heuristiques par leur niveau d'abstraction. Alors que les heuristiques sont destinées à résoudre un type de problème spécifique, une métaheuristique peut être adaptée à une large variété de problèmes. Ainsi la recherche locale ou encore les algorithmes génétiques sont considérés comme des métaheuristiques.

3.2.5 Dominance et Pareto optimalité (le problème de la relation d'ordre)

Comme cela a été présenté en introduction, il n'existe pas une solution unique à un problème multiobjectif mais un ensemble de solutions de meilleur compromis. L'absence d'ordre total entre les solutions d'un problème multiobjectif a conduit à élaborer une relation d'ordre partielle. Ainsi Pareto [Pareto 1896] proposa une approche hiérarchique permettant de déterminer si une solution est de meilleur compromis. Dans les définitions suivantes nous considérons un problème de maximisation.

Définition 10 (Pareto Dominance) Soit $z, z' \in Z$ deux solutions dans l'espace des critères. On dit que z domine z' ssi $\forall i \in \{1, 2, \dots, n\}, z_i \geq z'_i$ et $\exists j \in \{1, 2, \dots, n\}$ tel que $z_j > z'_j$. La relation " z domine z' " sera notée $z \succ z'$. Cette relation de dominance est étendue à l'espace de décision : soient $x, x' \in X$ deux solutions de l'espace de décision vérifiant $f(x) \succ f(x')$. On dira alors que x domine x' ou encore $x \succ x'$.

Définition 11 (Solution non-dominée ou Pareto optimale) Soit $z \in Z$ une solution dans l'espace des critères. On dit que z est non-dominée ssi $\forall z' \in Z, z' \not\succ z$.

De même soit $x \in X$ une solution dans l'espace de décision. On dit que x est non-dominée ssi $\forall x' \in X, x' \not\succ x$. De telles solutions sont dites Pareto optimales.

Définition 12 (Incomparabilité) Soient $z, z' \in Z$ deux solutions dans l'espace des critères. z et z' sont dites incomparables ssi $z \not\succ z'$ et $z' \not\succ z$. Par extension on dira que x et x' (éléments de X) sont incomparables ssi $f(x)$ et $f(x')$ sont incomparables.

Définition 13 (Ensemble Pareto optimal) Sous les notations précédentes, l'ensemble Pareto optimal est défini comme suit : $X_P = \{x \in X \mid \nexists x' \in X, x' \succ x\}$.

Définition 14 (Front de Pareto) Etant donné un ensemble Pareto optimal X_P , le front de Pareto est défini comme suit : $Z_P = \{f(x) \mid x \in X_P\}$.

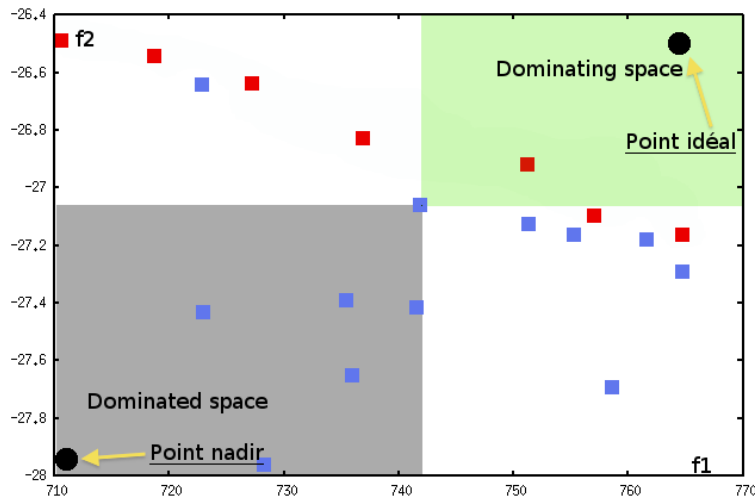


FIGURE 3.2 – Pareto dominance et incomparabilité

La Figure 3.2 illustre un ensemble de solutions, avec en rouge les solutions non-dominées (dans le cas d'un problème de maximisation sur les objectifs f_1 et f_2), et en bleu les solutions dominées. A partir de la solution bleue au centre, le rectangle gris (en bas) contient les solutions qu'elle domine et celui vert clair (en haut) les solutions qui la dominent. Un tel ensemble fait apparaître deux points virtuels fréquemment utilisés, le point idéal, composé des meilleurs éléments de toutes les solutions et le point nadir, composé des plus mauvais éléments de toutes les solutions.

Définition 15 (Point idéal) Le point idéal, fréquemment noté z^* , est composé des meilleures valeurs de chaque objectif : $z^* = (z_1^*, \dots, z_n^*)$ tel que $z_i^* = \max_{x \in X} f_i(x), \forall i \in \{1, \dots, n\}$.

Définition 16 (Point nadir) *Le point nadir, fréquemment noté z^n , est composé des plus mauvaises valeurs de chaque objectif : $z^n = (z_1^n, \dots, z_n^n)$ tel que $z_i^n = \min_{x \in X} f_i(x), \forall i \in \{1, \dots, n\}$.*

3.3 Métaheuristiques pour les problèmes multiobjectifs

3.3.1 Classification des principales approches

Depuis que les MOP sont abordés par la communauté scientifique, différentes approches de résolution ont été proposées. Dans la littérature on retrouve fréquemment une classification des méthodes en trois grandes familles :

- les approches basées sur la transformation du MOP en un problème mono-objectif. On trouvera dans cette catégorie les méthodes basées sur l'agrégation des fonctions objectif f_i en une seule fonction f . Ceci impose, outre une bonne connaissance du problème à traiter, d'être en mesure de quantifier l'apport recherché pour chaque objectif.
- les approches non-Pareto. Il s'agit d'approches qui considèrent tous les objectifs mais de façon désagrégée, indépendamment les uns des autres.
- les approches Pareto. Pour cette classe d'algorithmes, toutes les fonctions objectif sont considérées en même temps en exploitant le principe de non-dominance.

Les différents travaux menés dans le domaine ont montré que les approches Pareto sont généralement les plus performantes.

Nous utiliserons ce découpage dans cette section, lequel est synthétisé dans la Figure 3.3.

De façon transversale à ces différentes classes, les approches hybrides offrent de nombreux avantages et permettent de bénéficier des mécanismes éprouvés dans les heuristiques classiques. Bien souvent, tout comme dans le cas mono-objectif, les approches hybrides surpassent les métaheuristiques pures, comme dans [Knowles 2000], [Jaszkiewicz 2002] ou encore [Deb 2001].

C'est dans cette dernière classe que s'inscrit nos travaux présentés dans cette thèse.

3.3.2 Les approches par transformation mono-objectif

Ces approches fournissent des moyens souvent assez simples d'aborder un MOP. La complexité liée à la prise en compte simultanée de différents objectifs contradictoires

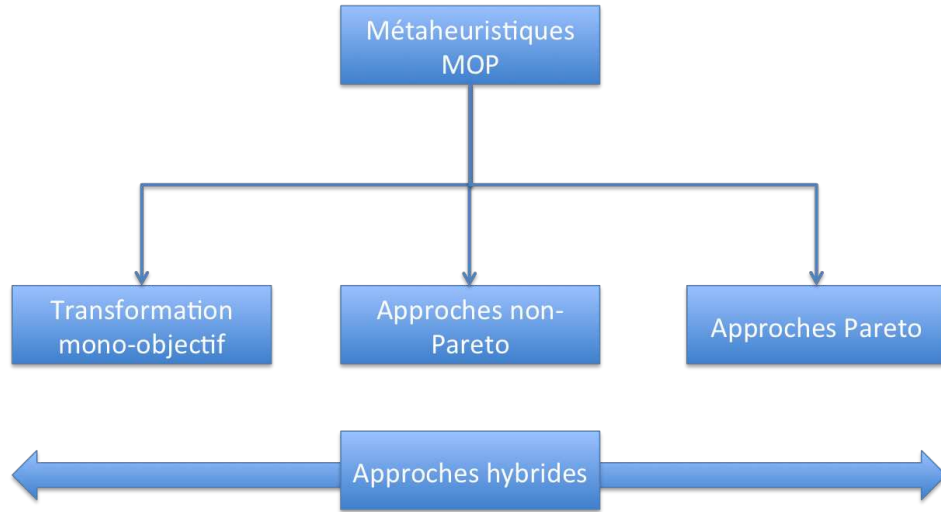


FIGURE 3.3 – Classification des méthodes de résolution d'un MOP

est ainsi éludée. Parmi les méthodes de transformation couramment rencontrées on remarquera particulièrement les méthodes d'agrégation, les méthodes ε -contrainte, et les méthodes de programmation par but. Il s'agit encore aujourd'hui d'approches intéressantes pour commencer à étudier un MOP. De plus, de part leur simplicité tant dans l'implémentation que l'exécution, ces méthodes fournissent une bonne base pour élaborer des méthodes hybrides, notamment en les couplant avec des algorithmes génétiques.

3.3.2.1 La méthode d'agrégation

Il s'agit incontestablement d'une des premières méthodes utilisée pour approcher les MOP. Ici le problème initial (MOP) est transformé en un problème (MOP_λ) où les fonctions objectif f_i sont combinées pour aboutir à une seule fonction F . Le plus souvent il s'agit d'une agrégation linéaire :

$$F(x) = \sum_{i=1}^n \lambda_i f_i(x) \quad (3.4)$$

où les poids λ_i sont tels que $\lambda_i \in [0..1]$ et $\sum_{i=1}^n \lambda_i = 1$.

Remarquons qu'une des principales caractéristiques de la méthode d'agrégation provient du fait que la solution optimale dépend directement de la pondération choisie. La Figure 3.4 illustre ce principe pour deux objectifs f_1 et f_2 .

Le vecteur de poids λ définit un hyperplan (en pointillé sur la figure) de progression

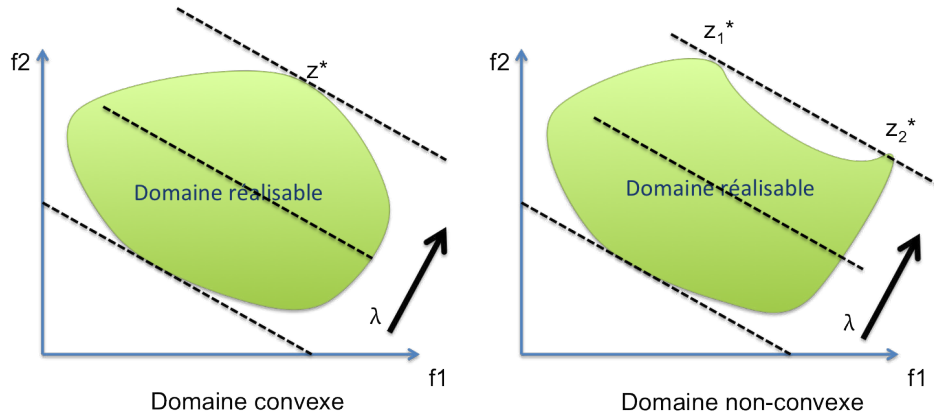


FIGURE 3.4 – Approche par agrégation des objectifs

dans l'espace objectif (une droite dans un espace à deux dimensions). La solution optimale z^* est obtenue au point de tangence entre cet hyperplan et l'espace réalisable du problème. Lorsque le front de Pareto est concave, plusieurs solutions optimales peuvent exister.

Comme nous l'avons déjà évoqué, la principale difficulté liée à cette approche réside dans le choix du vecteur λ . Cela suppose de bien connaître le problème traité et d'être capable d'exprimer et de quantifier des préférences associées aux objectifs.

De plus, le fait de fixer λ restreint fortement et prématurément l'espace de recherche. C'est pourquoi cette méthode est souvent utilisée en faisant varier les valeurs de λ . En particulier dans [Ishibuchi 1998] la méthode d'agrégation est utilisée dans un algorithme hybride mêlant génétique et recherche locale. A chaque sélection de parents le vecteur λ est régénéré aléatoirement.

Le dernier point à remarquer à propos de cette méthode concerne la commensurabilité des objectifs. En effet, pour éviter qu'un objectif soit sur-représenté ou sous-représenté dans la fonction agrégée, il importe de s'assurer que tous les objectifs soient dans la même plage de valeurs. Aussi la fonction d'agrégation peut selon les besoins être transformée sous la forme :

$$F(x) = \sum_{i=1}^n c_i \lambda_i f_i(x) \quad (3.5)$$

où les c_i sont des constantes de mise à l'échelle des objectifs. Pour fixer les valeurs de ces constantes on a souvent recours au vecteur idéal.

Une fois les objectifs agrégés, toutes les approches habituellement rencontrées pour résoudre des problèmes mono-objectif peuvent être appliquées. On rencontre

notamment la recherche locale, simple avec la méthode TPLS [Paquete 2003], mais aussi impliquant des mécanismes Tabou [Gandibleux 1997] et l'introduction de la méthode Tabou MOTS [Hansen 1997]. La méthode MOSA [Ulungu 1999] repose quant à elle sur un recuit simulé. Enfin de nombreux algorithmes génétiques sont également utilisés conjointement à la méthode d'agrégation, tels que [Yang 1994] pour le problème bi-critère de transport.

3.3.2.2 La méthode ε -contrainte

Avec cette approche, un seul objectif est retenu pour être optimisé. Par contre des contraintes sont appliquées aux autres objectifs pour assurer un niveau de qualité suffisant. Le problème initial (*MOP*) devient alors :

$$MOP_k(\varepsilon) \begin{cases} \max & z = f_k(x) \\ \text{tel que :} & x \in X \\ \text{s.c.} & f_j(x) \geq \varepsilon_j, j = 1, \dots, n \text{ où } j \neq k \end{cases} \quad (3.6)$$

ε est ainsi le vecteur des contraintes sur chaque objectif : $\varepsilon = (\varepsilon_1, \dots, \varepsilon_{k-1}, \varepsilon_{k+1}, \varepsilon_n)$

Nous illustrons sur la Figure 3.5 le cas d'un problème de maximisation. La contrainte ε_2 est appliquée à la fonction objectif 2 : la solution ε acceptable ne peut être inférieure à cette valeur. De fait la solution optimale sur f_1 devient la solution z^* .

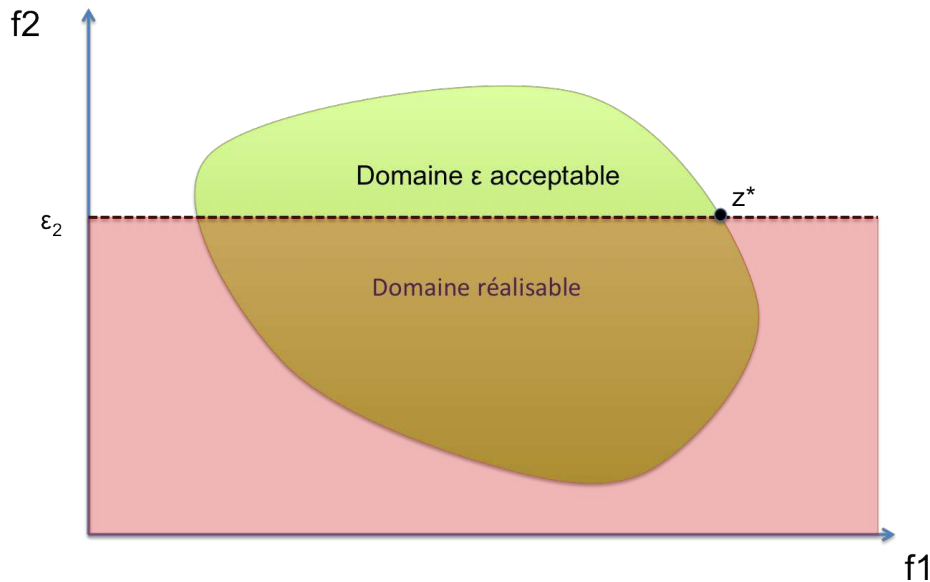


FIGURE 3.5 – Approche par methode ε -contrainte

Avec cette approche l'important est de pouvoir classer les objectifs par ordre de priorité et surtout d'être capable de fixer les valeurs ε souhaitées. Une démarche peut être décrite comme suit :

- résoudre le problème pour la fonction objectif principale f_i seule ;
- définir la valeur ε_i associée. Elle correspond au taux de dégradation acceptable (ex. on veut au moins 80% de l'optimalité) ;
- résoudre le problème pour la fonction objectif secondaire f_j avec la contrainte ε_i sur f_i ;
- définir la valeur ε_j associée. Ainsi de suite...

Les travaux ayant recours à la méthode ε -contrainte sont nombreux. Là encore, le MOP étant ramené à un problème mono-objectif, les approches courantes telles que les algorithmes génétiques, les recherches locales, etc. sont fréquemment utilisées.

3.3.2.3 La programmation par but

Ici le principe est de définir un point artificiel à atteindre dans l'espace objectif. Ainsi le décideur fixe le but visé pour chaque objectif. Le MOP initial est ensuite transformé en un problème mono-objectif qui intègre, dans la formulation du problème lui-même, la notion de but à atteindre. A titre d'exemple un nouveau problème peut être de minimiser la déviation au but. Le problème devient alors :

$$MOP(z^B) \left\{ \begin{array}{l} \min \left(\sum_{j=1}^n |f_j(x) - z_j^B|^p \right)^{\frac{1}{p}} \\ \text{tel que : } x \in X \end{array} \right. \quad (3.7)$$

Dans cette nouvelle formulation du MOP le but à atteindre est défini par z^B . La distance au but utilisée ici est la métrique de Tchebycheff avec $1 \leq p \leq \infty$. Dans le cas où $p = 2$ on retrouve la distance euclidienne. De même lorsque $p \rightarrow \infty$ on retrouve un Min-Max.

La Figure 3.6 illustre le déplacement d'une solution vers le but z^B . Là encore le choix des paramètres est déterminant pour la bonne marche de l'algorithme. Un but mal positionné peut altérer considérablement le résultat obtenu. En particulier si le but est situé à l'intérieur du domaine réalisable, l'algorithme ne pourra pas trouver une solution Pareto optimale. Or le choix du but dépend du décideur et nécessite une très bonne connaissance du problème traité.

Souvent le but est fixé en utilisant là encore le point idéal. Parmi les travaux utilisant la méthode programmation par but on citera l'algorithme génétique de [Coello Coello 1998] qui utilise une fonction min-max relativement au point idéal.

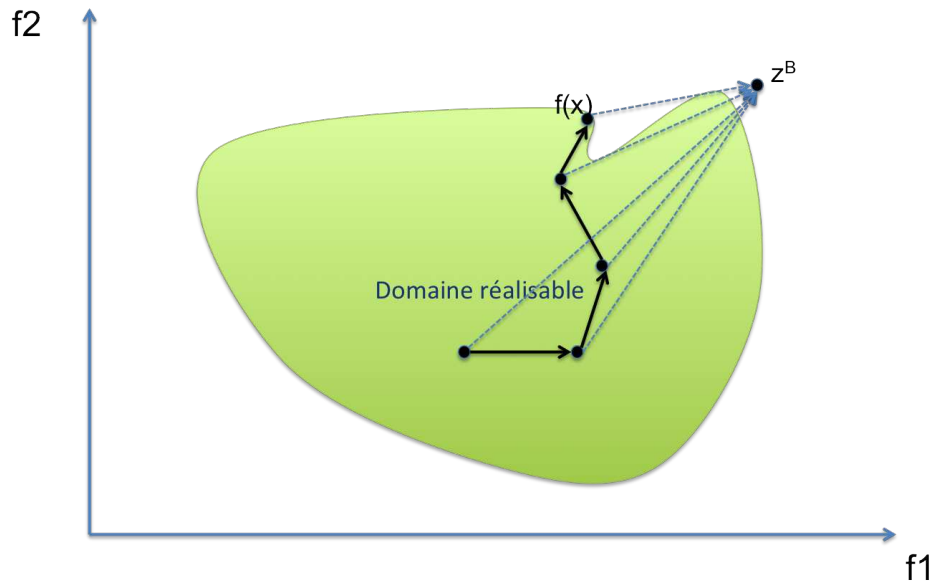


FIGURE 3.6 – Approche par but

3.3.3 Les approches non-Pareto

Cette catégorie d'approches regroupe les méthodes qui considèrent tous les objectifs dans le même algorithme, sans passer par la transformation du MOP en un problème mono-objectif. Par contre, contrairement aux approches Pareto qui sont bien plus répandues, les approches non-Pareto traitent chaque fonction objectif indépendamment les unes des autres.

On ne présentera ici que l'algorithme VEGA (Vector Evaluated Genetic Algorithm) [Schaffer 1985], qui est le premier AG réellement multiobjectif. Dans cet algorithme, l'opérateur de sélection est basé sur un objectif seulement, indépendamment des autres. A chaque génération la population est divisée en autant de groupes qu'il y a de fonctions objectif. Chaque groupe est évalué suivant l'objectif correspondant. Ainsi les individus du groupe i sont évalués selon la fonction objectif f_i . L'affectation de chaque individu à un groupe se fait aléatoirement et à chaque génération. Ainsi chaque individu appartiendra à un groupe potentiellement différent à chaque génération, et sera évalué sur chaque objectif.

Une des principales limites liées à cette approche vient du fait même que chaque objectif est considéré indépendamment des autres. La sélection ne tient en effet compte que d'un objectif, ce qui a tendance à favoriser les meilleurs individus pour cet objectif. De fait les solutions de compromis sont pénalisées. Pour pallier ce problème, la prise en compte simultanée de tous les objectifs s'avère être nécessaire. C'est ce que permettent les approches de type Pareto.

3.3.4 Les approches Pareto

Ces approches utilisent la relation de dominance à l'intérieur même des algorithmes. Ainsi les solutions ne sont plus évaluées vis à vis d'une fonction objectif mais vis à vis d'un compromis impliquant toutes les fonctions objectif. Depuis les premiers travaux basés sur la notion de Pareto, où [Goldberg 1989] introduit la notion de rang de dominance d'une solution, de nombreux algorithmes sont apparus. Du fait que la notion même de Pareto implique la prise en compte d'une population, on retrouve naturellement de nombreux AG appartenant à cette famille. Mais comme nous le verrons, il y a également des recherches locales basées sur la notion de Pareto, ainsi que des approches hybrides. Plus récemment, les méthodes à base d'indicateurs sont apparues. Au lieu de considérer les solutions de la population, c'est la population elle-même qui est évaluée au travers d'indicateurs choisis.

3.3.4.1 Les algorithmes génétiques Pareto

Les opérateur de "ranking" Dès lors qu'une solution n'est plus évaluée par rapport à ses fonctions objectif, il importe de définir une métrique capable de représenter le niveau de qualité d'une solution. Ainsi plusieurs méthodes ont été proposées et sont souvent associées à des AG couramment utilisés. Les trois principales répertoriées dans [Zitzler 2004] sont présentées ici.

Le rang de dominance : Il est défini, pour une solution s donnée, par le nombre de solutions qui domine s incrémenté de 1 [Fonseca 1993]. Ainsi les solutions de rang 1 sont les solutions non-dominées et plus une solution est dominée par un nombre important de solutions de la population, plus son rang se détériore. Le rang de dominance a été appliqué à l'origine dans la méthode MOGA (Multiobjective Genetic Algorithm) [Fonseca 1993].

Le compte de dominance : Ici, à l'inverse du rang de dominance, c'est le nombre de solutions dominées par la solution s qui détermine son rang. Ainsi un rang élevé traduit une bonne qualité de la solution.

La profondeur de dominance : Dans cette approche les solutions sont regroupées par fronts de Pareto [Goldberg 1989], les solutions non-dominées de la population reçoivent la valeur 1 puis sont retirées de la population. Les nouvelles solutions non-dominées reçoivent alors la valeur 2 puis sont également retirées de la population. Et ainsi de suite jusqu'à avoir classé toutes les solutions. Comme on le verra plus tard, toutes les solutions d'un même front ayant la même valeur impose de recourir à des mécanismes complémentaires pour les distinguer. Cette définition de la profondeur de dominance est utilisée dans les algorithmes NSGA (Non-dominated Sorting Genetic Algorithm) [Srinivas 1994] et NSGA-II [Deb 2002].

Remarque 1 : On notera qu'il est possible de combiner plusieurs évaluations de la dominance. C'est ce qui est notamment fait dans SPEA [Zitzler 1999] et SPEA2 [Zitzler 2001] où sont utilisés conjointement la profondeur et le rang de dominance.

Remarque 2 : Des travaux récents introduisent d'autres types de dominance, tels l' ε -dominance dans [Deb 2005], ou encore la g-dominance [Molina 2009].

Nous abordons maintenant quelques algorithmes en détail parmi les plus connus.

NSGA-II (Non-dominated Sorting Genetic Algorithm) [Deb 2002] est l'algorithme génétique le plus souvent rencontré dans la littérature pour résoudre des MOP. Cette deuxième version de NSGA [Srinivas 1994] utilise deux critères pour affecter la valeur de fitness aux solutions. Dans un premier temps les solutions sont regroupées par front en utilisant la profondeur de dominance [Goldberg 1989]. Ce premier critère reflète la qualité des solutions vis à vis de la convergence. Dans un deuxième temps, chaque solution d'un front est évaluée en fonction de la densité de solutions se trouvant à proximité, grâce à la fonction de *crowding*. Ce deuxième critère reflète la qualité des solutions vis à vis de la divergence.

Ainsi NSGA-II peut être caractérisé par :

- un tri rapide des solutions non-dominées selon la profondeur de dominance
- une estimation de la densité liée à l'environnement de chaque solution, généralement appelée distance de *crowding*
- un tri à deux niveaux : selon la profondeur de dominance d'abord, puis pour les solutions d'un même niveau selon la distance de *crowding*

Pour ce qui est du fonctionnement, deux populations sont utilisées, celle des parents P et celle des enfants Q , toutes deux de taille N . On notera R la population qui résulte de l'union de P et Q , avec $|R| = 2N$. Les éléments de chaque population évoluant au cours du temps, on notera P_t , Q_t , et R_t , l'état des ensembles à l'instant t . Les populations sont utilisées comme suit :

1. Réduction de R_t . Les solutions de R_t sont réparties dans les fronts de dominance F_i , puis les solutions de chaque front sont ordonnées selon la distance de *crowding*. Les N meilleures solutions du meilleur au plus mauvais front sont conservées dans P_{t+1} .
2. Sélection des parents. Les parents sont choisis par tournoi binaire avec remise parmi les solutions de P_{t+1} , basé sur la double fitness établie.
3. Création de Q_{t+1} . Les enfants sont générés par croisement et mutation à définir selon le problème traité.

L'algorithme 1, qui sera utilisé dans la suite de cette thèse, formalise ces différentes étapes.

Algorithm 1 NSGA-II

Require: population size N , generation number $nbIter$ /* $nbIter$ can be replaced by a time limit*/

- 1: $P \leftarrow \text{init}(N)$ /*init population P with N random individuals*/
- 2: $Q \leftarrow \emptyset$ /*init an empty children population*/
- 3: $\text{eval}(P)$ /*eval objectives for each individual*/
- 4: **for** $i = 1$ to $nbIter$ **do**
- 5: $P \leftarrow P \cup Q$
- 6: $\text{assignRank}(P)$ /*based on Pareto dominance*/
- 7: **for** each non-dominated front $f \in P$ **do**
- 8: $\text{setCrowdingDist}(f)$
- 9: **end for**
- 10: $\text{sort}(P)$ /*by rank and in each rank by the crowding dist */
- 11: $P \leftarrow P[0 : N]$
- 12: $Q \leftarrow \text{buildNextGeneration}(P)$ /*Binary Tournament Selection, Recombination and Mutation*/
- 13: **end for**

Remarque : La distance de *crowding* d'une solution s est calculée comme la longueur moyenne des côtés de l'hypercube (cf. Figure 3.7) défini par les plus proches solutions suivantes et précédentes dans l'espace des critères et appartenant au même front.

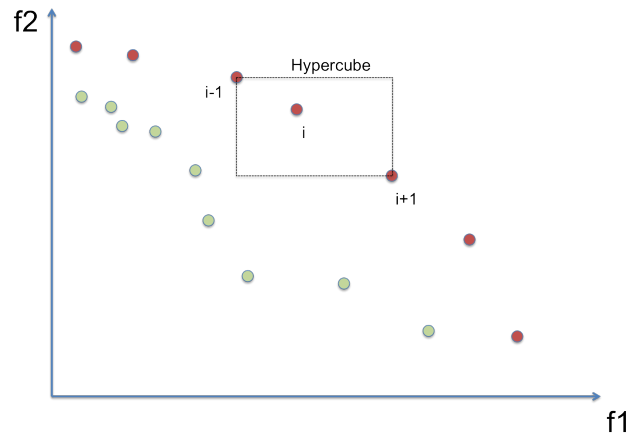


FIGURE 3.7 – Distance de *crowding* : moyenne des côtés de l'hypercube

SPEA2 (Strenght Pareto Evolutionary Algorithm) [Zitzler 2001] est une version améliorée de SPEA. En plus de la population P servant à sélectionner les parents de l'AG, cette approche utilise une archive complémentaire A , de taille fixe, pour stocker N_A solutions non-dominées au cours de la recherche. Au cas où le nombre de

solutions non-dominées est inférieur à N_A , l'archive est complétée avec les meilleurs individus dominés de la population P , en utilisant le critère de compte de dominance. Pour ce faire, à chaque solution s de la population $P \cup A$ est attribuée la valeur $g(s)$ correspondant au nombre de solutions de $P \cup A$ que s domine :

$$g(s) = |\{s' : s' \in P \cup A \text{ et } s \succ s'\}| \quad (3.8)$$

La valeur de fitness $h(s)$ attribuée à la solution s est égale à la somme des $g(s')$ pour tout s' dominant s :

$$h(s) = \sum_{s' \in P \cup A, s' \succ s} g(s') \quad (3.9)$$

Ainsi une solution s non dominée est telle que $h(s) = 0$. Inversement, plus $h(s)$ augmente, moins bonne est la solution s .

Bien que le compte de dominance ait tendance à permettre une certaine diversité, il met surtout en avant l'intensification. Pour distinguer les solutions ayant le même compte de dominance, SPEA2 intègre un deuxième mécanisme dédié au "sharing" pour favoriser la diversité. Il s'agit d'une adaptation de la méthode du k ème plus proche voisin [Silverman 1986], où la densité d'un point est inversement proportionnelle à la distance le séparant du k ème plus proche voisin. La densité d'une solution s est décrite par :

$$D(s) = \frac{1}{\sigma_s^k + 2} \quad (3.10)$$

avec : σ_s^k la distance entre la solution s et son k ème plus proche voisin.

On notera que k est généralement fixé à $k = \sqrt{N + N_A}$, où N est la taille de la population P et N_A la taille de l'archive A . La valeur de fitness globale utilisée dans SPEA2 est alors :

$$F(s) = h(s) + D(s) \quad (3.11)$$

ε -MOEA (ε -MultiObjective Evolutionary Algorithm). Cet algorithme proposé dans [Deb 2003] se rapproche de SPEA, avec l'intégration du principe d' ε -dominance. Ainsi, tout comme avec IBEA que l'on verra plus loin [Zitzler 2004], l'idée est d'enrichir la notion de rang en y ajoutant la notion de diversité. L'espace de recherche est découpé en sous-espaces, définis comme des hyper-boîtes, comportant chacun au

plus une solution. Ce principe permet de garantir une bonne couverture de l'espace de recherche.

Là encore deux populations évoluent parallèlement. La population P pour l'algorithme génétique lui-même, et la population archive A pour stocker les solutions non-dominées rencontrées. A chaque itération deux individus a et p sont sélectionnés, respectivement de A et P , et sont utilisés pour créer k enfants c_i , avec $i = 1 \dots k$. On notera que la sélection de a se fait aléatoirement alors que pour p un tournoi binaire est utilisé.

Les enfants c_i ont la possibilité d'intégrer les deux populations, selon des critères différents. Le critère retenu pour intégrer P est la dominance au sens de Pareto. Par contre pour intégrer l'archive le critère est plus sélectif : c'est l' ε -dominance qui est utilisée.

Dès lors on distingue trois cas pour l'intégration de c_i dans P :

- si c_i domine au moins une solution de P , elle remplace dans P aléatoirement une des solutions dominées
- si c_i est dominée par au moins une solution de P , elle est rejetée
- sinon, c_i remplace une solution de P choisie aléatoirement

De même on distingue trois cas pour l'intégration de c_i dans A :

- si c_i domine au moins une solution de A , elle remplace dans A toutes les solutions dominées
- si c_i est dominée par au moins une solution de A , elle est rejetée
- sinon, c_i remplace une solution de A choisie aléatoirement

SEEA (Simple Elitist Evolutionary Algorithm) proposé par [Liefvooghe 2010] est une approche génétique qui ne nécessite pas de fixer de paramètres. Elle est particulièrement adaptée pour résoudre des problèmes pour lesquels l'évaluation des solutions ne demande que peu de temps. Dans ce cas, elle présente l'avantage de converger rapidement.

Dans cette approche les parents sont directement issus de l'archive A . Il n'y a donc pas d'autres populations que l'archive. A chaque génération, N solutions de A sont choisies aléatoirement. Elles fournissent les parents pour y appliquer les opérateurs de croisement et de mutation.

Les enfants c_i obtenus sont intégrés à l'archive selon le critère de dominance :

- si c_i est dominé par un élément de A , il est rejeté
- si c_i domine des solutions de A , il remplace dans A toutes les solutions dominées
- sinon c_i est intégré à A

3.3.4.2 Les approches à base de recherche locale

Contrairement aux AG très répandus pour résoudre des MOP, les recherches locales sont plus rares, du moins pour ce qui est des approches Pareto. En effet, comme nous l'avons vu la notion même de Pareto suppose de considérer une population ce qui la rend particulièrement adaptée aux AG. Mais il est tout de même possible d'effectuer une recherche locale, donc à solution unique, sur la base d'une relation de dominance, donc par rapport à une population.

Les recherches locales, dans leur mode de fonctionnement, requièrent une structure de voisinage.

Définition 17 (Structure de voisinage) *Une structure de voisinage N est une fonction de $X \rightarrow 2^N$ qui, à toute solution $x \in X$, attribue un certain nombre de solution $N(x) \subset X$. De fait, $N(x)$ est le voisinage de x , $x' \in N(x)$ est un voisin de x .*

Les approches par transformation d'un MOP en problème mono-objectif présentées en section 3.3.2 peuvent être utilisées conjointement à l'archivage des solutions non-dominées rencontrées. De ce fait on pourra faire évoluer la recherche locale mono-objectif en fonction de la construction de l'archive. Par exemple, les poids permettant de combiner les fonctions objectif dans la méthode par agrégation peuvent être fixés pour permettre à la recherche de s'éloigner de l'archive déjà constituée.

Nous nous focaliserons ici uniquement sur les recherches locales qui intègrent la notion de dominance. Il s'agit de la classe de problèmes appelée dans [Liefoghe 2012] DMLS (Dominance-based Multiobjective Local Search).

Les approches de type DMLS présentent la caractéristique commune de reposer sur une archive de solutions non-dominées que l'on cherche à améliorer. D'une façon générale, l'idée est de considérer les solutions de l'archive, d'explorer leurs voisinages, et de mettre à jour l'archive avec les nouvelles solutions non-dominées rencontrées. L'archive est donc ici de taille variable. L'algorithme s'arrête spontanément lorsqu'il ne reste plus de solution de l'archive avec de nouveau voisin non-dominé.

PLS (Pareto Local Search) [Paquete 2004] a une archive A qui est initialisée avec une première solution aléatoire, marquée comme non-exploitée. A chaque itération, une solution $s \in A$ est choisie aléatoirement parmi les solutions non-exploitées. Le voisinage complet de s est alors exploré, toutes les solutions $s' \in V(s)$ sont candidates pour l'intégration à l'archive. Si aucune solution de A ne domine s' , s' est ajoutée à A . Toute solution de A dominée par s' est alors retirée de l'archive.

Cette recherche locale particulièrement répandue pour aborder les MOP est

utilisée dans le chapitre 4. Elle servira de référence pour résoudre un problème de localisation de stations d'auto-partage. Aussi nous détaillons dans l'algorithme 2 son fonctionnement.

Algorithm 2 PLS

```

1:  $t \leftarrow 0$ 
2:  $s_t \leftarrow \text{GenerateInitSol}()$  /*init the solution  $s_t$  with a random individual*/
3:  $A \leftarrow \{s_t\}$ 
4: repeat
5:   for all  $s'_t \in \text{Neighborhood}(s_t)$  do
6:      $\text{eval}(s'_t)$ 
7:     if  $s'_t$  is not dominated by any  $r \in A$  then
8:        $A \leftarrow \text{UpdateArchive}(s'_t)$  /*add  $s'_t$  in  $A$  and remove all dominated solutions*/
9:     end if
10:  end for
11:  mark  $s_t$  as explored
12:   $t \leftarrow t + 1$ 
13:   $s_t \leftarrow \text{select}(A)$  /*select randomly a not explored solution from  $A$  */
14: until all solutions in  $A$  are explored

```

PAES (Pareto Archived Evolution Strategy) [Knowles 1999] présente une recherche locale relativement différente de PLS. Là encore une solution initiale est générée aléatoirement, évaluée pour chaque objectif, et intégrée à l'archive A . Ici l'archive est de taille bornée. A chaque cycle l'algorithme choisit une solution $s \in A$, la recopie dans s' , puis applique un opérateur de mutation à s' . A ce titre cet algorithme est parfois classé dans la famille des AG. Toutefois il n'y a ni croisement, ni sélection de parents. En fait la mutation définit l'opérateur de la recherche locale. Il y a alors une comparaison entre s et s' .

Si s' domine s , elle la remplace dans A . Si s domine s' , s' est rejetée. Par contre si s et s' sont non-comparables, s' est comparée à l'archive. Si elle domine une solution de A elle est intégrée. Si elle est dominée par une solution de A elle est rejetée. Sinon (il s'agit d'un meilleur compromis), l'introduction ou non dans l'archive se fait sur la base d'une hyper-grille, de façon à favoriser la diversité de l'archive.

3.3.4.3 Optimisation à base d'indicateur

Les approches à base d'indicateur constituent incontestablement une des orientations phare de ces dernières années. Alors que les indicateurs sont souvent utilisés pour évaluer le niveau de performance des ensembles de solutions (cf section 3.4), ici ils sont incorporés au sein même de l'algorithme. Les approches de ce type sont

particulièrement appréciées car elles permettent de prendre en compte, dans les mécanismes de la recherche, les préférences de l'utilisateur. Parmi les heuristiques à base d'indicateur, il existe des approches génétiques ou locales. Nous les avons toutefois regroupées dans une section à part en raison de leur originalité.

IBEA (Indicator Based Evolutionary Algorithm) proposé par [Zitzler 2004] est un algorithme évolutionnaire à base d'indicateur. Jusque-là nous avons vu des approches génétiques basées sur des opérateurs de *ranking* qui étaient très efficaces en matière d'intensification, mais qui nécessitaient des mécanismes complémentaires pour bien gérer la diversité des solutions. Ici le problème est réglé en considérant un indicateur binaire I . Il est à définir en fonction des orientations à donner à l'algorithme, pour équilibrer selon les besoins les phases d'intensification et de diversification. I représente ainsi le but global que l'on souhaite viser. Il permet de déterminer pour chaque solution x de la population P ce qu'elle apporte pour atteindre cet objectif. De ce fait la fitness $F(x)$ attribuée à x correspond à la perte de qualité qu'engendrerait son absence, plus $F(x)$ est grande plus x a de l'importance dans la population. La formulation de $F(x)$ retenue dans IBEA est :

$$F(x) = \sum_{x' \in P \setminus \{x\}} -e^{-I(x', x)/k} \quad (3.12)$$

avec : k un facteur d'échelle.

L'opérateur de sélection des parents est un tournoi binaire basé sur la fitness F . Après génération des enfants, la population est réduite pour retrouver sa taille initiale N . Cette réduction se fait par suppression successive des plus mauvais individus, avec mise à jour des valeurs de fitness après chaque suppression.

Bien que la méthode n'impose pas d'indicateur particulier, dans [Zitzler 2004] est proposé l'indicateur ε -additif :

$$I_{\varepsilon+}(z, z') = \max_{i \in \{1, \dots, n\}} \{z_i - z'_i\} \quad (3.13)$$

Cet indicateur calcule ainsi la translation minimale nécessaire pour que z domine faiblement z' . Un autre indicateur est également proposé : l'indicateur hypervolume différentiel I_{HD} [Zitzler 1999]. Mais il est à noter que tout autre indicateur peut être utilisé selon l'effet recherché.

IBMOLS (Indicator-Based evolutionary Multi-Objective Local Search) s'appuie sur la méthode IBEA [Zitzler 2004] pour définir une recherche locale à base d'indicateur,

celui-ci n'étant pas imposé par l'algorithme. C'est donc à chacun de définir, selon le but à atteindre, l'indicateur approprié.

IBMOLS prend en entrée un indicateur I et une population initiale P de taille k (qui peut être composée de k solutions aléatoires). A chaque itération, le voisinage complet de toutes les solutions de P est exploré. Tous les voisins sont candidats à l'intégration dans P . Les solutions non-dominées par P y sont ajoutées. Les solutions de P dominées par une solution nouvellement intégrée sont supprimées. A ce stade la taille de la population peut avoir considérablement grossi. Une réduction de P est réalisée en retirant successivement la solution ayant la plus mauvaise contribution au sens de l'indicateur. Cette réduction est faite jusqu'à ce que P soit à nouveau de taille k .

L'algorithme IBMOLS s'arrête spontanément lorsqu'aucun voisin de l'archive ne peut plus être ajouté. Les auteurs précisent que cet algorithme décrit un cadre général, mais qu'il est possible de définir des variantes.

Approches à base de l'indicateur hypervolume Parmi les indicateurs actuellement très utilisé, l'indicateur hypervolume ressort tout particulièrement. Dans [Brockhoff 2008], une analyse approfondie de ces approches est proposé. On renverra notamment à SMS-EMOA [Beume 2007]. Dans cet AG, l'opérateur de sélection est basé sur l'indicateur hypervolume (maximisation de l'hypervolume dominé) combiné à un tri selon la non-domination. On trouvera également l'approche MO-CMA-ES (MultiObjective-Covariance Matrix Adaptation-Evolution Strategy) [Igel 2007], qui s'appuie sur l'algorithme mono-objectif CMA-ES [Hansen 2001]. Enfin [Zitzler 2007] propose une méthodologie permettant de définir une mesure de la qualité qui s'appuie sur la mesure de l'hypervolume.

3.3.5 Les approches hybrides

Afin de résoudre un MOP, nous avons présenté jusqu'à présent les méthodes par transformation d'un problème mono-objectif, les méthodes d'évolution parallèle, également appelées non-Pareto, et les approches Pareto basées sur la relation de dominance. Ces approches dites "pures" ont été les seules métaheuristiques développées jusque dans les années 1990. Les premières approches hybrides apparurent à cette époque. Elles consistent à incorporer à l'intérieur d'un de ces algorithmes des éléments provenant d'une autre méthode. De ce fait elles ont la particularité de fournir des algorithmes pouvant appartenir à une des classes présentées, mais aussi éventuellement à plusieurs. De plus il est également possible de coupler une approche de type heuristique avec une méthode exacte.

Avec les approches à base d'indicateurs, les approches hybrides connaissent aujourd'hui un véritable attrait. En effet, l'hybridation d'heuristiques a montré dans

le contexte mono-objectif les améliorations qu'elle pouvait apporter, et maintenant également pour aborder les MOP. Lorsque l'hybridation consiste à intégrer une recherche locale à l'intérieur d'un algorithme génétique, souvent à la place de l'opérateur de mutation, mais pas nécessairement, on parle d'algorithmes mémétiques. Dans le chapitre 4 un algorithme de ce type sera proposé.

Dans cette section nous listons uniquement les classes de modèles rencontrées et connues pour être efficaces. Le livre [Talbi 2013] consacre notamment un chapitre à ce sujet. Nous suivrons la classification que Talbi propose.

3.3.5.1 Low-Level Relay Hybrids (LRH)

Seules quelques hybridations de ce type existent. Il s'agit ici d'incorporer une métaheuristique dans une S-métaheuristique (métaheuristique à base de solution unique telle une recherche locale). Comme nous l'avons vu plus haut, les approches à base de solution unique ne sont pas bien adaptées à la résolution d'un MOP qui admet un ensemble de solutions. Aussi ce type d'hybridation n'est pas souvent utilisé

3.3.5.2 Low-Level Teamwork Hybrids (LTH)

Ici au contraire une métaheuristique est intégrée dans une P-métaheuristique (métaheuristique à base de population tel un AG). Il s'agit de l'hybridation la plus utilisée. En effet, comme nous l'avons vu dans la section 3.3.4, les approches à base de population telles que les approches Pareto présentent le réel avantage de considérer la population complète lors de la recherche. Elles ont toutefois le défaut de leur qualité : elles convergent plus lentement que les S-métaheuristicques. Ainsi l'hybridation LTH vise à rassembler les avantages des P-métaheuristicques et des S-métaheuristicques au sein d'un même algorithme, les premières apportant la diversité alors que les deuxièmes apportent l'intensité.

Les S-métaheuristicques utilisées dans ce cadre peuvent être de différentes natures, allant de la recherche locale au recuit simulé, en passant par une Recherche Tabou. Les algorithmes mémétiques sont donc de cette nature.

Parmi les plus connus on rencontre [Ishibuchi 1998]. Il s'agit d'un AG par transformation du MOP en mono-objectif par agrégation. Sa particularité est que le vecteur de poids servant à l'agrégation est généré aléatoirement et pour chaque sélection de parents. Ainsi deux parents ne sont pas sélectionnés sur la base de la même fonction agrégée. Par la suite une recherche locale sur un voisinage réduit est appliquée à chaque enfant.

On citera également [Jaszkiewicz 2002] qui repose sur un mécanisme particulier.

A chaque itération une fonction d'utilité F_t est construite aléatoirement. Elle permet de sélectionner une sous-population composée des meilleurs individus au sens de cette définition de la fitness. C'est sur cet ensemble qu'est appliqué l'AG. Chaque enfant est alors intensifié par une recherche locale.

Dans [Basseur 2003] une approche hybride est elle-même hybridée avec un AG. En fait cela se fait sous la forme d'une alternance entre une phase d'AG (utilisant une mutation adaptative), suivi d'une phase mémétique.

Enfin une généralisation des approches mémétiques est proposée dans [Talbi 2001], où plusieurs stratégies de sélection et de maintien de la diversité sont présentées.

3.3.5.3 High-Level Relay Hybrids (HRH)

Ces approches reposent sur une succession de métaheuristiques. Il y a ainsi une alternance de différentes méthodes, ce qui permet de faire ressortir plusieurs comportements. Souvent les HRH sont composées d'une phase d'intensification avec une S-métaheuristique, suivie d'une phase de diversification avec une P-métaheuristique.

3.3.5.4 High-Level Teamwork Hybrids (HTH)

Les algorithmes considérés ici reposent également sur plusieurs métaheuristiques mais qui évoluent en parallèle tout en coopérant. En effet, nous voyons bien que les processeurs ont tendance à stagner en vitesse de calcul. Les principales marges de progression concernent la mise en parallèle de plusieurs unités de calcul. Cela peut se faire à plusieurs niveaux, en faisant coopérer des machines sous forme de clusters, en faisant coopérer les unités de calcul d'une même machine, ou encore en utilisant les coeurs massivement parallélisés des GPU (Graphics Processing Unit) en ayant recours au GPGPU (General-purpose processing on graphics processing units).

Le domaine de la recherche opérationnelle étant souvent gros consommateur en ressources de calcul, et donc en temps, la parallélisation des calculs constitue un vaste champ d'exploration. L'intégration du GPGPU dans la plate-forme ParadisEO [Melab 2013] au travers de l'environnement ParadisEO-MO-GPU en témoigne.

3.4 Evaluation des algorithmes multiobjectifs : une approche à base d'indicateurs

Nous l'avons vu, une des principales sources de difficulté lorsque l'on aborde un MOP vient du fait que le problème admet potentiellement un très grand nombre

de solutions. C'est une difficulté pour résoudre le problème d'optimisation et pour évaluer le niveau de performance d'un algorithme à travers l'évaluation de l'ensemble des résultats. Alors qu'en mono-objectif l'efficacité d'un algorithme découle immédiatement de l'évaluation de la solution générée, lors de la résolution d'un MOP ce n'est généralement pas le cas. Prenons deux ensembles A et B , les seuls cas triviaux sont lorsque toutes les solutions de A sont dominées par des solutions de B ou inversement. Dans tous les autres cas, où des solutions de A sont non-dominées par des solutions de B ou des solutions de B sont non-dominées par des solutions de A , il apparaît difficile de déterminer quel ensemble est le meilleur, c'est-à-dire quel ensemble apporte les solutions les plus intéressantes. Le recours à des indicateurs de performance est donc primordial.

On distinguera deux types d'indicateurs : les indicateurs unaires (ou absolus), capables d'attribuer directement une note à un ensemble, et les indicateurs binaires (ou relatifs), qui prennent deux ensembles en paramètres et attribuent une note à un ensemble par rapport à un autre. Nous présentons les indicateurs principaux de la littérature.

3.4.1 Indicateur de contribution

Il s'agit peut-être de l'indicateur le plus intuitif. L'indicateur binaire de contribution [Meunier 2000], généralement noté $I_C(A, B)$, détermine quel ensemble de A ou B apporte le plus de solutions non-dominées par l'autre ensemble. Il rejoint la métrique C proposée par [Zitzler 1999] peu de temps avant.

Plus formellement $I_C(A, B)$ revient à calculer :

$$I_C(A, B) = \frac{\frac{|A \cap B|}{2} + |W_A| + |N_A|}{|ND(A \cup B)|} \quad (3.14)$$

avec : W_A l'ensemble des solutions de A qui dominent au moins une solution de B et N_A l'ensemble des solutions de A non-comparables avec B .

L'élément particulièrement intéressant avec cet indicateur est que $I_C(A, B) + I_C(B, A) = 1$ et si $A = B$ alors $I_C(A, B) = I_C(B, A) = 0.5$. De plus il est rapide à calculer. On notera également que plus A est bon, plus l'indicateur $I_C(A, B)$ est élevé.

Par contre comme nous le montrerons dans le chapitre 4, I_C n'est pas un bon marqueur de la diversité.

3.4.2 Indicateur ε -additif

Cet indicateur binaire proposé dans [Zitzler 2003] est particulièrement intéressant. Il détermine la translation minimale à appliquer à toutes les solutions de A , et sur tous les objectifs, pour dominer B . Formellement ça revient à calculer :

$$I_{\varepsilon+}(A, B) = \min_{\varepsilon \in \mathbb{R}} \{ \forall x \in B, \exists x' \in A : x \preceq_{\varepsilon+} x' \} \quad (3.15)$$

L'intérêt de $I_{\varepsilon+}$ est de refléter la qualité d'un ensemble par rapport à un autre en terme d'intensité d'abord, mais également en terme de diversité. Ici plus $I_{\varepsilon+}$ est petit, meilleur est l'ensemble A par rapport à B .

Remarques :

1) Pour que cet indicateur ait un sens il est indispensable de normaliser les fonctions objectif. En effet, si l'objectif f_1 prend ses valeurs dans l'intervalle $[0, 10]$ et l'objectif f_2 dans l'intervalle $[0, 1000]$, les variations de f_1 seront écrasées par celles de f_2 .

2) Si une seule solution de B domine assez nettement A et que toutes les autres solutions de B sont légèrement dominées par A , alors B sera considéré comme meilleur. Cela revient à dire que cet indicateur considère la translation du pire des cas. Une variante qui semble intéressante serait de considérer la translation moyenne à appliquer.

3.4.3 Indicateur hypervolume

Il s'agit d'un indicateur incontournable depuis sa proposition par [Zitzler 1999]. L'indicateur hypervolume I_H présente en effet l'avantage de refléter la qualité d'un ensemble, tant en terme d'intensité que de diversité.

Cet indicateur connaît deux variantes, une version unaire et une version binaire. La version unaire de cet indicateur permet de fournir, pour un ensemble de solutions, son niveau de qualité sans nécessiter le recours à un ensemble de référence. Toutefois deux conditions sont indispensables à une bonne interprétation :

- de disposer d'une normalisation des objectifs pour qu'un objectif ne soit pas sur-représenté ou sous-représenté, tout comme avec l'indicateur ε -additif
- d'identifier un point de référence pour le calcul de l'hypervolume

La version binaire permet, par différence, de déterminer ce qu'apporte un ensemble et qui n'est pas apporté par un autre ensemble.

La version unaire nous intéressera tout particulièrement puisque c'est elle qui

est utilisée dans le chapitre 4.

La Figure 3.8 illustre ce que représente l'hypervolume, ramené à une surface dans le cas d'un problème bi-objectif. Dans cet exemple d'un problème de maximisation le point de référence est ramené à l'origine du repère. Il peut toutefois être pris ailleurs. Il peut être défini comme le point nadir d'un ensemble de référence, point virtuel constitué de la moins bonne valeur sur chaque objectif des de l'ensemble.

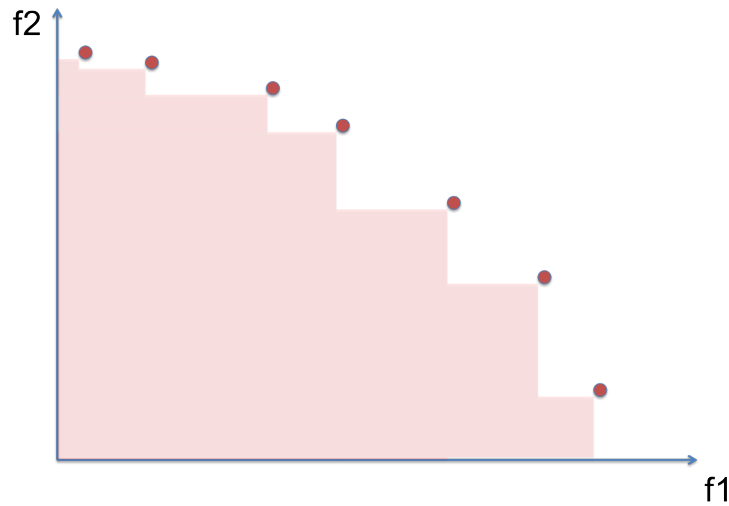


FIGURE 3.8 – Indicateur hypervolume I_H

L'indicateur I_H est particulièrement utilisé aujourd'hui en raison de son niveau de pertinence pour représenter la qualité d'un ensemble, autant en diversité qu'en intensité. Par contre, il présente le gros désavantage d'être complexe à calculer : dès que le nombre d'objectifs et de solutions est important, il nécessite un temps de calcul significatif. Plusieurs algorithmes ont été proposés dans la littérature, dont des méthodes d'approximation, pour accélérer réellement les calculs.

Cet indicateur que nous utilisons dans le chapitre 4 est calculé de façon exacte en utilisant l'algorithme HSO (Hypervolume by Slicing Objectives) [While 2006]. L'algorithme 3 présente les étapes principales de HSO.

Algorithm 3 HSO : Hypervolume by Slicing Objectives

Require: ps /* ps is a Pareto set*/

Require: n /*the number of objective functions*/

```

1:  $pl \leftarrow \text{sort}(ps, 1)$  /*sort  $ps$  worsening on objective 1*/
2:  $s \leftarrow \{(1, pl)\}$  /*the ordered Pareto set  $pl$  has level 1*/
3: for  $k \leftarrow 1$  to  $n - 1$  do
4:    $s' \leftarrow \{\}$ 
5:   for each  $(x, ql)$  in  $s$  do
6:     /*at first loop  $ql$  contains  $pl$ , at second loop  $ql$  contains successively each
       slice as described in Figure 3.9, and so on... */
7:     for each  $(x', ql')$  in  $\text{slice}(ql, k)$  do
8:       /*each set is sliced for the next level and added into  $s'$ */
9:        $\text{add}(x * x', ql')$  into  $s'$ 
10:    end for
11:  end for
12:   $s \leftarrow s'$ 
13: end for
14:  $vol \leftarrow 0$ 
15: for each  $(x, ql)$  in  $s$  do
16:    $vol \leftarrow vol + x * |\text{head}(ql)[n] - \text{refPoint}[n]|$ 
17: end for
18: return  $vol$ 

```

La fonction **slice** prend une liste pl de points pour les objectifs $k \dots n$ et renvoie un ensemble de listes de points issus de pl pour les objectifs $k + 1 \dots n$. La Figure 3.9 illustre cette décomposition. La liste des solutions a, b, c, d est triée par ordre décroissant selon l'axe des x (objectif f_1). On aperçoit en-haut de la figure à gauche l'hypervolume total. La construction de la première tranche $slice1$ est réalisée à partir de a jusqu'à b (sur l'axe des x). Elle est représentée en-bas à droite. La deuxième tranche part de b et de la projection de a dans le plan normal à f_1 contenant b , jusqu'à c . L'algorithme continue à partir du plan contenant c jusqu'au plan contenant d , puis du plan contenant d jusqu'au plan contenant le point de référence.

Ce mécanisme est appliqué à nouveau mais selon l'axe y (objectif f_2). Il permet de décomposer chaque ensemble de points issu de la décomposition selon f_1 .

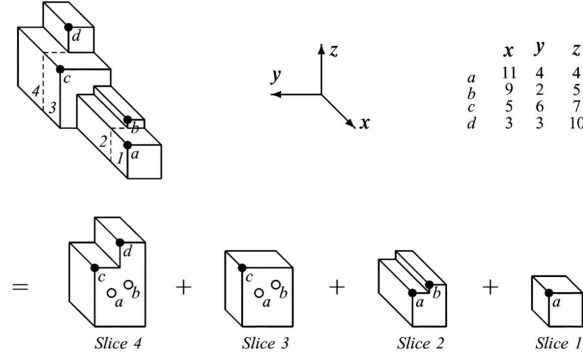


FIGURE 3.9 – Décomposition d'hypervolume de dimension 3 par HSO (extrait de [While 2006])

3.4.4 Indicateur de distance générationnelle

Cet indicateur binaire nécessite un ensemble de référence, définissant le but à atteindre. Dans le meilleur des cas cet ensemble de référence est l'ensemble Pareto optimal ce qui n'est évidemment jamais le cas lorsque l'on aborde des problèmes réels.

L'indicateur de distance générationnelle I_G proposé dans [Collette 2002] est décrit de la façon suivante :

$$I_G(A, B) = \frac{\left(\sum_{i=1}^n d_i^p \right)^{1/p}}{n} \quad (3.16)$$

avec : n le nombre d'éléments de A et d_i la distance entre la i -ème solution de A et la solution de B la plus proche. On notera que généralement on choisi $p = 2$, soit la distance euclidienne.

3.5 Synthèse

Dans ce chapitre ont été présentés les nombreux concepts permettant d'appréhender la résolution d'un problème multiobjectif (MOP). La complexité que soulève cette classe de problème impose en effet de définir une formalisation mathématique appropriée, ainsi que des méthodes de résolution efficaces.

Dans un premier temps, nous avons présenté les principales notions impliquées dans la résolution d'un MOP. En particulier la notion de Pareto optimalité apparaît

comme un élément essentiel de cette classe de problème. En effet, les objectifs à optimiser étant généralement contradictoires (l'amélioration de f_1 se fait souvent au détriment de f_2), le résultat d'un MOP est constitué d'un ensemble de solutions de meilleur compromis. L'absence d'ordre total entre les solutions fait apparaître la relation de dominance, qui a été détaillée dans cette première partie.

Après avoir défini ces notions de base, nous avons proposé un inventaire des métaheuristiques les plus couramment utilisées. Nous avons ainsi identifié quatre classes de méthode.

La première classe consiste à transformer le MOP en problème mono-objectif. Par exemple, la méthode par agrégation des fonctions objectif selon une pondération choisie est souvent utilisée pour réduire l'ensemble des fonctions objectif en une seule. Cette transformation modifie le problème initial en le simplifiant. L'apparition d'une seule fonction objectif à optimiser présente l'avantage de permettre l'utilisation des heuristiques classiques, disponibles en contexte mono-objectif. Toutefois cette approche soulève un problème essentiel, elle ne reflète pas la diversité des solutions situées dans le front de Pareto.

La deuxième classe de méthode que nous avons présentée concerne les approches non-Pareto. Il s'agit d'heuristiques considérant toutes les fonctions objectif mais séparément les unes des autres. L'algorithme génétique VEGA [Schaffer 1985] en est une illustration. Par rapport aux méthodes de la classe précédente, les approches non-Pareto permettent d'explorer une plus grande partie de l'espace de recherche. Elles ont toutefois tendance, en optimisant une fonction objectif à la fois, à pénaliser les solutions de compromis.

La troisième classe concerne les approches Pareto. Il s'agit d'heuristiques qui s'appuient sur la relation de dominance entre les solutions. Elles visent ainsi à construire le meilleur ensemble de solutions toutes non-dominées entre elles. La prise en compte simultanée des fonctions objectif permet une très bonne exploration de l'espace de recherche. Parmi ces approches de type Pareto, les algorithmes génétiques tel que NSGA-II [Deb 2002] ou SPEA-II [Zitzler 2001] tiennent une place importante. Mais les travaux les plus récents concernent des heuristiques à base d'indicateur. SMS-EMOA [Beume 2007] propose par exemple une approche génétique basée sur l'indicateur hypervolume.

La quatrième classe présentée dans ce chapitre est la classe des approches hybrides. Elles ont la particularité de combiner des méthodes issues des classes précédentes afin de bénéficier des avantages de chacune. Selon le type d'heuristiques combinées, et la façon dont elles le sont, quatre familles d'approche hybride ont été présentées. Le niveau de performances des heuristiques hybrides se révèle être en général très bon.

Enfin, la dernière partie de ce chapitre concerne l'évaluation des algorithmes dans

le domaine des MOP. En effet, le fait de disposer d'un ensemble de solutions, avec une évaluation différente sur chaque objectif, rend la comparaison des algorithmes délicate. Les approches à base d'indicateur répondent à cette problématique, les principales étant décrites dans ce chapitre. Nous avons présenté notamment l'indicateur hypervolume qui rend bien compte de la performance tant en intensité qu'en diversité d'un ensemble.

Comme nous l'avons vu, lorsque l'on aborde un MOP, les approches hybrides fournissent de très bons résultats. En particulier, pour traiter un problème combinatoire de grande taille, elles permettent souvent d'atteindre un ensemble de solutions non-dominées de bonne qualité, tant en diversité qu'en intensité, dans un temps fixé. Le chapitre suivant propose une telle approche hybride, à base de recherche locale et d'algorithme génétique, appliquée à l'optimisation de l'implantation de stations pour un service d'auto-partage électrique qui est modélisé sous la forme d'un MOP. L'algorithme hybride proposé, appelé MEMO pour Memetic Evolutionary MultiObjective, sera comparé à des approches génétiques et locales classiques.

Algorithmes mémétiques et optimisation multiobjectif

Ce chapitre est consacré à l'élaboration d'une métaheuristique multiobjectif hybride. Elle a été expérimentée sur un problème réel de localisation de stations pour un service d'auto-partage.

Dans un premier temps, une présentation des caractéristiques propres aux algorithmes mémétiques est faite. Ensuite, nous détaillons l'élaboration et le fonctionnement de l'algorithme proposé MEMO. La troisième partie de ce chapitre traite d'un cas d'étude : la localisation de stations d'auto-partage. Dans cette partie est présenté le problème, ainsi que sa modélisation sous la forme d'un problème multiobjectif. La partie suivante fournit une analyse comparative détaillée entre les approches de référence NSGA-II, PLS et FLSMO, et l'algorithme MEMO. Enfin une dernière partie aborde le problème de l'aide à la décision en contexte multiobjectif, avec la proposition de la plate-forme logicielle geoLogic.

Sommaire

4.1	Introduction	112
4.2	Schéma directeur d'un algorithme mémétique	113
4.3	MEMO : un algorithme mémétique multiobjectif	117
4.4	Optimisation appliquée à la planification urbaine : localisation de stations d'auto-partage	125
4.5	Application de l'algorithme mémétique MEMO	134
4.6	De l'optimisation à l'aide à la décision multiobjectif	158
4.7	Synthèse	169

4.1 Introduction

Après avoir présenté un état de l'art des approches permettant de résoudre des problèmes multiobjectifs (Multiobjective Optimization Problem - MOP), puis la façon de modéliser la manière dont les personnes se déplacent sur un territoire, nous nous attachons dans ce chapitre à définir une heuristique de type mémétique pour résoudre des problèmes multiobjectifs. Le cas d'étude retenu ici est la localisation de stations pour un service d'auto-partage, problème mêlant modélisation de la mobilité et optimisation.

L'équation 4.1 rappelle la formulation générale d'un problème multiobjectif déjà présentée en section 3.2.2 :

$$MOP \begin{cases} \max & z = f(x) \\ \text{où} & f(x) = (f_1(x), f_2(x), \dots, f_n(x)) \\ \text{tel que :} & x \in X \end{cases} \quad (4.1)$$

Parmi les nombreuses difficultés que posent les MOP, une des principales est liée au fait que cette classe de problèmes donne lieu à une variété de solutions. En effet, il existe rarement une solution optimale unique, qui soit la meilleure sur tous les objectifs. Le résultat est donc un ensemble de solutions, dites de meilleur compromis. Dès lors il apparaît indispensable de distinguer deux étapes dans la résolution d'un MOP : la construction de l'ensemble des solutions non-dominées (processus d'optimisation multicritère) puis le choix de la solution à retenir (processus d'aide à la décision multicritère). Dans ce chapitre nous nous focaliserons dans un premier temps sur la construction de l'ensemble Pareto [Coello 2004], ensemble qui a pour image dans l'espace des critères le front de Pareto (Pareto Front - FP). La façon d'aider le décideur à choisir la solution à implémenter fait l'objet de la fin de ce chapitre.

Les approches de type hybride tel que les algorithmes mémétiques ont déjà fait leurs preuves dans le domaine mono-objectif en permettant de résoudre de nombreux problèmes. Le concept même de mémétique fut introduit par Pablo Moscato [Moscato 1989] pour parler d'une heuristique à base de population incluant une recherche locale. Parmi les problèmes où l'approche mémétique fournit de très bon résultats, on citera l'exemple du problème académique de coloriage de graphe, connu pour être très difficile, et qui est résolu avec de très bonnes performances [Galinier 2013].

Les résultats obtenus par les approches mémétiques en mono-objectif ont incité les chercheurs à transposer ce mécanisme d'hybridation au contexte multiobjectif [Talbi 2001][Jaszkiewicz 2002]. Le mécanisme consistant à introduire une recherche locale dans un algorithme génétique (AG) est alors généralement utilisé pour accélérer le déroulement de l'algorithme, mais pas uniquement. L'hybridation joue également un rôle dans la façon d'explorer l'espace de recherche ; le but étant de fournir de

nouvelles propriétés permettant d'améliorer l'ensemble des solutions trouvées. En effet les AG utilisés seuls ont souvent tendance à faire progresser la population dans certaines régions de l'espace de recherche, ce qui a pour conséquence de créer des clusters. L'introduction de la recherche locale permet quant à elle de se focaliser sur une direction, souvent choisie vers le front de Pareto (meilleure intensification), mais parfois aussi de façon à obtenir une meilleure répartition sur le front de Pareto (meilleure diversification). La recherche locale dans un AG permet ainsi de guider la progression de la population pour garantir les propriétés souhaitées. Dès lors que sont hybridées deux métaheuristiques, recherche locale et algorithme génétique, connues pour être très performantes toutes les deux, se pose la question de l'apport de chacune d'elles. Aussi nous nous intéresserons également dans ce chapitre à identifier si, à travers l'hybridation, l'algorithme mémétique bénéficie des avantages de chaque métaheuristique, ou si de nouvelles propriétés sont engendrées.

Nous présenterons tout d'abord le principe fondateur des algorithmes mémétiques. Dans une deuxième partie, l'algorithme hybride MEMO que nous proposons sera détaillé. Ensuite nous introduirons le problème auquel nous avons appliqué l'algorithme MEMO, un problème de localisation de stations pour un service d'auto-partage de type autoLib'. La section suivante sera consacrée à une comparaison entre les résultats obtenus avec l'approche hybride et les résultats issus de métaheuristiques pures, que ce soit des approches génétiques ou des recherches locales. Enfin, une dernière partie introduira le problème de l'aide à la décision dans un contexte multiobjectif, avec une présentation de la plate-forme geoLogic.

4.2 Schéma directeur d'un algorithme mémétique

Un algorithme mémétique, dans un contexte multiobjectif, s'attache à construire le "meilleur" ensemble de solutions non-dominées. Dans le cas général la notion d'ensemble Pareto est alors utilisée pour faire référence à l'ensemble des solutions réalisables non-dominées, couvrant la totalité de l'espace de recherche. Or dans la plupart des cas, dès lors que l'on a à faire à un problème combinatoire réel, un tel ensemble composé de toutes les solutions non-dominées sur la totalité de l'espace de recherche n'est pas connu et ne peut être connu. Aussi par abus de langage nous parlerons par la suite de front de Pareto pour décrire l'ensemble des solutions non-dominées connues à un moment donné de l'exécution d'un algorithme.

Selon la littérature plusieurs acronymes sont utilisés pour parler d'algorithmes mémétiques en contexte multiobjectif. Nous retiendrons le terme de MOGLS (Multi-Objective Genetic Local Search) utilisé notamment par Talbi dans [Talbi 2013]. Le déroulement d'un MOGLS s'articule autour des deux principaux opérateurs : l'opérateur évolutionnaire, source d'exploration, et la recherche locale, garantissant l'intensification. L'algorithme 4 décrit ce fonctionnement général.

Algorithm 4 Algorithme mémétique générique MOGLS

Require: population size N , generation number $nbIter$ /* $nbIter$ can be replaced by a time limit or any termination condition*/

- 1: $P_0 \leftarrow \text{init}(N)$ /*init population P with N random individuals*/
- 2: $P'_0 \leftarrow \emptyset$ /*init an empty children population*/
- 3: $\text{eval}(P_0)$ /*eval objectives for each individual*/
- 4: **for** $i = 1$ to $nbIter$ **do**
- 5: Select P'_i from P_{i-1} based on fitness value
- 6: Apply genetic operator on $P'_i \rightarrow P''_i$
- 7: Apply local search on $P''_i \rightarrow P'''_i$
- 8: $\text{eval}(P''_i, P'''_i)$
- 9: Select P_i from $(P_{i-1} \cup P''_i \cup P'''_i)$
- 10: **end for**

La façon dont intervient la recherche locale dans l'AG peut varier selon les approches. Ainsi la recherche locale peut, comme souvent la mutation, survenir aléatoirement avec une certaine probabilité, ou encore après un certain nombre de générations. Mais il est également possible de l'appliquer à chaque nouvel individu résultant d'un croisement. L'équilibre entre recherche globale (génétique) et recherche locale est primordial pour obtenir de bons résultats.

4.2.1 La recherche locale dans un contexte multiobjectif

De nombreuses heuristiques permettent d'effectuer une recherche locale dans un contexte multiobjectif. Un aperçu des possibilités est donné dans le chapitre 3. En particulier les recherches locales de type Pareto DLMS (Dominance-Based Local Search Approaches for Multiobjective Combinatorial Optimization) [Liefoghe 2012] fournissent de très bons résultats. On retrouve ainsi parfois ces approches dans des algorithmes mémétiques. Par exemple PAES est adapté pour donner l'algorithme mémétique M-PAES [Knowles 2000]. Toutefois lorsque la recherche locale est associée à un algorithme évolutionnaire, il est possible d'envisager une recherche locale plus simple, tant en termes d'implémentation que d'exécution. En effet une recherche locale simple peut être rendue très performante dès lors qu'on y ajoute des mécanismes génétiques. Ainsi la version la plus simple consiste à considérer successivement chaque objectif, sans tenir compte des autres objectifs. Mais il est possible de complexifier les traitements en prenant en compte l'ensemble de la population pour orienter la recherche locale vers le front de Pareto ou pour favoriser certaines propriétés de la population. En particulier dans le cas d'une optimisation avec a priori sur le type de solutions souhaitées, le décideur peut choisir de focaliser la recherche dans un sous-espace de l'espace de recherche complet.

La recherche locale est susceptible de fournir un certain nombre de nouvelles

solutions non-dominées par la population. Se pose alors la question de savoir ce qui doit être fait de ces solutions non-dominées. Elles sont d'une part archivées pour être retournées par l'algorithme à la fin de l'exécution, mais qu'en est-il de leur introduction dans la population courante ? Dans certains travaux elles sont toutes intégrées à la population. Mais cela risque très vite de poser des problèmes de diversité de la population. Aussi d'autres approches introduisent des critères complémentaires, tels que l'ajout de la dernière solution trouvée uniquement. Ou encore des méthodes imposent une distance minimale avec les solutions présentes dans la population pour accepter l'ajout d'une nouvelle solution.

Enfin la durée de la recherche locale peut être un élément variant d'une méthode à l'autre. Ainsi certains algorithmes ont recours à quelques itérations seulement de recherche locale. Dans d'autres cas, la recherche locale se fait de façon complète, jusqu'à l'optimum local, ou encore jusqu'à ce que la progression ne soit plus suffisamment significative.

4.2.2 L'opérateur génétique dans un contexte multiobjectif

Les nombreuses approches génétiques définies dans un contexte multiobjectif (Multi-Objective Evolutionary Algorithm - MOEA) peuvent être hybridées avec une recherche locale. Ainsi les métaheuristiques souvent utilisées seules, telles que SPEA [Zitzler 2001] ou encore NSGA-II [Deb 2002], se retrouvent également fréquemment utilisées dans des algorithmes mémétiques. Dans leur livre [Coello 2007] Coello et al. présentent de façon détaillée un certain nombre d'approches et la façon dont les MOEA sont utilisés.

D'une façon générale le MOEA fournit un cadre permettant à la recherche locale de s'appliquer dans un contexte de population. En particulier on retrouve dans tous les cas les opérateurs de sélection, de croisement et de remplacement. La mutation est sans doute l'opérateur le plus généralement remis en cause, puisqu'il est souvent remplacé par la recherche locale. On retrouve donc un principe très semblable à ce qui se fait en mono-objectif, avec des opérateurs adaptés au contexte multiobjectif.

4.2.3 L'algorithme mémétique, une recherche locale améliorée ?

Après avoir vu ce que peuvent être la recherche locale et l'opérateur génétique dans le contexte mémétique, il apparaît intéressant de s'interroger sur la façon dont agit chaque opérateur sur le fonctionnement global de l'algorithme. En effet, l'hybridation de deux métaheuristiques connues pour être très performantes permet-elle de bénéficier des avantages de chacune ? Ou au contraire est-ce que de nouveaux comportements apparaissent ?

TABLE 4.1 – Rôles des opérateurs selon le type d’algorithme

	Recherche Locale			Algorithme évolutionnaire			
	Descente	Tabou	Recuit	Sélec.	Crois.	Mut.	Rempl.
LS	int	div	div	-	-	-	-
AG	-	-	-	int	int	div	int
Mémétique	int	div	div	int	div	int	int

Pour répondre à ces questions revenons au fonctionnement de chaque métaheuristique prise séparément afin d’étudier leur fonctionnement. En particulier le point central de tout algorithme d’optimisation, que ce soit en mono-objectif ou en multiobjectif, est la façon dont s’équilibrent l’intensification et la diversification, concepts également connus sous les termes d’exploitation et d’exploration respectivement.

Nous tentons dans le tableau 4.1 une synthèse de l’influence des différentes composantes des AG et recherches locales, en matière de diversification et intensification. Chaque ligne représente une heuristique (Locale, Génétique, Mémétique), et chaque colonne un opérateur (de la famille des recherches locales à gauche et des AG à droite). Les cases contiennent le rôle dominant correspondant à l’opérateur de la colonne pour l’algorithme de la ligne, en matière d’intensification (int) et de diversification (div).

Ainsi dans la recherche locale la descente joue un rôle d’intensification, alors que les paramètres de la Recherche Tabou ou du Recuit Simulé jouent un rôle de diversification.

De même dans le cas d’un AG, la sélection, le croisement et le remplacement tendent à intensifier la population et la mutation à la diversifier. Bien entendu ce tableau n’est pas exhaustif mais traduit une tendance généralement constatée. En particulier il est possible de nuancer l’influence des opérateurs, en introduisant une certaine dose d’aléas dans la sélection par exemple.

L’élément particulièrement intéressant à remarquer est l’inversion des rôles entre croisement et mutation selon que l’on ait à faire à un AG ou à un algorithme mémétique. En effet, alors que le croisement est l’opérateur d’intensification clé pour un AG (c’est généralement suite à un croisement que de nouvelles bonnes solutions émergent), il a un rôle de diversification dans le contexte mémétique. Inversement la mutation qui introduit souvent de la diversité en génétique, lorsqu’elle est remplacée par une descente en mémétique, elle devient un opérateur d’intensification.

Ainsi nous sommes tentés de considérer qu’un algorithme mémétique n’est ni plus ni moins qu’une recherche locale à base de population, dans laquelle le croisement joue un rôle d’opérateur de diversification. Ce n’est pas le croisement qui permet de trouver de meilleures solutions mais la recherche locale. Dans un AG en général

ce n'est pas la mutation qui permet de trouver de meilleures solutions mais le croisement.

4.3 MEMO : un algorithme mémétique multiobjectif

Dans cette section nous proposons un nouvel algorithme hybride, l'algorithme Memetic Evolutionary MultiObjective (MEMO). Le fait même que l'algorithme soit mémétique fournit une trame générale bien identifiée, comme nous l'avons vu, mais il reste un certain nombre de caractéristiques à définir :

- choix de la recherche locale
- choix de l'opérateur génétique
 - ⇒ quelle sélection des parents ?
 - ⇒ quel croisement ?
 - ⇒ quelle gestion de population ? Quelle solution remplacer lors de la mise à jour de population ? Quelles solutions archiver ?
- choix de l'enchaînement de ces deux opérateurs : effectuer une recherche locale après chaque croisement, après chaque génération, ou de façon plus espacée (une fois de temps en temps)
- génération de la population initiale : aléatoirement, à l'aide d'un algorithme glouton, etc.
- définition d'un critère d'arrêt : nombre d'itérations, temps de calcul, etc.

Par ailleurs les résultats obtenus dépendront également fortement d'une série de paramètres, tels que la taille de la population, et qui seront à définir.

Rappelons que l'élément clé qui doit guider le choix des opérateurs à l'intérieur du schéma général de la recherche repose sur l'équilibre entre :

- une exploration de l'espace de recherche grâce à l'AG
- une intensification de la population obtenue à travers la stratégie locale

4.3.1 NSGA-II comme opérateur génétique

L'AG choisi pour servir de base à l'algorithme MEMO est NSGA-II [Deb 2002] proposé par Deb et al. Il s'agit certainement de l'AG le plus utilisé pour résoudre des problèmes multi-objectifs, bien qu'il existe des algorithmes plus efficaces. NSGA-II présente l'avantage de proposer un très bon compromis entre performance et simplicité de fonctionnement. Il s'agit ainsi encore aujourd'hui d'un algorithme souvent utilisé comme référence, pour comparer les performances des nouvelles approches proposées.

Notons que NSGA-II servira non seulement à construire l'opérateur génétique de MEMO, mais également seul pour fournir des éléments de comparaison afin d'évaluer les apports de l'hybridation que nous proposons.

L'algorithme 5 rappelle le fonctionnement de NSGA-II présenté section 3.3.4.1 du chapitre 3.

Algorithm 5 NSGA-II

Require: population size N , generation number $nbIter$ */*nbIter can be replaced by a time limit*/*

- 1: $P \leftarrow \text{init}(N)$ */*init population P with N random individuals*/*
- 2: $Q \leftarrow \emptyset$ */*init an empty children population*/*
- 3: $\text{eval}(P)$ */*eval objectives for each individual*/*
- 4: **for** $i = 1$ to $nbIter$ **do**
- 5: $P \leftarrow P \cup Q$
- 6: $\text{assignRank}(P)$ */*based on Pareto dominance*/*
- 7: **for** each non-dominated front $f \in P$ **do**
- 8: $\text{setCrowdingDist}(f)$
- 9: **end for**
- 10: $\text{sort}(P)$ */*by rank and in each rank by the crowding distance*/*
- 11: $P \leftarrow P[0 : N]$
- 12: $Q \leftarrow \text{buildNextGeneration}(P)$ */*Binary Tournament Selection, Recombination and Mutation*/*
- 13: **end for**

Définissons maintenant les différents opérateurs intervenant dans NSGA-II.

4.3.1.1 La sélection des parents

Elle fait partie des éléments définis dans l'algorithme NSGA-II. Il s'agit d'un tournoi binaire permettant de favoriser les meilleures solutions sans pour autant éliminer totalement les autres solutions. Le tournoi est basé sur le rang des solutions dont le calcul est détaillé dans le chapitre 3. NSGA-II implique le recours à un tournoi binaire, mettant en compétition deux solutions uniquement. On notera toutefois que plus généralement il est possible de mettre plus de deux solutions en compétition. La Figure 4.1 montre la sélection de la solution s_2 suite à la mise en tournoi de s_2 et s_6 . Dans cette illustration plus le cercle est grand, meilleure est la solution. Dans le cas d'un tournoi, les individus mis en compétition sont quant à eux choisis aléatoirement. On voit ainsi que s_2 est choisie sans être la meilleure, simplement parce qu'elle est meilleure que s_6 .

La sélection peut ainsi être vue comme un opérateur d'intensification, dans la

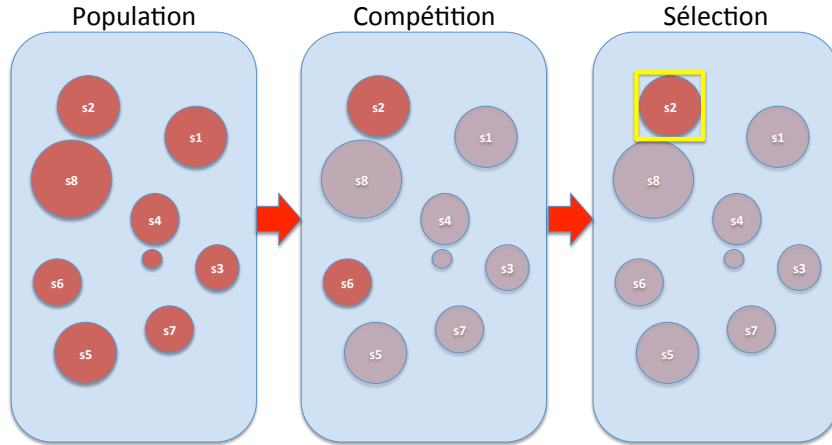


FIGURE 4.1 – Sélection par tournoi binaire

mesure où elle favorise les meilleures solutions. Il apparaît de façon immédiate que plus on met de solutions en compétition, plus la sélection sera élitiste. Le cas extrême consiste à mettre en compétition toute la population, ce qui reviendrait à choisir la meilleure solution (élitisme extrême).

4.3.1.2 Le croisement

Contrairement à la sélection des parents, l'opérateur de croisement n'est pas défini dans NSGA-II. Ceci s'explique par le fait que le croisement peut être étroitement lié au problème et dépend donc du cas de figure.

Sans connaître le problème à résoudre il apparaît ainsi délicat de préconiser un opérateur de croisement plutôt qu'un autre. C'est pourquoi nous proposons ici d'envisager deux opérateurs de croisement, avec deux niveaux d'élitisme. Le choix d'un opérateur ou de l'autre se fera au regard des résultats obtenus dans les deux cas.

Un croisement élitiste Nous choisissons pour ce croisement de prendre les meilleurs gènes de chaque parent. Toutefois dès lors que l'on parle de meilleur gène il faut être en mesure d'attribuer une valeur à chaque gène, ce qui soulève deux problèmes. D'une part cela suppose que l'on puisse identifier la contribution de chaque gène à l'évaluation globale de la solution, ce qui n'est pas forcément le cas. D'autre part nous sommes dans un contexte multiobjectif. Il n'existe donc pas de relation d'ordre total entre les solutions et par conséquent entre les gènes contributeurs à la solution, chaque gène possédant une contribution à chacun des objectifs.

Il apparaît alors nécessaire de définir la façon dont on va comparer les gènes

les uns aux autres, ce qui suppose de rétablir un ordre total. Pour ce faire nous appliquerons une méthode d'agrégation des objectifs. Le croisement se fera ainsi dans une direction $\omega = (\omega_1, \dots, \omega_k)$, avec k le nombre de fonctions objectifs, tel que $\omega_i \geq 0$ et $\sum_{i=1}^k \omega_i = 1$. La pondération des objectifs n'est pas figée. Elle est tirée aléatoirement lors de la création de chaque enfant. On notera qu'ici aucune direction n'est privilégiée. On ne cherchera donc pas à guider les solutions issues du croisement vers le front de Pareto ou vers une meilleure répartition des solutions le long du front. Au contraire, on laissera l'aléa jouer son rôle et apporter une plus grande diversité.

Le recours à cette agrégation f des objectifs sous forme d'une somme pondérée des f_i présente l'avantage de fournir temporairement une valeur de fitness unique à chaque solution, et dans notre cas une valeur de contribution unique de chaque gène à la solution globale.

$$f = \sum_{i=1}^k \omega_i * f_i \quad (4.2)$$

De plus la recherche locale présentée section 4.3.2 utilisera également ce procédé d'agrégation, ce qui permettra d'utiliser les mécanismes simples disponibles dans le contexte mono-objectif.

La figure 4.2 illustre la façon dont chaque parent fournit ses meilleurs gènes à l'enfant, ici chacun pour moitié. Selon la façon dont est codé le problème, le découpage se fera différemment. Un exemple sera présenté dans la section 4.4.3.3.

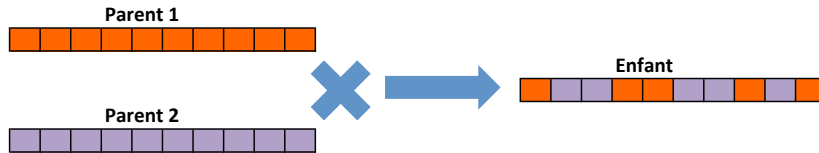


FIGURE 4.2 – Croisement élitiste : chaque parent fournit ses meilleurs gènes

Malgré le caractère élitiste du croisement proposé, on notera que cette opération dans le contexte mémétique peut être vue comme un opérateur de diversification. En effet le croisement a pour effet de déstructurer les solutions pour fournir un nouveau point de départ pour la recherche locale.

Un croisement aléatoire Du fait qu'il n'est pas toujours possible de définir un croisement élitiste, et parce que l'élitisme peut impliquer des convergences prématurées de l'algorithme, nous gardons comme alternative également le croisement aléatoire. Cela peut être un croisement à un point classique, si la séquence des gènes a une importance dans le codage de la solution. C'est le cas par exemple pour des

problèmes d'ordonnancement où chaque gène code une tâche à exécuter et où la solution décrit la succession de l'ensemble des tâches.

Mais dans d'autres cas, il peut être au contraire nécessaire de casser la succession des gènes. En particulier dans les codages binaires des solutions où chaque gène représente toujours la même chose, avec l'information "on prend" ou "on ne prend pas", l'aléa comme nous le verrons plus loin peut être obtenu de manière adaptée.

4.3.1.3 La gestion de population

NSGA-II définit la façon dont la population évolue, en conservant les meilleures solutions selon le rang de dominance d'abord, et selon la distance de crowding ensuite. Toutefois dans le cadre de l'algorithme Memetic Evolutionary MultiObjective (MEMO), nous considérons deux populations. La population *PG* de taille fixe qui est celle utilisée par l'algorithme génétique. Et une deuxième population *PA* de taille variable qui sert d'archive. Afin de limiter la taille de l'archive de nombreux mécanismes sont souvent utilisés, pour supprimer notamment les individus trop proches les uns des autres. Les fonctions de partage (sharing) permettent ainsi de conserver une archive de taille raisonnable avec une bonne diversité d'individus. Dans le cas présent nous n'imposons pas de taille limite pour archiver les solutions rencontrées. Nous acceptons ainsi toute nouvelle solution qui n'est dominée par aucune solution déjà présente dans l'archive.

A la fin de la recherche, l'algorithme retourne un ensemble contenant toutes les solutions non-dominées rencontrées au cours de la recherche. Le fait de ne pas borner cet ensemble risque d'engendrer un nombre considérable de solutions, éventuellement souvent proches les unes des autres. Nous présenterons dans le chapitre 4.6 une façon d'aider le décideur à choisir la solution qui l'intéresse parmi toutes les possibles.

4.3.1.4 Et la mutation...

Comme nous l'avons indiqué certains algorithmes mémétiques conservent en plus de la recherche locale un opérateur de mutation. Toutefois dans de nombreux cas de figure la recherche locale vient en remplacement de la mutation. Ici la mutation généralement présente dans l'algorithme NSGA-II est simplement remplacée par la recherche locale.

4.3.2 L'opérateur de recherche locale

Afin de répondre aux attentes présentées précédemment, la recherche locale doit fournir un équilibre entre rapidité et exhaustivité. Nous proposons ici une approche

de type First Improvement Hill Climbing (FIHC) qui accepte successivement le premier voisin de meilleure fitness, à partir d'un point de départ tiré aléatoirement. De cette façon une exploration partielle du voisinage est effectuée à chaque itération. La condition d'arrêt de la recherche est constituée de l'optimum local, solution n'admettant aucune solution de meilleure fitness dans son voisinage.

Plus formellement la recherche locale est définie par le couple (Ω, V) où Ω est l'ensemble des solutions réalisables (il s'agit de l'espace de recherche) et V la structure de voisinage $V : \Omega \rightarrow 2^\Omega$ qui attribue à chaque solution $s \in \Omega$ un ensemble de voisins $V(s)$. De fait, la condition d'arrêt de la recherche locale est l'optimum local s^* , atteint si et seulement si $\forall s \in V(s^*), f(s) \leq f(s^*)$ (dans le cas d'un problème de maximisation).

On notera que dans le cadre de la recherche locale le problème est ramené à un problème mono-objectif, la valeur de fitness d'une solution étant la somme pondérée des différents critères dans la direction choisie aléatoirement et définie par le vecteur de poids, tel que décrit section 4.3.1. L'algorithme 6 présente la façon dont la recherche est réalisée.

Algorithm 6 Local Search : FIHC

Require: s the child coming from crossover operator

- 1: $f_s \leftarrow \sum_{i=1}^k \omega_i * f_i(s)$
 - 2: **repeat**
 - 3: $s' \leftarrow \text{selectNeighbor}(s)$ */*select randomly the first neighbor of s such that $f'_s > f_s^*$ */*
 - 4: **if** $s' \neq \emptyset$ **then**
 - 5: $s \leftarrow s'$
 - 6: $f_s \leftarrow \sum_{i=1}^k \omega_i * f_i(s)$
 - 7: $\text{addNotDominated}(s)$ */*add s in the archive if not dominated inside it and remove all dominated solutions from the archive*/*
 - 8: **end if**
 - 9: **until** s is a local optimum (i.e. $s' = \emptyset$)
-

Chaque solution rencontrée lors de la recherche locale qui n'est pas dominée par l'archive est ajoutée à cette dernière. L'archive est ainsi construite et alimentée par la recherche locale.

La solution s à la fin de la recherche locale, qui correspond à l'optimum local, est ajoutée à la population utilisée pour la partie génétique. Elle sert donc à construire les générations suivantes.

On notera qu'il aurait également été possible de choisir n'importe quelle autre solution non-dominée rencontrée au cours de la recherche locale. En effet, en termes de dominance la dernière solution n'est a priori pas meilleure que n'importe quelle

autre solution non-dominée. Il aurait alors été possible de définir d'autres critères tel que celui de diversité en appliquant une fonction de sharing pour choisir la solution à intégrer à la population pour l'AG. Ici nous avons choisi de donner la priorité à la direction choisie qui n'admet qu'une meilleure solution : la dernière trouvée.

La recherche locale choisie ici repose donc sur l'agrégation des fonctions objectif. Elle présente l'avantage d'être simple et rapide. Il aurait également été envisageable de choisir une recherche locale de type Pareto, telle qu'un algorithme DMLS [Liefvooghe 2012]. Ceci aurait permis de conserver le caractère multi-objectif du problème. Toutefois ce type d'approches, dans le cadre d'une hybridation, pose la question de l'archive servant de référence. En effet, si la recherche locale appliquée à chaque enfant part d'une archive vide, la recherche locale basée sur le critère de dominance prendrait un temps considérable, bien souvent pour ré-explore des parties de l'espace de recherche déjà explorées précédemment. A l'inverse, si l'archive initiale de chaque recherche locale est constituée de l'archive globale de l'algorithme mémétique, le risque est que les enfants soient directement dominés par une solution de l'archive ce qui arrêterait la recherche avant même qu'elle n'ait commencé.

4.3.3 Déroulement de l'algorithme mémétique

La figure 4.3 illustre les trois étapes principales : la sélection des parents p_1 et p_2 , leur croisement pour donner naissance à la solution s , suivi d'une recherche locale dans la direction donnée par le vecteur $\vec{v}$ pour aboutir à la solution s^* . On notera s_i les solutions successives de la recherche locale. La direction de $\vec{v}$ est directement déterminée par le vecteur des pondérations $\omega = (\omega_1, \dots, \omega_k)$.

Remarque : Le vecteur directeur de la progression de la recherche locale autorise toutefois un certain degré de liberté dans l'espace de recherche. En effet le vecteur $\vec{v}$ et la solution courante s_i définissent un hyperplan orthogonal à $\vec{v}$ et passant par s_i . Lors de la recherche locale toutes les solutions s_{i+1} se trouvant au-dessus de cet hyper plan améliorent la fitness agrégée et sont donc acceptées.

L'algorithme 7 présente ces étapes.

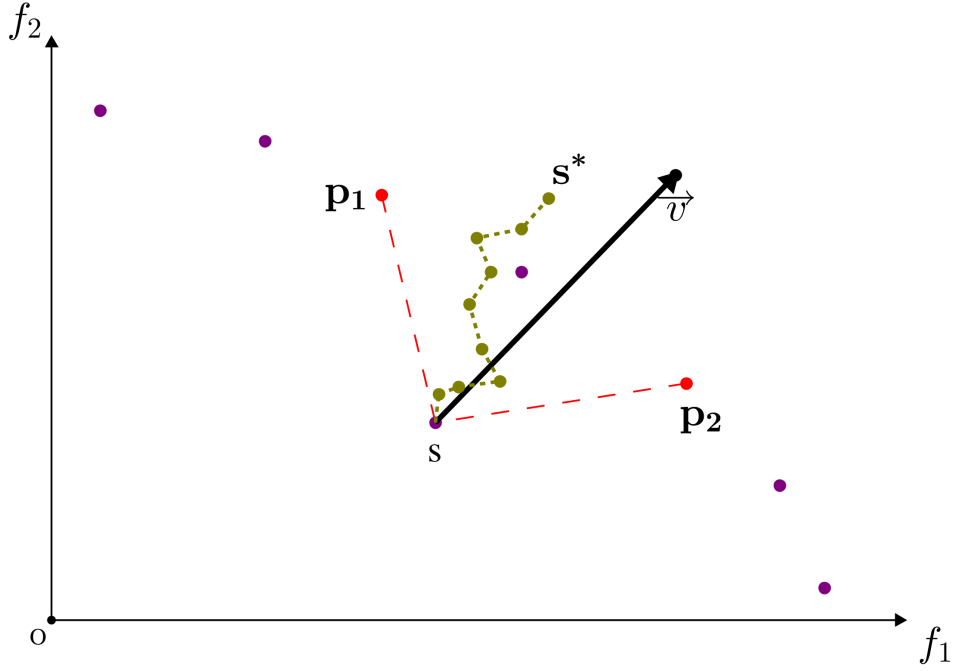


FIGURE 4.3 – Sélection et croisement suivis d'une recherche locale

Algorithm 7 MEMO

Require: population size N , generation number $nbIter$ /* $nbIter$ can be replaced by a time limit*/

- 1: $P \leftarrow \text{init}(N)$ /*init population P with N random individuals*/
- 2: $Q \leftarrow \emptyset$ /*init an empty children population*/
- 3: $A \leftarrow \emptyset$ /*init an empty archive*/
- 4: $\text{eval}(P)$ /*eval objectives for each individual*/
- 5: **for** $i = 1$ to $nbIter$ **do**
- 6: $P \leftarrow P \cup Q$
- 7: $\text{assignRank}(P)$ /*based on Pareto dominance*/
- 8: **for** each non-dominated front $f \in P$ **do**
- 9: $\text{setCrowdingDist}(f)$
- 10: **end for**
- 11: $\text{sort}(P)$ /*by rank and in each rank by the crowding distance*/
- 12: $P \leftarrow P[0 : N]$
- 13: $Q \leftarrow \text{buildChildren}(P)$
- 14: $\text{improveChildren}(Q)$ /*following the local search FIHC*/
- 15: **end for**

La principale caractéristique de cet algorithme par rapport à NSGA-II est l'introduction

de la fonction *improveChildren*. Elle a en effet la charge d'appliquer la recherche locale à l'ensemble des individus de la population des enfants Q . De fait dans cet algorithme la fonction *buildChildren* effectue la sélection et le croisement mais pas la mutation qui est remplacée par la recherche locale.

L'approche algorithmique étant présentée, nous allons passer à un cas d'application avec le problème de l'auto-partage.

4.4 Optimisation appliquée à la planification urbaine : localisation de stations d'auto-partage

4.4.1 Introduction

Le service de l'auto-partage a été introduit en 1940 [Shaheen 2008]. Depuis, plusieurs approches ont été proposées dont les trois principales sont :

- le service à station unique : les voitures sont mises à disposition dans une station, et le retour se fait dans la même station.
- le service sans station : les véhicules sont déployés dans la ville sur des places de stationnement classique. Un véhicule est alors emprunté où il se trouvait et déposé n'importe où dans la zone d'utilisation.
- le service à stations multiples : les voitures sont réparties dans des stations, sont empruntées dans une station et restituées dans n'importe quelle autre station du service.

Le cadre d'application qui nous intéresse est cette troisième version. Il s'agit notamment du mode de fonctionnement adopté par autolib' à Paris, ainsi que pour le partage de vélo avec vélib' à Paris ou véloV' à Lyon. Un des enjeux majeur est alors de positionner les stations afin de fournir le meilleur service au moindre coût.

Un déséquilibre important du service est certainement l'élément principal pouvant poser problème. On distingue les deux situations extrêmes suivantes :

- le cas où une station est vide. Une demande de retrait de véhicule à cette station sera alors rejetée.
- le cas où une station est pleine. Une demande de dépôt sera également rejetée.

Les demandes non satisfaites posent deux problèmes : d'une part elles constituent un manque à gagner direct pour l'exploitant, d'autre part elles entraînent une perte de confiance dans le service, ce qui est à terme préjudiciable.

Divers travaux ont donné lieu à plusieurs approches pour optimiser le fonctionnement de l'auto-partage et réduire les situations de blocage. Notamment dans [Barth 2002] l'accent est mis sur la participation de l'utilisateur lui-même pour redéployer les véhicules,

en l'incitant à aller vers certaines stations. Mais l'efficacité de l'auto-partage passe certainement par une conception optimale en amont. Ainsi des travaux ont déjà été réalisés pour répondre à cette problématique de localisation des stations en appliquant des méthodes d'optimisation exactes. [de Almeida Correia 2012] propose de placer les stations en résolvant un Problème Linéaire Mixte. Nous notons toutefois que le principal inconvénient de ces approches est qu'elles considèrent une version très simplifiée du problème. En particulier le découpage du territoire se fait en grande zone de 1km^2 . En effet, la complexité du problème n'autorise pas de conserver des approches exactes dès lors que l'on souhaite analyser finement le territoire. Le recours à une heuristique s'avère alors être nécessaire.

4.4.2 Enjeu de la localisation des stations

Avant toute formalisation d'un problème d'optimisation, il apparaît nécessaire de comprendre le problème lui-même afin d'identifier les éléments clés qui définissent les fonctions objectifs. En particulier dans le cas présent, cela revient à caractériser les éléments ayant l'impact le plus significatif sur le fonctionnement du service d'auto-partage.

Qu'est-ce qu'un service d'auto-partage qui fonctionne bien ? Afin de rendre compte des divers facteurs traduisant le "bon" fonctionnement du service, nous identifions que les principaux éléments à optimiser simultanément doivent permettre de :

- répondre à un maximum de besoins en termes de déplacements
- nécessiter la plus faible intervention humaine pour redéployer les véhicules en cas de saturation du service
- assurer une utilisation du service tout au long de la journée et non uniquement aux heures de pointe

De plus le niveau de précision géographique semble être primordial. En effet indiquer qu'une station serait efficace dans une certaine zone à 1km près ne contient pas le même degré d'information que d'indiquer une position à 100m près. En corollaire on imagine bien qu'appliquer un algorithme d'optimisation sur un maillage à 1km ou à 100m n'implique pas la même complexité de l'espace de recherche lui-même. On voit alors que le choix du découpage territorial est un enjeu non négligeable.

4.4.3 Description formelle du problème

4.4.3.1 Données d'entrée : description et modélisation

L'implantation de stations dans un environnement urbain suppose la prise en compte d'un grand nombre de paramètres liés tant à la disponibilité du foncier qu'à une connaissance fine du besoin en termes de mobilité.

Caractérisation de la disponibilité foncière La disponibilité du foncier peut être vue comme un coût associé à chaque emplacement pouvant recevoir techniquement une station. L'introduction de ce coût présente comme avantage d'apporter un degré de finesse important dans la caractérisation de chaque emplacement. Dans cette optique un emplacement ne pouvant héberger une station se verrait associer un coût infini. Toutefois l'inconvénient majeur de cette approche est qu'elle impose de calculer une fonction coût généralisée intégrant le prix du foncier, le gain financier apporté par l'emplacement et le gain en matière de service rendu.

A l'opposé la disponibilité du foncier peut être vue comme une fonction binaire indiquant les emplacements disponibles pour recevoir une station. Bien que dans ce cas il n'y ait plus la finesse apportée par l'association d'un coût à chaque emplacement, ce qui a pour effet de ramener tous les emplacements au même niveau, une telle approche présente l'avantage d'être beaucoup plus souple à mettre en oeuvre. En particulier la localisation des stations dans ce cas se fait en considérant la demande (le besoin en mobilité) sans lui associer un gain monétisé qui est souvent une question d'appréciation.

C'est cette dernière caractérisation du foncier que nous avons pris en compte dans la modélisation du problème afin de limiter les erreurs susceptibles d'être introduites par une approche monétisée. Toutefois, afin d'intégrer le fait que des emplacements puissent être plus contraignants que d'autres, puissent imposer des travaux plus lourds et plus coûteux que d'autres, nous proposons de laisser la possibilité au décideur, lorsqu'une solution lui est proposée, de rejeter un emplacement choisi. Le rejet a posteriori permet au décideur de voir toutes les solutions proposées par l'outil d'optimisation, d'analyser le gain apporté par chaque emplacement d'une solution, et éventuellement de rejeter un emplacement et de relancer l'algorithme pour trouver un emplacement de substitution.

La figure 4.4 montre un découpage du territoire en mailles de 300m de côté, avec en grisé les mailles où il n'est pas possible de disposer une station. Dans les mailles non grisées il n'est pas forcément possible de mettre une station n'importe où mais il existe des possibilités. Le choix de 300m est guidé par les méthodes de déploiement de transports en commun dont la zone de capture des piétons autour d'un arrêt est en général d'un rayon de 300m.



FIGURE 4.4 – Illustration des sites pouvant accueillir une station

Identification de la demande Après l'identification des sites pouvant recevoir un aménagement tel qu'une station pour véhicules électriques partagés, la deuxième donnée d'entrée nécessaire à la localisation des stations est une connaissance précise du besoin. Cela implique d'être en mesure de :

- caractériser et modéliser l'ensemble des flux de déplacements tels qu'ils s'opèrent sur le territoire
- filtrer parmi tous les flux ceux qui sont susceptibles d'être effectués en véhicule partagé

Modéliser l'ensemble des flux nécessite de disposer d'un maximum de données d'entrée pour se trouver au plus près de la réalité. Les données peuvent provenir des administrations territoriales elles-mêmes et sont constituées de comptages aux carrefours, aux gares, d'Enquêtes Ménages/Déplacements (EMD), d'enquêtes diverses, etc. Elles sont également construites au niveau national, notamment par l'INSEE à travers le recensement, ou encore des grandes enquêtes telles que l'Enquête Nationale Transports et Déplacements (ENTD), la dernière remontant à 2008.

Les données issues de la téléphonie mobile permettent de compléter celles déjà disponibles en fournissant un échantillon de taille très importante et réactualisable en continu.

L'étape consistant à filtrer les utilisateurs potentiels présente de nombreuses difficultés. Il ne s'agit plus de modéliser au plus juste une situation connue mais d'anticiper sur ce que va être l'utilisation des véhicules en auto-partage. Pour ce faire il faut être en mesure d'estimer qui utilisera le service en fonction de critères tels que le type de personne (âge, sexe, catégorie socio-professionnelle, etc.) ainsi que le motif du déplacement.

A ceci s'ajoutent les déplacements dits "induits" par l'introduction de ce nouveau service. En effet toutes les expériences passées (Praxitèle à La Rochelle [M.-H. 2000], etc.) ont montré que l'introduction d'un nouveau mode de déplacement entraînait non seulement un report modal des autres modes vers ce dernier, ceci pour des déplacements qui étaient déjà effectués, mais en plus que ce nouveau mode générerait de nouveaux déplacements.

Afin de ne pas biaiser les résultats en introduisant de mauvaises hypothèses, nous avons laissé la possibilité à l'utilisateur d'introduire ses propres hypothèses sur les critères de filtres, caractérisés par une pondération des déplacements actuels : par tranche d'âge, sexe et moyen de transport avant introduction du nouveau mode. Ceci permet de considérer différents scénarios et d'étudier les répercussions sur les solutions fournies par l'algorithme de localisation.

Cette étape laissée aux mains de l'utilisateur aurait également pu être intégrée à l'algorithme d'optimisation. Une approche robuste ou encore stochastique pourrait être envisagée dans ce cas de figure mais ne sera pas considérée ici.

Ces deux étapes de modélisation des déplacements puis filtrage des utilisateurs pressentis aboutissent à la création de la demande.

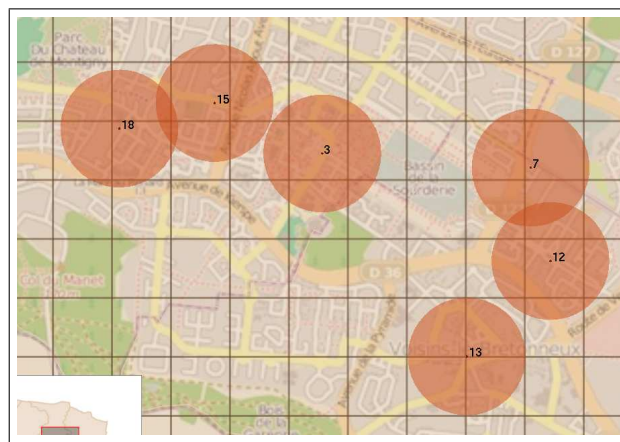


FIGURE 4.5 – Stations et zone de capture

La figure 4.5 présente la localisation des stations (centre du cercle), ainsi que la zone de capture (aire du disque). Nous considérons donc ici une aire de capture d'un rayon de 300m, pouvant toutefois être paramétrée par l'utilisateur. Le flux des déplacements pouvant être servi par une station sera ainsi l'ensemble des déplacements ayant pour origine ou destination la zone de capture associée. On notera que les flux peuvent être pris partiellement par une station (maille non entièrement couverte) et/ou répartis entre plusieurs stations (zones d'intersection entre plusieurs aires).

4.4.3.2 Modélisation des fonctions objectifs

Afin de répondre aux enjeux posés précédemment, nous proposons trois objectifs à optimiser. Ces fonctions ont pour rôle d'identifier parmi les sites candidats lesquels retenir afin d'y placer une station.

Définition 18 (Site candidat) *On appelle site candidat une maille carrée pouvant*

recevoir une station. Cette maille est elle-même caractérisée par une position (x, y) et un pas de maillage (longueur de côté).

Définissons les variables intervenant dans les fonctions objectifs.

Ω : l'ensemble des solutions réalisables

s : une solution élément de l'ensemble Ω . Notons qu'une solution réalisable est une façon de placer n stations sur le territoire (choix de n sites candidats), sachant qu'un site candidat ne peut accueillir qu'une station.

st_i : la station i du service. Son domaine de définition est l'ensemble des sites candidats du territoire.

T : l'ensemble des plages horaires d'une journée pour prendre ou déposer un véhicule, $T = \{t_1, \dots, t_p\}$.

t : une plage horaire, élément de T . On prendra généralement des plages de 15 minutes, soit $p = 96$.

$f(st_i, st_j)$: nombre de personnes allant de st_i à st_j au cours de la journée

$f(st_i, st_j, t)$: nombre de personnes allant de st_i à st_j au cours de la période t

$f(st_i, t)$: nombre de déplacements ayant pour origine ou destination st_i au cours de la période t

$\bar{f}(st_i)$: nombre moyen de déplacements ayant pour origine ou destination st_i au cours de la journée

$f_{min}(st_i) = \sum_t \min \left(\sum_{st_j \in s \setminus \{st_i\}} f(st_i, st_j, t), \sum_{st_j \in s \setminus \{st_i\}} f(st_j, st_i, t) \right)$ est le cumul du minimum entre flux entrant et flux sortant tout au long de la journée pour st_i .

$f_{max}(st_i) = \sum_t \max \left(\sum_{st_j \in s \setminus \{st_i\}} f(st_i, st_j, t), \sum_{st_j \in s \setminus \{st_i\}} f(st_j, st_i, t) \right)$ est le cumul du maximum entre flux entrant et flux sortant tout au long de la journée pour st_i . La fonction correspond au flux demandé total.

Maximiser les flux de déplacement Il s'agit de l'objectif qui semble être le plus intuitif. En effet les stations doivent être positionnées de façon à être à l'origine et à la destination d'un maximum de besoins en termes de déplacement. Il ne s'agit donc pas simplement de mettre les stations dans les zones attractives comportant le plus de personnes mais de façon à satisfaire un maximum d'échanges origine/destination. La Figure 4.6 montre que mettre simplement les stations dans les zones les plus attractives ne suffit pas forcément. On y voit en particulier des zones (A à F) avec les échanges existants entre chaque zone. L'épaisseur de la flèche traduit le nombre de déplacements existants. Les zones rouges sont les zones sélectionnées pour mettre une station.

Dans cet exemple il s'agit de placer deux stations dans les zones les plus intéressantes. Sur la figure de gauche on place les stations sur les zones les plus attractives (A et

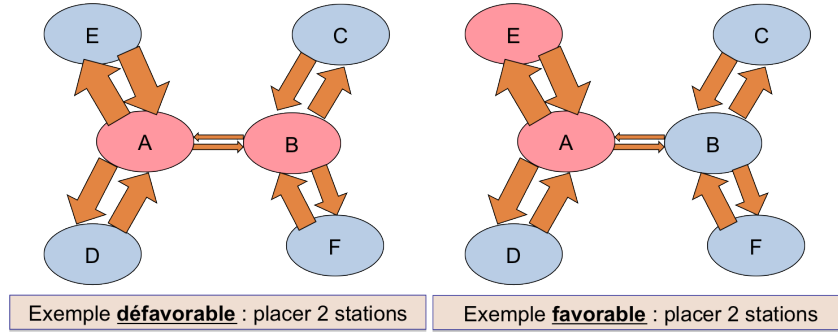


FIGURE 4.6 – Implantation des stations selon l'attraction vs les flux

B). On remarque toutefois que bien que l'attractivité soit forte, les échanges entre *A* et *B* sont très faibles ce qui est défavorable au fonctionnement de l'auto-partage. Sur la figure de droite au contraire les stations sont positionnées de façon à maximiser les flux (*A* et *E* sont choisies). Cela montre un fonctionnement favorable.

Cela revient à étudier un graphe totalement connecté, non pas en termes de sommet (attractivité) mais en termes d'arêtes (déplacements). Ce problème est déjà de nature combinatoire.

On définit alors la fonction f_1 à maximiser qui a pour objectif de favoriser le flux de déplacements entre les stations :

$$f_1(s) = \sum_{st_i \in s} \sum_{st_j \in s \setminus \{st_i\}} f(st_i, st_j) \quad (4.3)$$

Maximiser l'équilibre des stations Ce deuxième objectif vise à réduire les situations de saturation. On cherche ici à éviter qu'une station soit sans voiture ou au contraire sans place disponible pour y déposer une voiture. Il semble évident que le fonctionnement du service est amené à varier d'un jour à l'autre. Il est en effet impossible d'identifier à coup sûr le nombre de voitures qui arrivent à chaque instant ou partent vers une autre station. Aussi plutôt que de chercher à minimiser le nombre de rejets sur un scénario de simulation, nous souhaitons maximiser l'équilibre entre les flux entrants et les flux sortants de chaque site choisi pour y mettre une station, tout au long de la journée.

La Figure 4.7 illustre un cas favorable à l'instant t pour la station s . En effet les flux entrants (verts) et sortants (oranges) sont parfaitement équilibrés.

La station s conserve un nombre stable de voitures et de places disponibles.

On définit la fonction f_2 à maximiser pour favoriser la proportion de flux équilibrés

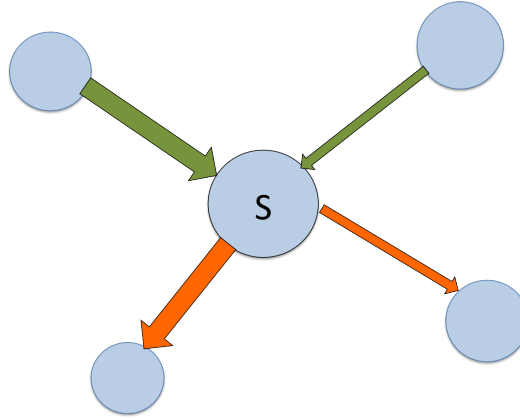


FIGURE 4.7 – Equilibre des flux entrants / flux sortants à l’instant t pour la station s

au cours du temps.

$$f_2(s) = \sum_{st_i \in s} \frac{f_{min}(st_i)}{f_{max}(st_i)} \quad (4.4)$$

Maximiser l’utilisation tout au long de la journée Cette dernière fonction objectif a pour rôle de favoriser le fonctionnement de l’auto-partage tout au long de la journée. La principale motivation à ceci est que l’auto-partage n’a de sens que s’il permet de répondre à de nombreuses demandes, et donc d’être utilisé au cours d’un maximum de plages horaires. L’auto-partage est en effet une solution de transport alternative qui part du principe que les voitures personnelles ne sont que très peu utilisées au cours d’une journée. Il s’agit donc de mutualiser les moyens de transport. Pour que ce soit efficace il est indispensable de disposer également d’un fonctionnement en dehors des heures de pointe. Par ailleurs pour répondre à des demandes en heures de pointe, les approches de type transport en commun semblent plus appropriées.

La Figure 4.8 illustre deux niveaux d’utilisation au long de la journée, en vert une situation favorable vis-à-vis de cet objectif et en rouge un cas défavorable.

Pour formaliser cette fonction f_3 nous proposons de minimiser l’écart type du nombre de déplacements ayant pour origine ou destination chaque station, au cours de la journée. Ainsi cette fonction sera le reflet du niveau de stabilité. On notera également que pour éviter l’effet de facteur d’échelle entre les différents sites, nous utiliserons un écart type relatif à la moyenne des déplacements. De plus, afin d’obtenir une fonction f_3 à maximiser (comme f_1 et f_2), nous considérerons pour f_3 l’opposé de l’écart type.

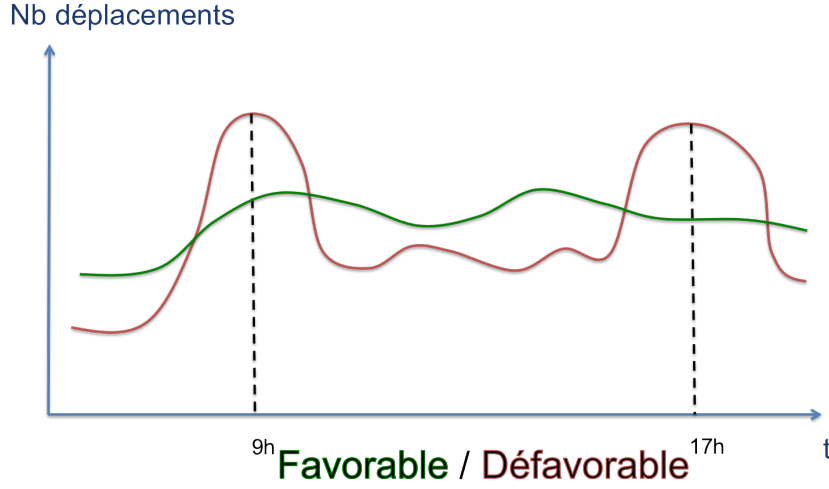


FIGURE 4.8 – Deux fonctionnements d'une station au cours de la journée

$$f_3(s) = - \sum_{st_i \in s} \sqrt{\frac{1}{|T|} \sum_t (f(st_i, t) - \bar{f}(st_i))^2} \quad (4.5)$$

4.4.3.3 Modélisation du problème d'optimisation

La démarche menée pour résoudre ce problème suit une approche bi-niveaux. Lors du placement d'une station la première étape consiste à sélectionner le site candidat, en tenant compte des fonctions objectifs décrites précédemment. Une fois le site candidat identifié, la seconde étape consiste à placer la station plus finement à l'intérieur de la zone couverte par le site candidat (que nous appellerons site fin).

Ce traitement en deux temps permet d'accélérer considérablement les calculs algorithmiques en réduisant de façon significative la quantité de données à manipuler. En effet dans la pratique nous choisissons pour sites candidats des mailles carrées de 300m de côté, alors que le site fin est défini par un maillage de 25m de côté. Il y a donc 144 sites fins dans un site candidat, soit 144 fois plus d'emplacements pour positionner une station lors de la recherche. De plus la matrice de flux contenant la demande en mobilité de site à site augmente avec le carré du nombre de sites, soit d'un facteur $144^2 = 20736$. En considérant une ville d'une taille moyenne la quantité de données à considérer ne peut être traitée par les machines actuelles.

Alors que le choix des sites candidats d'une solution au problème se fait par optimisation combinatoire (en optimisant les fonctions objectifs définies précédemment), le positionnement fin de la station à l'intérieur du site candidat est fait de façon à ce que la station soit à proximité d'un maximum de personnes. Il s'agit dans

cette deuxième phase d'appliquer un simple algorithme de p -médian. Aussi dans la suite nous ne tiendrons pas compte de cette deuxième phase mais uniquement de l'approche combinatoire permettant de choisir n sites candidats pour positionner n stations.

Soit n stations à déployer sur le territoire. Une solution s admissible est un ensemble de n sites candidats soumis à la contrainte suivante :

$$s = \{st_1, \dots, st_n\} \text{ tel que } \forall i, j \in [1, n], st_i = st_j \Rightarrow i = j \quad (4.6)$$

Comme défini précédemment on notera Ω l'ensemble des solutions admissibles.

Par suite le problème d'optimisation *LOC* visant à localiser les stations devient :

$$LOC \begin{cases} \max_{s \in \Omega} & z = f(s) \\ \text{où} & f(s) = (f_1(s), f_2(s), f_3(s)) \\ \text{tel que :} & s = \{st_1, \dots, st_n\} \in \Omega \\ & \forall i, j \in [1, n] st_i = st_j \Rightarrow i = j \end{cases} \quad (4.7)$$

Le problème d'optimisation étant posé, il reste à définir la relation de voisinage à utiliser pour la recherche locale. Ainsi le voisinage d'une solution s est composé de toutes les solutions qui peuvent être atteintes en déplaçant une station quelconque vers n'importe quel autre site disponible. La relation de voisinage s'écrit donc :

$$v(s) = \{s' \in \Omega | dist(s, s') = 1\} \quad (4.8)$$

La fonction *dist* désigne la distance de Hamming entre s et s' .

4.5 Application de l'algorithme mémétique MEMO

Dans cette section sont présentés les résultats obtenus ainsi que les méthodes permettant de positionner l'algorithme MEMO par rapport aux approches existantes. En effet, la comparaison des performances d'algorithmes multiobjectifs est particulièrement délicate. Pour cause, le résultat n'étant pas une solution unique mais un ensemble de solutions non-dominées entre elles, il importe de définir des éléments de comparaison, tels que des indicateurs qui permettent d'attribuer une valeur de fitness à un ensemble de solution.

La non-comparabilité de deux ensembles est illustrée avec la figure 4.9. On voit dans cet exemple que toutes les solutions d'un ensemble ne sont pas dominées par

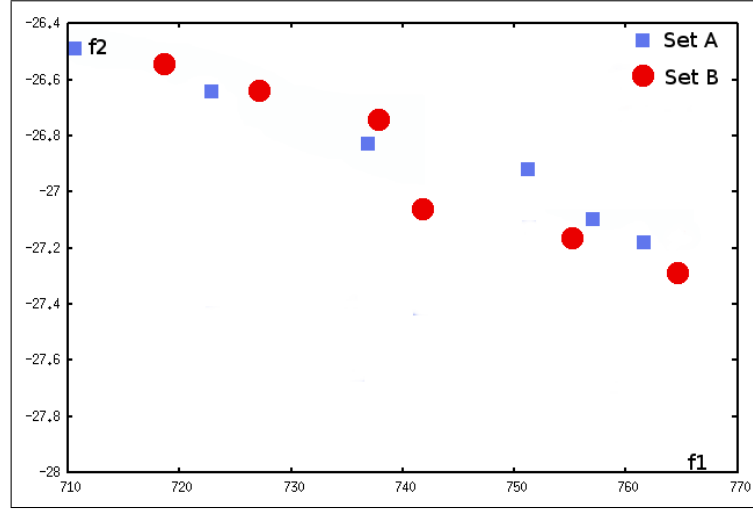


FIGURE 4.9 – Problème multiobjectif de maximisation

des solutions de l'autre ensemble, seule condition permettant de considérer qu'un ensemble est meilleur que l'autre.

Nous présenterons tout d'abord les indicateurs retenus pour évaluer notre algorithme, ce qui ne nous permettra pas de dire qu'une solution est meilleure que l'autre en général mais qu'elle est meilleure relativement à l'indicateur considéré.

4.5.1 Introduction des indicateurs de comparaison

Parmi les indicateurs présentés au chapitre 3, nous en avons retenu trois pour étudier les résultats fournis par MEMO :

- l'indicateur ε additif [Zitzler 2003]
- l'indicateur de contribution [Meunier 2000]
- l'indicateur hypervolume [Zitzler 1999]

On notera également le nombre de solutions non-dominées trouvées par les algorithmes étudiés, ce qui constituera également un indicateur.

4.5.1.1 Indicateur ε -additif

Cet indicateur donne la translation minimale ε à appliquer dans chaque dimension à l'ensemble A pour qu'il domine l'ensemble B .

$$I_{\varepsilon+}(A, B) = \min_{\varepsilon \in \mathbb{R}} \{ \forall x \in B, \exists x' \in A : x \preceq_{\varepsilon+} x' \} \quad (4.9)$$

Etant donné que la valeur ε est la même dans chaque dimension (une dimension est une fonction objectif), il apparaît immédiatement que l'indicateur $I_{\varepsilon+}$ est sensible aux plages de valeurs de chaque objectif. En effet si un objectif admet comme valeurs possible $v \in [0, 1]$ alors qu'un autre objectif a pour domaine de définition $[0, 10000]$, on remarquera que le deuxième objectif risque de masquer totalement le premier. C'est pourquoi une phase de normalisation des valeurs est indispensable avant de calculer l'indicateur $I_{\varepsilon+}$.

Une valeur de $I_{\varepsilon+}$ proche de 0 est le reflet de deux ensembles de qualité semblable. Si la valeur devient négative alors l'ensemble mesuré A est significativement meilleur que l'ensemble de référence B .

4.5.1.2 Indicateur de contribution

Lors de la comparaison de deux ensembles résultats A et B , l'indicateur de contribution noté $I_C(A, B)$ détermine la proportion de solutions non dominées apportée par A dans $ND(A \cup B)$, avec ND retournant les solutions non-dominées.

$$I_C(A, B) = \frac{\frac{|A \cap B|}{2} + |W_A| + |N_A|}{|ND(A \cup B)|} \quad (4.10)$$

W_A désigne l'ensemble des solutions de A qui dominent au moins une solution de B et N_A l'ensemble des solutions de A non-comparables avec B .

De façon évidente une valeur de I_C la plus grande possible est souhaitée.

4.5.1.3 Indicateur hypervolume

Il s'agit d'un indicateur incontournable depuis sa proposition par Zitzler et Thiele [Zitzler 1999]. L'indicateur hypervolume I_H présente en effet l'avantage de refléter la qualité d'un ensemble, tant en termes d'intensité (bonnes solutions) que de diversité (bonne couverture du front de Pareto). De plus cet indicateur présente l'avantage d'être unaire : il fournit le niveau de qualité d'un ensemble de solutions sans nécessiter le recours à un ensemble de référence. L'indicateur doit fournir l'hypervolume le plus vaste possible pour une bonne population. Toutefois deux conditions sont indispensables à une bonne interprétation :

- de disposer d'une normalisation des objectifs pour qu'un objectif ne soit pas sur-représenté ou sous-représenté
- d'identifier un point de référence pour le calcul de l'hypervolume

Pour normaliser les objectifs nous pouvons déterminer les valeurs extrêmes sur chaque objectif pour les ensembles Pareto obtenus au cours de nombreux runs de diverses méthodes. De cette façon chaque objectif normalisé est compris entre 0 et 1. Il n'est pas exclu de trouver une valeur normalisée < 0 (solution pire que la pire solution trouvée au cours des nombreux runs) ou encore > 1 (on améliore la meilleure solution trouvée) lors d'une nouvelle exécution.

En théorie, pour l'objectif n , $\forall s \in S, f_{n_{min}} \leq f_n(s) \leq f_{n_{max}}$. Il s'en suit que nous avons la normalisation :

$$0 \leq \frac{f_n(s) - f_{n_{min}}}{f_{n_{max}} - f_{n_{min}}} \leq 1 \quad (4.11)$$

4.5.2 Définition d'un ensemble de référence avec NSGA-II, PLS et FLSMO

Il ne fait aucun doute que les résultats obtenus lors d'une exécution d'un algorithme stochastique ne permettent pas de conclure sur les performances de l'algorithme en question. En effet la nature aléatoire de ces algorithmes fait qu'une exécution peut être bonne et la suivante médiocre. Pour pallier ce problème, une moyenne est généralement calculée sur plusieurs exécutions. Dans notre situation nous avons réalisé pour chaque test une série de 5 exécutions afin de déterminer des valeurs moyennes, que ce soit pour l'évaluation des fonctions objectifs ou des indicateurs.

Bien que nous jugerons de la performance d'un algorithme sur les valeurs moyennes obtenues au cours des 5 exécutions, on remarquera que d'autres critères d'évaluation pourraient être utilisés. En particulier il serait possible de considérer l'écart type des résultats obtenus afin de déterminer le niveau de stabilité de l'algorithme, ou encore le meilleur résultat sur 5 exécutions si l'on souhaite privilégier un algorithme très efficace de temps en temps, et en acceptant qu'il soit peu performant le reste du temps.

L'étude de la performance d'un algorithme suppose de définir des critères de comparaison comme nous l'avons présenté en introduction, mais aussi de disposer d'un ensemble de référence auquel se comparer. Pour permettre d'avoir le meilleur recul sur les résultats fournis par différentes approches, nous allons comparer les résultats obtenus avec un grand nombre de configurations d'algorithmes, pour des recherches locales et des approches évolutionnaires, avant d'intégrer notre algorithme mémétique.

Etant donné que MEMO repose dessus, l'algorithme génétique de référence sera NSGA-II. Pour la recherche locale nous allons comparer deux algorithmes. PLS qui présente l'avantage de nécessiter peu de paramétrage et qui fait référence en optimisation combinatoire multi-objectif, et FLSMO qui est un algorithme que nous

avons proposé et qui est également simple à mettre en oeuvre.

Remarque : Un algorithme d'optimisation a pour vocation de livrer une solution (ou un ensemble de solutions). Il est alors nécessaire qu'il soit fini, ce qui passe par la définition d'un critère d'arrêt. Lors d'un algorithme de descente un tel critère d'arrêt peut être l'optimum local, solution n'admettant aucune solution améliorante dans son voisinage. Dans le cas multi-objectif il s'agirait d'une solution n'admettant aucun voisin non-dominé par une solution précédemment trouvée. Avec des approches évolutionnaires il est toujours possible de croiser des individus pour construire de nouvelles générations. Le critère d'arrêt n'est donc pas directement inclus dans l'algorithme de recherche. Il faut alors définir une porte de sortie. Le critère que nous choisissons est un critère de temps d'exécution. Nous remarquons qu'en moyenne l'algorithme PLS de référence atteint un optimum local après 2 heures d'exécution. Nous choisissons donc cette durée comme durée de référence pour arrêter notre algorithme MEMO ainsi que NSGA-II.

Conditions de test des algorithmes :

- chaque configuration est testée sur 5 runs.
- le temps limite d'exécution d'un run est fixé à 2h.

4.5.2.1 Référence avec NSGA-II

Plusieurs facteurs agissent sur les performances de cet algorithme. Il apparaît alors nécessaire de fixer au mieux ces éléments. Nous avons agi sur les trois principaux leviers :

- la taille de la population : 400, 600, 800, 1000.
- l'opérateur de croisement : aléatoire, élitiste.
- l'opérateur de mutation : aléatoire, premier voisin améliorant, meilleur voisin améliorant.

Nous avons testé les six combinaisons d'opérateurs pour chaque taille de population. Chacune des configurations a été exécutée cinq fois afin de fournir la performance moyenne, le minimum, le maximum et l'écart type.

Les résultats sont présentés à chaque fois dans un tableau indiquant les valeurs sur les trois objectifs, le nombre de solutions non-dominées trouvées, et l'indicateur hypervolume.

Croisement aléatoire / mutation aléatoire Il s'agit là de l'AG dans ses mécanismes les plus simples. Seuls les mécanismes de NSGA-II (sélection des meilleurs parents et remplacement des moins bons individus) agissent pour faire converger la population vers le front de Pareto. Le croisement aléatoire prend successivement un gène aléatoire non encore pris de chaque parent. De même la mutation qui intervient avec une

probabilité de 10% prend aléatoirement un gène (une station) et le fait muter vers une valeur aléatoire (emplacement disponible).

Le meilleur choix en termes d'hypervolume moyen est la population de taille 1000 comme le montre le tableau 4.2. De manière générale, on note que la qualité des résultats vue par I_H augmente avec la taille de la population. De même le nombre de solutions du front augmente avec la taille de la population.

TABLE 4.2 – NSGA-II : croisement aléatoire et mutation aléatoire

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 400	min	22209.086	1184.555	-27.141	3316	0.280
	max	22787.135	1196.755	-26.033	3871	0.290
	mean	22517.829	1191.554	-26.639	3565	0.286
	std	-	-	-	-	0.004
Pop 600	min	21975.203	1179.223	-26.693	3865	0.284
	max	22787.135	1225.583	-26.115	4115	0.292
	mean	22448.830	1196.253	-26.402	3979	0.289
	std	-	-	-	-	0.003
Pop 800	min	21300.516	1189.853	-27.115	4070	0.289
	max	22787.135	1360.691	-26.013	4370	0.324
	mean	22339.576	1226.781	-26.508	4265	0.297
	std	-	-	-	-	0.014
Pop 1000	min	21975.203	1184.091	-26.413	4031	0.285
	max	22718.160	1379.408	-25.870	4608	0.331
	mean	22285.594	1264.523	-26.119	4406	0.306
	std	-	-	-	-	0.021

Croisement aléatoire / mutation premier voisin améliorant Le croisement aléatoire est conservé alors que la mutation joue un rôle plus élitiste. La mutation intervient toujours avec une probabilité de 10% mais cette fois-ci le gène ayant la plus faible contribution à l'évaluation globale est choisi pour être muté. Pour définir la plus faible contribution les valeurs fitness sont agrégées selon une pondération aléatoire. Une fois le gène à muter choisi il prend la première valeur proposant une amélioration de la fitness générale selon la même pondération aléatoire des objectifs. On notera que pour ne favoriser aucune agrégation, et donc aucune direction, pour chaque enfant une nouvelle valeur aléatoire est retenue. Cet opérateur de mutation revient à prendre le premier voisin améliorant.

Le meilleur hypervolume moyen observé dans le tableau 4.3 est là encore obtenu avec la population de taille 1000 et les comportements sont semblables à ceux observés précédemment.

TABLE 4.3 – NSGA-II : croisement aléatoire et mutation premier voisin améliorant

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 400	min	22787.135	1031.162	-26.765	2958	0.236
	max	22787.135	1196.755	-25.840	4113	0.292
	mean	22787.135	1163.636	-26.478	3828	0.281
	std	-	-	-	-	0.022
Pop 600	min	22787.135	1196.755	-27.002	4281	0.292
	max	22787.135	1196.755	-26.167	4395	0.293
	mean	22787.135	1196.755	-26.396	4358	0.293
	std	-	-	-	-	0.000
Pop 800	min	22787.135	1196.200	-26.272	4537	0.292
	max	22787.135	1196.755	-25.956	4701	0.294
	mean	22787.135	1196.644	-26.164	4615	0.293
	std	-	-	-	-	0.000
Pop 1000	min	22787.135	1196.755	-26.389	4706	0.293
	max	22787.135	1379.408	-26.018	5150	0.339
	mean	22787.135	1269.816	-26.169	4863	0.311
	std	-	-	-	-	0.022

Croisement aléatoire / mutation meilleur voisin améliorant Pour cette série de tests le croisement reste inchangé, mais la mutation transforme le gène à la plus mauvaise contribution vers la valeur fournissant la meilleure fitness agrégée possible. L'évaluation des solutions se fait également par agrégation des fitness selon une pondération aléatoire. Cela revient à choisir le meilleur voisin.

Le tableau 4.4 montre que les meilleurs résultats sont obtenus avec la population de taille 800 ainsi qu'avec la population de taille 400. Toutefois aucun cas ne permet d'atteindre le seuil de 0,3 pour I_H alors que les populations de taille 1000 l'atteignaient en moyenne dans les deux versions précédentes.

On remarque que les écarts types sont réduits pour toutes les tailles de population ; le choix du meilleur voisin en mutation rend l'algorithme plus déterministe et donc plus prévisible, mais moins bon en termes d'exploration.

TABLE 4.4 – NSGA-II : croisement aléatoire et mutation meilleur voisin améliorant

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 400	min	22787.135	1196.755	-26.932	3356	0.291
	max	22787.135	1225.583	-26.691	3571	0.297
	mean	22787.135	1202.520	-26.767	3492	0.292
	std	-	-	-	-	0.003
Pop 600	min	22787.135	1196.755	-27.123	3634	0.291
	max	22787.135	1196.755	-26.366	3922	0.292
	mean	22787.135	1196.755	-26.765	3771	0.291
	std	-	-	-	-	0.000
Pop 800	min	22787.135	1189.853	-26.686	3931	0.291
	max	22787.135	1196.755	-25.933	4187	0.293
	mean	22787.135	1195.374	-26.281	4061	0.292
	std	-	-	-	-	0.000
Pop 1000	min	22787.135	1196.755	-26.863	3855	0.290
	max	22787.135	1196.755	-26.362	4206	0.292
	mean	22787.135	1196.755	-26.614	3963	0.291
	std	-	-	-	-	0.001

Croisement élitiste / mutation aléatoire Le croisement élitiste est présenté en section 4.3.1.2. Contrairement au croisement aléatoire dans lequel chaque parent transmet au hasard la moitié de ses gènes à l'enfant, dans le cas élitiste les gènes transmis sont choisis en fonction de leur qualité. Pour déterminer ce niveau de qualité, une agrégation de la contribution des gènes à chaque objectif est réalisée, suivant la pondération choisie aléatoirement. La mutation aléatoire est identique à celle présentée précédemment.

Cette combinaison s'avère moins performante que la méthode avec croisement aléatoire et mutation aléatoire (cf. 4.2) sur I_H , avec en plus un algorithme plus instable (écart type de I_H plus grand).

TABLE 4.5 – NSGA-II : croisement élitiste et mutation aléatoire

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 400	min	22787.135	1077.372	-27.001	2004	0.239
	max	22787.135	1238.517	-26.562	3180	0.293
	mean	22787.135	1163.751	-26.742	2590	0.271
	std	-	-	-	-	0.025
Pop 600	min	22787.135	1077.372	-26.768	2125	0.240
	max	22787.135	1184.555	-26.590	3223	0.289
	mean	22787.135	1143.831	-26.708	2781	0.269
	std	-	-	-	-	0.023
Pop 800	min	22787.135	1068.192	-26.978	2385	0.239
	max	22787.135	1266.107	-26.478	3618	0.297
	mean	22787.135	1201.474	-26.799	3197	0.284
	std	-	-	-	-	0.022
Pop 1000	min	22787.135	1087.381	-27.016	2484	0.244
	max	22787.135	1379.408	-26.689	3963	0.337
	mean	22787.135	1223.906	-26.847	3408	0.292
	std	-	-	-	-	0.030

Croisement élitiste / mutation premier voisin améliorant Les résultats sont issus de la combinaison des opérateurs de croisement et mutation définis au-dessus.

Là aussi, la combinaison de ces opérateurs est moins performante que la même mutation combinée avec le croisement aléatoire. L'instabilité de l'algorithme se confirme sur l'analyse des écarts types de I_H .

TABLE 4.6 – NSGA-II : croisement élitiste et mutation premier voisin améliorant

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 400	min	22787.135	1031.162	-27.230	2251	0.237
	max	22787.135	1256.312	-26.790	3524	0.296
	mean	22787.135	1172.886	-26.940	3151	0.282
	std	-	-	-	-	0.023
Pop 600	min	22787.135	1028.761	-26.891	2246	0.237
	max	22787.135	1217.010	-26.590	3682	0.294
	mean	22787.135	1122.358	-26.722	2882	0.262
	std	-	-	-	-	0.026
Pop 800	min	22787.135	1077.942	-26.872	2370	0.243
	max	22787.135	1256.312	-26.362	3819	0.298
	mean	22787.135	1187.050	-26.708	3232	0.276
	std	-	-	-	-	0.026
Pop 1000	min	22787.135	1031.162	-26.962	2661	0.237
	max	22787.135	1266.107	-26.355	3877	0.297
	mean	22787.135	1185.167	-26.785	3426	0.283
	std	-	-	-	-	0.023

Croisement élitiste / mutation meilleur voisin améliorant Là encore les résultats sont issus de la combinaison des opérateurs de croisement et mutation définis au-dessus.

Cette combinaison offre à peu près les mêmes performances qu'avec le croisement aléatoire, mais avec plus d'instabilité sur l'écart type de I_H .

TABLE 4.7 – NSGA-II : croisement élitiste et mutation meilleur voisin améliorant

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 400	min	22787.135	1196.200	-27.206	2540	0.291
	max	22787.135	1237.517	-26.846	3097	0.298
	mean	22787.135	1211.265	-27.015	2733	0.294
	std	-	-	-	-	0.003
Pop 600	min	22787.135	1184.555	-26.923	2749	0.292
	max	22787.135	1256.312	-26.729	3010	0.299
	mean	22787.135	1211.669	-26.845	2891	0.295
	std	-	-	-	-	0.003
Pop 800	min	22787.135	1031.162	-27.039	2063	0.237
	max	22787.135	1256.312	-26.435	3201	0.297
	mean	22787.135	1185.407	-26.817	2893	0.283
	std	-	-	-	-	0.023
Pop 1000	min	22787.135	1087.381	-26.885	2222	0.244
	max	22787.135	1256.312	-26.651	3369	0.299
	mean	22787.135	1205.781	-26.768	2982	0.287
	std	-	-	-	-	0.021

Synthèse des résultats NSGA-II Il apparait après observation des résultats que les meilleures performances sont obtenues avec un croisement aléatoire, combiné avec une mutation vers le premier voisin améliorant pour une population de 1000 individus, et combiné avec la mutation aléatoire pour une population de 800 individus.

Pour les populations de petite taille (400 et 600 individus), la méthode la plus performante est le croisement élitiste avec une mutation vers le meilleur voisin améliorant (cf. tableau 4.7), mais elle reste inférieure au croisement aléatoire appliqué sur 1000 individus. On voit donc que la taille de la population a un impact très fort sur la méthode. Les grandes populations compensent les marches aléatoires du fait de la diversité des solutions ; au contraire les petites populations ne sont efficaces qu'avec des marches plus déterministes car il n'y a pas assez de diversité dans la population.

La Figure 4.10 présente la façon dont évolue l'indicateur I_H pendant la recherche pour différentes configurations de NSGA-II. Ceci permet de savoir comment la population se rapproche du front de Pareto au cours du temps, en fonction des opérateurs génétiques et pour la meilleure taille de population. On rappelle qu'un hypervolume grand est gage de qualité.

On remarque que la meilleure version de NSGA-II (croisement aléatoire, mutation vers le premier voisin améliorant et taille de population 1000) fournit les meilleurs résultats très rapidement. On voit notamment qu'après moins de 200.10^5 itérations

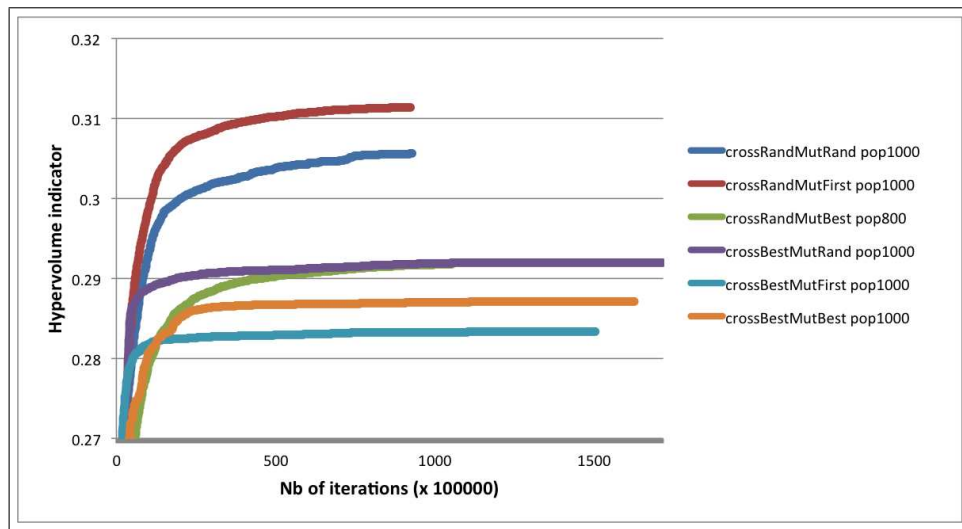


FIGURE 4.10 – Evolution des indicateurs hypervolumes moyens de 6 configurations NSGA-II

elle surpasse toutes les autres configurations.

4.5.2.2 Référence avec PLS

Cette approche locale à base de Pareto est décrite de façon détaillée dans le chapitre 3. Nous rappelons également en section 4.5.2.3 quelques éléments qui caractérisent PLS. Les résultats obtenus sont synthétisés dans le tableau 4.8. Le premier élément marquant concerne le nombre de solutions non-dominées trouvées par PLS. Alors que NSGA-II atteint une seule fois les 5000 solutions trouvées parmi toutes les configurations testées, PLS atteint les 6400 solutions. Par ailleurs, la meilleure valeur de l'indicateur I_H trouvée par PLS est significativement moins bonne que celle trouvée pour NSGA-II. Ainsi PLS permet de trouver beaucoup de solutions mais elles apportent souvent peu de diversité.

TABLE 4.8 – Résultats de la recherche locale PLS

		f_1	f_2	f_3	NbSol	Hypervolume
PLS	min	22787.135	1024.459	-25.184	4146	0.239
	max	22787.135	1417.612	-25.047	6401	0.264
	mean	22787.135	1126.099	-25.081	4808	0.245
	std	-	-	-	-	0.011

4.5.2.3 Référence avec FLSMO : un nouvel algorithme de recherche locale

L'algorithme FLSMO (Fast Local Search Multiobjective) est une recherche locale basée sur la notion de Pareto optimalité. Nous l'avons proposée début 2013 [Moalic 2013a] pour résoudre ce problème de l'auto-partage. Il s'agit d'un algorithme proche de PLS [Paquete 2004] et qui s'est révélé être particulièrement efficace. Comme nous le verrons sur les Figures 4.12 et 4.13, il présente l'avantage de converger plus rapidement que PLS.

Le principe est le suivant. L'algorithme démarre avec une première solution générée aléatoirement qui est marquée comme non exploitée. Elle est alors ajoutée à une population archive. Tant qu'il y a des solutions dans l'archive qui ne sont pas exploitées, l'algorithme en choisit une aléatoirement, notée s , et explore son voisinage jusqu'à trouver un voisin s' non-dominé par l'archive. s' est alors ajoutée à l'archive et devient la nouvelle solution courante. Son voisinage est exploré jusqu'à trouver un nouveau voisin non dominé par l'archive. Ce processus est appliqué récursivement jusqu'à être dans un optimum local. Toute solution non dominée par l'archive rencontrée au cours de cette descente est ainsi ajoutée à l'archive. Il s'agit de fait d'une descente réalisée à partir de s . L'algorithme 8 décrit formellement ce processus.

Algorithm 8 Fast Local Search for Multiobjective Problems

```

1:  $S \leftarrow \text{init}()$  /*init the solution set S with a random individual*/
2:  $s \leftarrow \text{select}(S)$  /*select randomly a not explored solution from S*/
3: while  $s \neq \emptyset$  do
4:   repeat
5:      $s' \leftarrow \text{selectNeighbor}(s)$  /*select randomly a neighbor of s not dominated by S*/
6:     if  $s' \neq \emptyset$  then
7:        $s \leftarrow s'$ 
8:        $\text{addNotDominated}(s)$  /*add s in S and remove all dominated solutions*/
9:     end if
10:  until  $s' = \emptyset$ 
11:  mark  $s$  as explored
12:   $s \leftarrow \text{select}(S)$  /*select randomly a not explored solution from S*/
13: end while

```

Cet algorithme présente l'avantage d'assurer une bonne intensification en exploitant les solutions non-dominées déjà trouvées, tout en assurant une bonne diversité.

Pour rappel PLS fonctionne de façon assez semblable. Toutefois au lieu d'effectuer une descente à partir d'une solution non exploitée de l'archive, PLS choisit une solution s de l'archive et explore son voisinage complet pour ajouter tous les voisins non-dominés à l'archive. La solution s est alors marquée comme exploitée.

FLSMO peut ainsi être vu comme un algorithme avec recherche en profondeur d'abord, puis en largeur. PLS par contre s'apparente à un parcours en largeur d'abord, puis en profondeur. La Figure 4.11 illustre ce principe pour les deux algorithmes.

Les résultats de la méthode FLSMO sur 5 exécutions sont reportés dans le tableau 4.9. Le nombre de solutions trouvées avec FLSMO est significativement plus faible qu'avec PLS, en moyenne, en minimum et en maximum. L'hypervolume est également plus petit mais très légèrement seulement. Les solutions de FLSMO apportent en moyenne plus de diversité que celles de PLS.

TABLE 4.9 – Résultats de la recherche locale FLSMO

		f_1	f_2	f_3	NbSol	Hypervolume
FLSMO	min	22787.135	1012.214	-25.168	3992	0.235
	max	22787.135	1077.372	-25.047	4405	0.243
	mean	22787.135	1044.340	-25.104	4246	0.239
	std	-	-	-	-	0.003

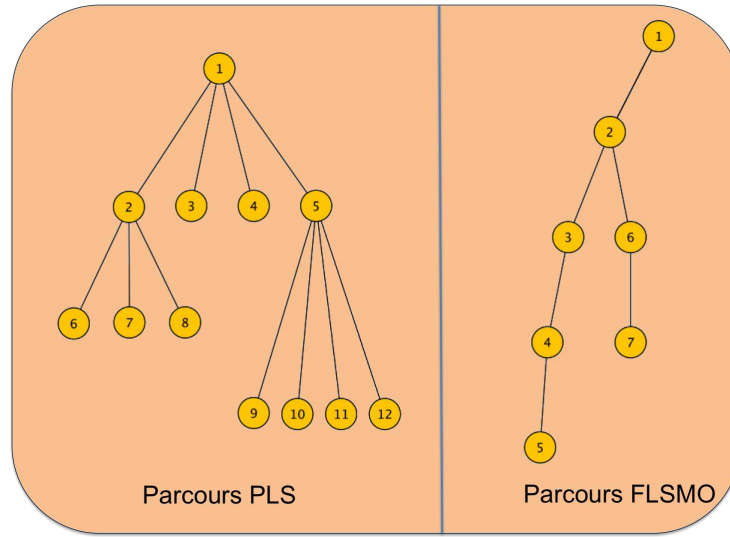


FIGURE 4.11 – Parcours PLS vs FLSMO

4.5.2.4 Synthèse des références calculées

Les résultats obtenus avec l'algorithme génétique NSGA-II, la recherche locale PLS, et la nouvelle recherche locale que nous proposons FLSMO, permettent d'extraire un certain nombre d'enseignements.

Premier enseignement : les approches locales sont significativement moins performantes que l'approche génétique. Même la moins bonne configuration génétique en matière d'hypervolume moyen (0.262 pour le croisement élitiste avec une mutation vers le premier voisin améliorant et une population de taille 600) reste meilleure que les résultats obtenus par les recherches locales.

Deuxième enseignement : on se propose de regarder la façon dont convergent les algorithmes. La Figure 4.12 montre comment en moyenne progressent les valeurs de I_H au cours du déroulement de l'algorithme. On remarque que FLSMO élabore très rapidement une bonne archive, et n'est rattrapé que tardivement par PLS. Ici la contrainte de temps de calcul n'est pas importante mais dans le cas où le temps est un critère à prendre en compte, FLSMO se révèle être une approche très compétitive. Le zoom sur les $2 \cdot 10^6$ premières itérations de la Figure 4.13 montre en effet que FLSMO est meilleur que NSGA-II jusqu'à 500.000 itérations et continue à être nettement meilleur que PLS après $2 \cdot 10^6$ itérations.

Pour confronter notre algorithme mémétique MEMO nous choisirons donc comme référence les solutions trouvées par NSGA-II. L'ensemble sera composé de toutes les solutions non-dominées rencontrées au cours des cinq exécutions de l'algorithme, ce qui en fait une référence a priori solide.

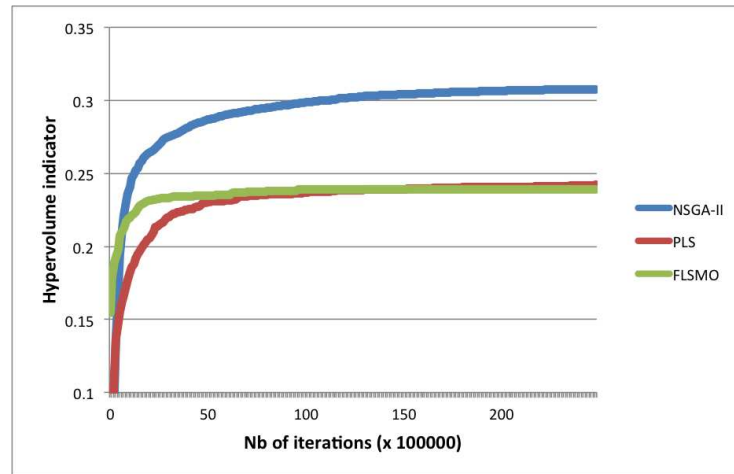


FIGURE 4.12 – Progression des algorithmes de référence

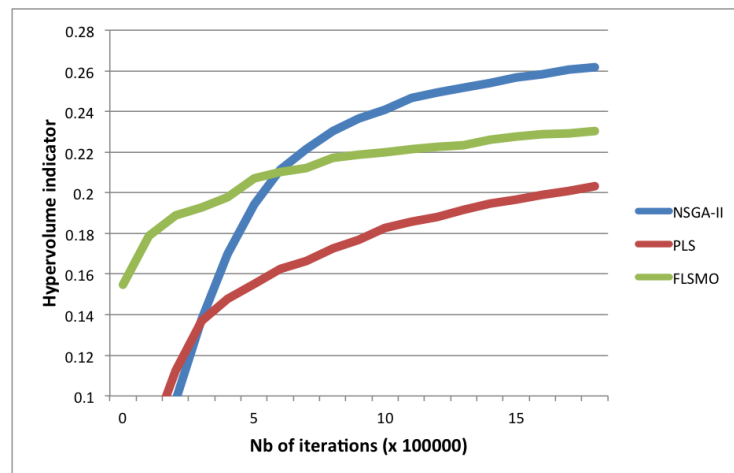


FIGURE 4.13 – Premières itérations des algorithmes de référence

4.5.3 Définition du paramétrage de MEMO

De la même façon que pour NSGA-II, nous analysons l'impact des éléments ayant une influence sur le déroulement de l'algorithme mémétique. En particulier les résultats suivants sont obtenus en faisant varier le croisement, qui comme avec NSGA-II sera soit élitiste soit aléatoire, ainsi que la taille de la population.

Remarque : comme souvent dans le cadre d'algorithmes mémétiques, la taille des populations est significativement plus petite que pour des algorithmes génétiques. Ceci s'explique par le rôle que joue la population dans les deux cas. Comme nous l'avons vu, dans le cadre génétique c'est le croisement qui joue un rôle d'intensificateur (on espère que les enfants soient globalement meilleurs que les parents). De ce fait

la population doit contenir une diversité suffisante pour limiter la "consanguinité" et éviter que la population converge trop vite vers un même individu. Au contraire dans le cas d'un algorithme mémétique le croisement joue un rôle de diversificateur. Or une population trop grande a pour effet de déstructurer ce que la recherche locale a fait (elle apporte trop de diversification). Il est ainsi fréquent de rencontrer des populations de taille comprise entre 5 et 20 selon les problèmes traités.

4.5.3.1 Croisement aléatoire / mutation recherche locale

Le tableau 4.10 montre ce que donne le croisement aléatoire en fonction du nombre d'individus dans la population.

Comme nous pouvons le voir les populations de taille 15 et 20 fournissent le meilleur indicateur hypervolume moyen. La population de taille 20 a trouvé les meilleurs résultats pour les 3 mesures de I_H , en minimum, maximum et moyenne. Mais nous avons globalement de très bons résultats moyens, tous supérieurs à 0,34 pour l'ensemble des populations. L'algorithme est donc relativement peu sensible pour des populations de 5 à 25 individus.

Remarquons également que le nombre de solutions du front est beaucoup plus petit avec MEMO qu'avec les trois autres approches. Le fait d'obtenir de meilleures valeurs pour I_H avec moins de solutions implique que les solutions trouvées sont réellement de meilleure qualité.

TABLE 4.10 – Résultats de MEMO avec croisement aléatoire

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 5	min	22787.135	1780.146	-25.571	2029	0.323
	max	22787.135	1874.759	-25.123	2652	0.350
	mean	22787.135	1838.390	-25.258	2443	0.345
	std	-	-	-	-	0.006
Pop 10	min	22787.135	1783.407	-25.432	2399	0.348
	max	22787.135	1874.759	-24.989	2749	0.352
	mean	22787.135	1857.137	-25.210	2619	0.350
	std	-	-	-	-	0.001
Pop 15	min	22787.135	1520.241	-25.540	2457	0.346
	max	22787.135	1874.759	-25.088	2862	0.352
	mean	22787.135	1832.427	-25.245	2700	0.351
	std	-	-	-	-	0.002
Pop 20	min	22787.135	1783.407	-25.561	2664	0.350
	max	22787.135	1874.759	-25.088	2853	0.353
	mean	22787.135	1853.190	-25.234	2758	0.351
	std	-	-	-	-	0.001
Pop 25	min	22787.135	1473.509	-25.522	2140	0.314
	max	22787.135	1874.759	-25.088	2908	0.353
	mean	22787.135	1825.475	-25.293	2763	0.350
	std	-	-	-	-	0.008

4.5.3.2 Croisement élitiste / mutation recherche locale

Le croisement élitiste avait donné pour NSGA-II des résultats moins bons que le croisement aléatoire. Nous souhaitons vérifier si cette observation se confirme dans le cas mémétique.

Nous voyons dans le tableau 4.11 que les tailles de populations optimales sont également 15 et 20. Par contre, l'indicateur I_H est globalement moins bon que lorsque le croisement se fait aléatoirement. Nous retiendrons donc le croisement aléatoire qui apporte un gain sur la diversification de la recherche.

TABLE 4.11 – Résultats de MEMO avec croisement élitiste

		f_1	f_2	f_3	NbSol	Hypervolume
Pop 5	min	22787.135	1517.531	-25.685	2076	0.309
	max	22787.135	1874.759	-25.088	2912	0.349
	mean	22787.135	1746.403	-25.332	2590	0.340
	std	-	-	-	-	0.014
Pop 10	min	22787.135	1520.241	-25.778	2435	0.339
	max	22787.135	1874.759	-25.192	3104	0.353
	mean	22787.135	1756.767	-25.334	2861	0.348
	std	-	-	-	-	0.003
Pop 15	min	22787.135	1313.459	-25.610	2418	0.301
	max	22787.135	1874.759	-25.116	3174	0.353
	mean	22787.135	1801.014	-25.316	2937	0.349
	std	-	-	-	-	0.011
Pop 20	min	22787.135	1415.791	-25.685	2338	0.338
	max	22787.135	1874.759	-25.022	3264	0.354
	mean	22787.135	1716.701	-25.337	3064	0.349
	std	-	-	-	-	0.004
Pop 25	min	22787.135	1473.509	-25.777	2629	0.311
	max	22787.135	1874.759	-25.033	3315	0.354
	mean	22787.135	1676.512	-25.390	3060	0.346
	std	-	-	-	-	0.011

4.5.3.3 Synthèse de l'algorithme mémétique MEMO

L'algorithme MEMO a montré que, tout comme NSGA-II, le meilleur fonctionnement est obtenu avec un croisement aléatoire. De plus une population de taille 20 fournit les meilleurs résultats pour les deux types de croisement. On voit ainsi qu'un croisement trop intensif nuit à la recherche globale de l'algorithme. Il s'agit d'un équilibre subtil à trouver entre diversification et intensification.

Pour la suite de l'analyse nous utiliserons donc cette configuration de l'algorithme en comparaison avec les algorithmes de référence.

4.5.4 Analyse des solutions de MEMO

La sélection de l'ensemble de référence et de la configuration MEMO est basée sur l'indicateur hypervolume. Toutefois comme nous l'avons vu d'autres indicateurs existent et demeurent intéressants à étudier. Dans le tableau 4.12 nous reprenons les résultats de PLS, NSGA-II et MEMO en considérant également les indicateurs

$I_{\varepsilon+}$ et I_C . Rappelons que $I_{\varepsilon+}$ est d'autant meilleur qu'il est petit, contrairement à I_H et I_C .

TABLE 4.12 – Synthèse des résultats obtenus avec les algorithmes présentés

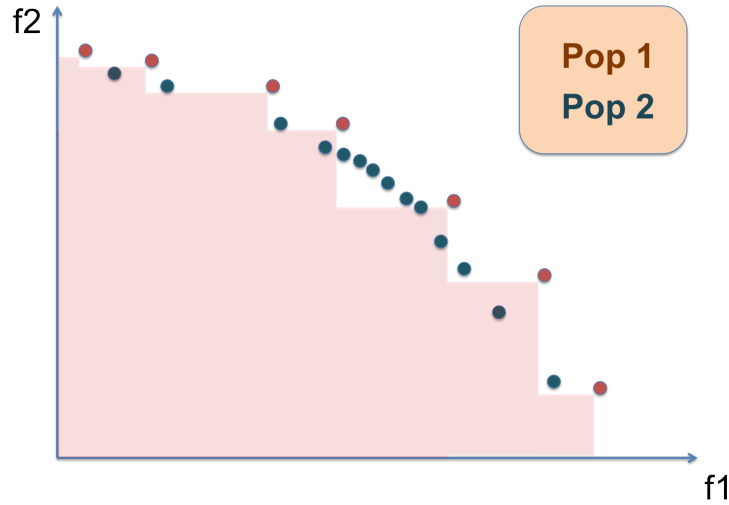
		NbSol	I_H		$I_{\varepsilon+}$		I_C	
PLS	min	4146	0.239	0%	40.276	0%	0.417	0%
	max	6401	0.264	0%	24.261	0%	0.499	0%
	mean	4808	0.245	0%	36.173	0%	0.442	0%
	std	-	0.011	-	6.878	-	0.033	-
FLSMO	min	3992	0.235	-1%	43.309	+7%	0.378	-9%
	max	4405	0.243	-8%	35.647	+47%	0.437	-12%
	mean	4246	0.239	-2%	39.185	+8%	0.416	-6%
	std	-	0.003	-	2.980	-	0.021	-
NSGA-II	min	4706	0.293	+22%	15.078	-62%	0.118	-71%
	max	5150	0.339	+28%	0.649	-97%	0.309	-38%
	mean	4863	0.311	+27%	9.330	-74%	0.189	-57%
	std	-	0.022	-	6.731	-	0.081	-
MEMO	min	2664	0.350	+46%	2.097	-94%	0.237	-43%
	max	2853	0.353	+33%	1.544	-93%	0.266	-46%
	mean	2758	0.351	+43%	1.815	-95%	0.248	-44%
	std	-	0.001	-	0.178	-	0.007	-

L'indicateur $I_{\varepsilon+}$ se révèle être très corrélé avec I_H . Ainsi l'algorithme mémétique obtient la meilleure valeur pour I_H et pour $I_{\varepsilon+}$.

Par contre l'indicateur de contribution est nettement meilleur avec PLS qu'avec MEMO ou même NSGA-II. Tentons une explication. La recherche locale a tendance à trouver beaucoup de solutions proches les unes des autres. C'est le principe même de la progression par étude du voisinage. Ainsi si beaucoup de solutions se trouvent dans une zone non-dominée par les solutions de référence cela a pour effet d'augmenter la valeur de I_C , même si ces solutions ne semblent pas très intéressantes.

La Figure 4.14 montre ainsi deux populations pour un problème bi-objectif de maximisation. Bien que la population rouge soit plus intéressante que la population verte, en tout cas en termes de qualité des solutions du front, la population verte a le meilleur indicateur de contribution. En effet, lorsque nous considérons toutes les solutions (rouges et vertes) non-dominées, il reste une majorité de solutions vertes. Dans notre cas PLS retient 4808 solutions non-dominées contre 2758 pour MEMO ; sur cet indicateur cela donne un avantage certain à PLS. Par contre, même avec 4863 sur le front, NSGA-II reste plus mauvais que MEMO sur cet indicateur.

Cela montre bien que le choix d'un indicateur pour étudier la qualité d'une population en optimisation multiobjectif est primordial. Si seul l'indicateur de contribution

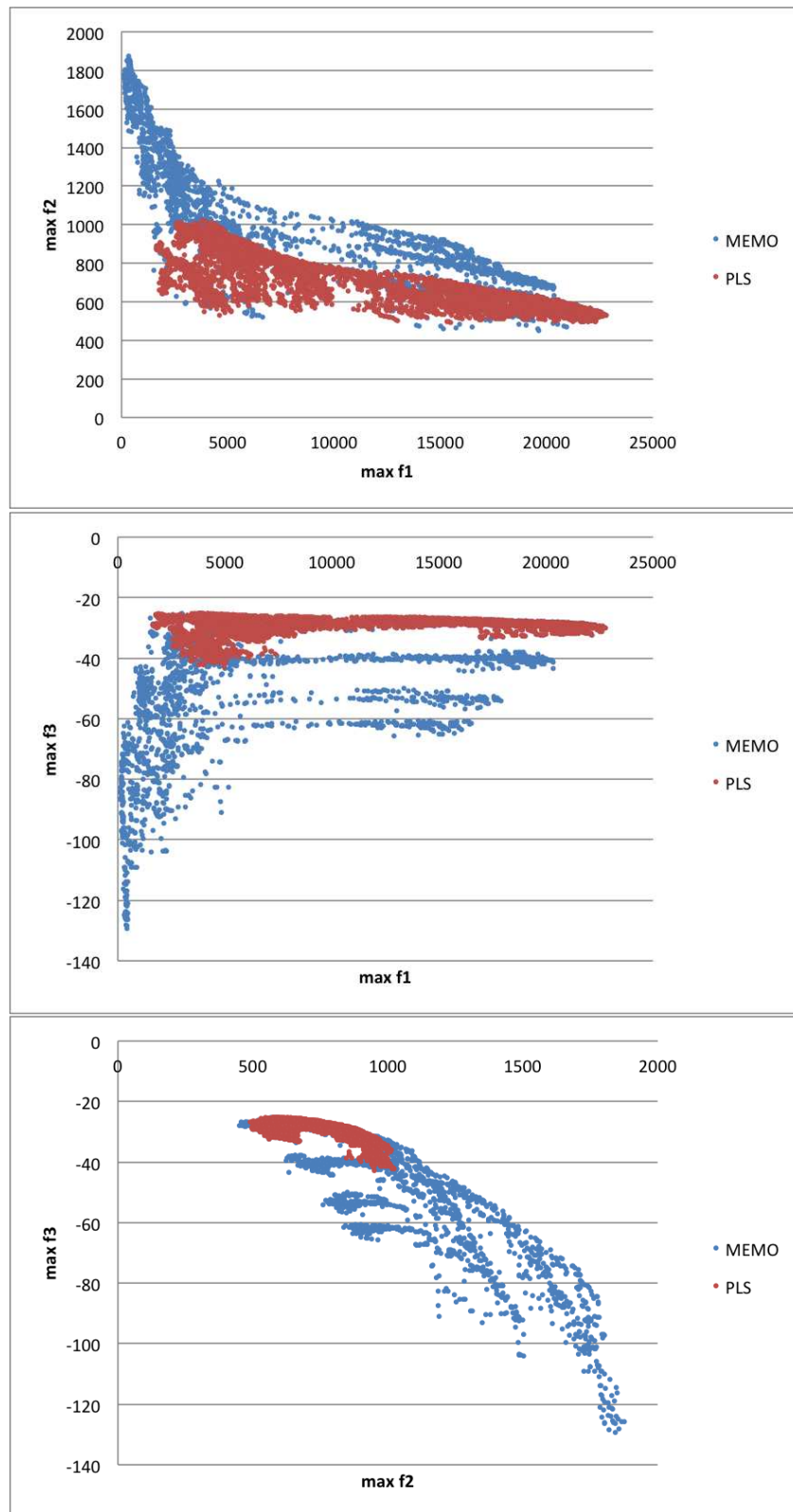
FIGURE 4.14 – Comparaison de deux ensembles pour I_C

avait été retenu on aurait conclu que PLS était l'approche la plus performante, alors qu'elle explore une partie moindre de l'espace de recherche et que sur tous les objectifs elle trouve globalement de moins bonnes valeurs.

La visualisation des nuages de solutions dans l'espace des critères de la Figure 4.15 montre qu'en effet MEMO est bien meilleur que PLS, en considérant les projections sur chacun des objectifs. Chaque couleur correspond à un algorithme : en bleu les résultats de MEMO et en rouge les résultats de référence obtenus avec PLS.

Il est ainsi possible de voir que PLS reste cantonné à une petite partie de l'espace de recherche alors que MEMO trouve des solutions bien mieux diversifiées : MEMO permet de trouver des solutions dans des sous-espaces de recherche inexplorés par PLS. De plus, en termes d'intensification PLS trouve plus de solutions proches les unes des autres mais de moins bonne qualité. Notamment pour f_2 MEMO trouve des solutions bien meilleures que PLS.

Comparons maintenant les projections des solutions de MEMO et de NSGA-II. Dans un premier temps la Figure 4.16 montre ce que révélaient les indicateurs $I_{\varepsilon+}$ et I_H , à savoir que MEMO construit un ensemble de solutions meilleur que NSGA-II, tant en termes de diversité que d'intensité. Mais en comparant également les solutions rouges des Figures 4.15 et 4.16 (PLS et NSGA-II), on vérifie bien que NSGA-II propose des solutions mieux diversifiées que PLS.

FIGURE 4.15 – Solutions PLS et MEMO projetées sur f_1 - f_2 ; f_1 - f_3 et f_2 - f_3

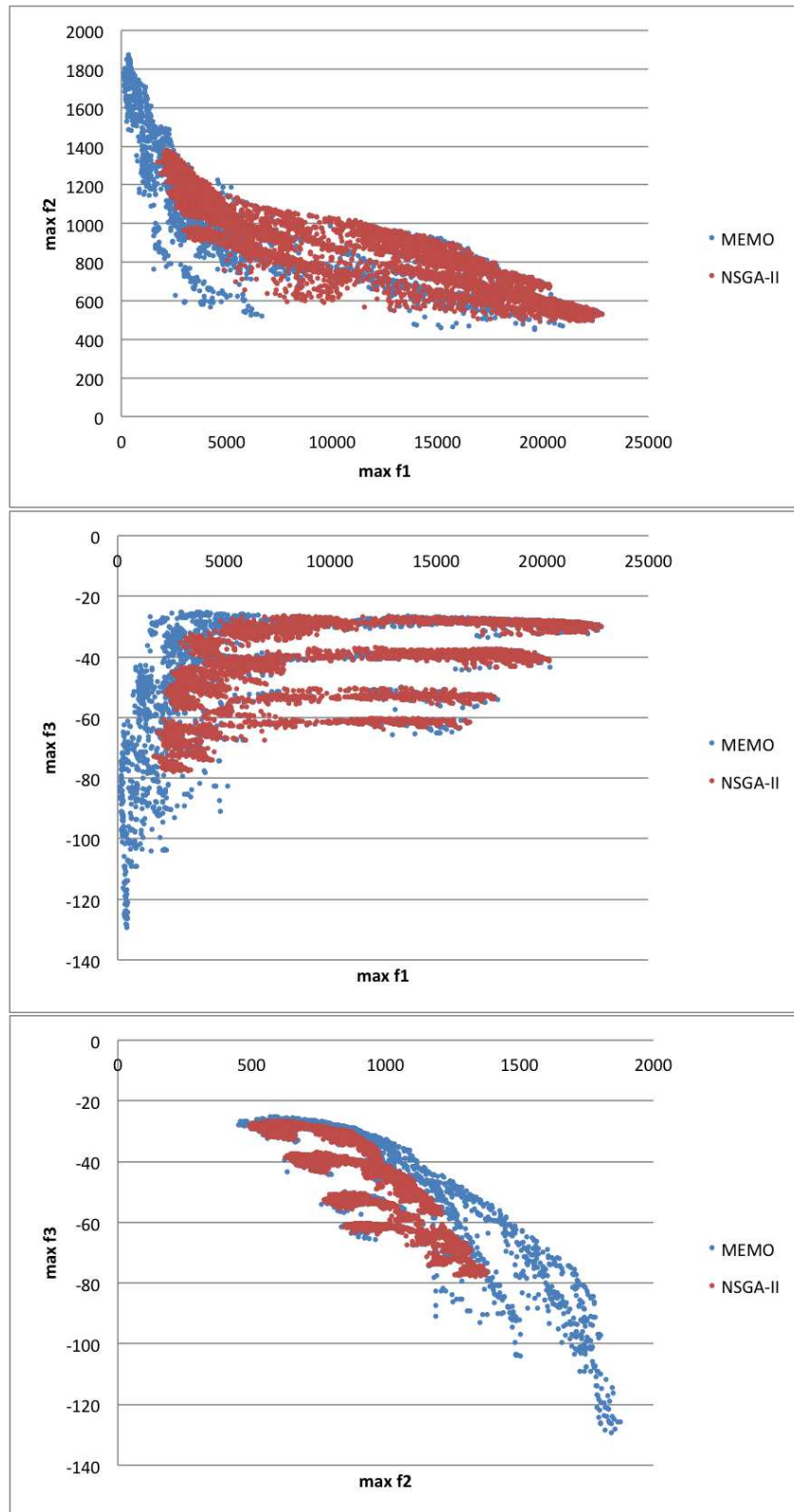


FIGURE 4.16 – Solutions NSGA-II et MEMO projetées sur f_1 - f_2 ; f_1 - f_3 et f_2 - f_3

4.5.5 Analyse de la progression de l'algorithme MEMO

Pour bien comprendre le fonctionnement d'un algorithme il est intéressant de disposer d'outils permettant de suivre sa progression. Dans cette section nous proposons d'étudier comment évolue la valeur de I_H au cours du déroulement de l'algorithme. Comme nous l'avons fait pour les algorithmes de référence, cela permettra d'identifier comment convergent les algorithmes vers l'ensemble final de solutions. Nous avons calculé la valeur I_H de l'archive des solutions non-dominées toutes les 100.000 itérations, durant l'exécution de l'algorithme. Précisons qu'une itération correspond à une configuration, c'est-à-dire au déplacement d'une station.

La Figure 4.17 traduit l'évolution de I_H pour PLS en bleu, NSGA-II en rouge et MEMO en vert. Comme nous l'avons déjà remarqué lors de l'analyse des résultats de référence, NSGA-II dépasse très rapidement PLS, après quelques minutes seulement. Mais ici nous voyons également l'apport de MEMO en comparaison à NSGA-II. Nous pouvons voir notamment qu'à partir de $20 \cdot 10^6$ itérations NSGA-II tend à stagner alors que MEMO continue à progresser.

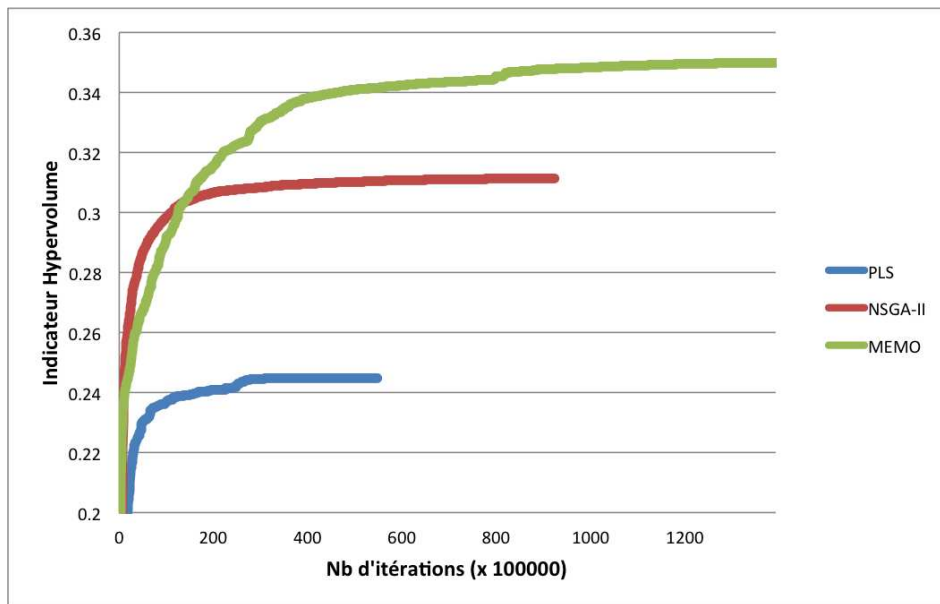


FIGURE 4.17 – Progression de PLS, NSGA-II et MEMO

De plus on peut constater une progression de MEMO plus continue que NSGA-II. En effet NSGA-II en bleu commence par être moins bon que MEMO, mais très rapidement passe légèrement au-dessus avant de repasser sous la courbe de MEMO.

Enfin il semble intéressant de remarquer que PLS converge en moyenne vers une valeur de 0,245 après $60 \cdot 10^6$ itérations. Ce niveau de $I_H = 0,245$ est atteint par MEMO après quelques itérations seulement, soit en quelques secondes contre plus

d'une heure pour PLS ce qui est une performance extrêmement intéressante.

4.6 De l'optimisation à l'aide à la décision multiobjectif

4.6.1 Contexte de l'aide à la décision

Lorsque l'on aborde un problème multiobjectif (MOP), trois phases peuvent être identifiées. Il y a tout d'abord la modélisation du problème, mathématiquement, avec la définition des fonctions objectifs. Ensuite, il y a la phase de recherche des solutions. C'est là que se concentre la grande majorité des travaux et publications scientifiques, à travers l'élaboration de métaheuristiques, capables de trouver une solution optimale ou plus souvent aujourd'hui un ensemble de solutions de meilleur compromis. Enfin, il y a la prise de décision qui doit conduire le décideur à faire le meilleur choix pour son problème.

On se rend compte que générer l'ensemble Pareto n'est pas suffisant. Lorsqu'un algorithme retourne 1.000 solutions ou plus, aussi bonnes soient-elles en termes de Pareto optimalité, la personne qui doit en retenir une peut rapidement être perdue. Choisir une solution est un processus complexe et dont le résultat est risqué.

Selon les approches, la prise en compte des préférences du décideur peut intervenir de différentes façons. On distingue ainsi généralement trois schémas : l'expression du choix *a priori*, l'expression du choix *a posteriori*, et enfin l'expression du choix *interactif*. Comme nous l'avons vu dans le chapitre 3, certaines approches de résolution des MOP intègrent directement les préférences du décideur. C'est le cas notamment lorsque le MOP est transformé en problème mono-objectif. Dans la méthode d'agrégation [Cohon 1978] par exemple, le décideur définit le vecteur de poids λ qui donne l'importance souhaitée à chaque objectif dans la fonction agrégée. Le choix est dans ce cas *a priori* et cela suppose non seulement que le décideur sait ce qu'il veut mais qu'en plus l'algorithme est capable de le fournir. Mais souvent dans les approches Pareto aucune préférence du décideur n'est intégrée. C'est ce qui fait la force de ces méthodes, mais c'est également ce qui peut poser certaines difficultés. En particulier un algorithme peut trouver de nombreuses solutions non-dominées, qui pourtant ne répondront peut-être pas aux attentes du décideur. Les préférences du décideur ne sont prises en compte qu'une fois l'ensemble des solutions généré. Le choix est exprimé *a posteriori*. Enfin, lorsque les préférences de l'utilisateur sont exprimées durant la recherche on parle de choix *interactif*.

Quelque soit le moment où intervient le décideur, la prise en compte de ses préférences pour l'aider à choisir une solution est un élément fondamental. Dans cette section nous proposons notre approche pour aborder l'aide à la décision. Pour plus de détails sur les approches classiques d'aide à la décision, nous renvoyons le

lecteur au différents ouvrages de Bernard Roy, et particulièrement [Roy 1985]. Nous présentons ici la plate-forme geoLogic que nous avons développée pour aiguiller le décideur dans sa compréhension d'un problème de mobilité et de déploiement d'infrastructures. Cet outil s'inscrit clairement dans la catégorie des approches dites *a posteriori*. En effet, l'algorithme de recherche que nous avons proposé s'attache à construire l'ensemble Pareto le plus exhaustif. Par contre le décideur n'intervient pas pour guider la construction de cet ensemble.

4.6.2 La plate-forme geoLogic

Elle constitue un socle commun à de nombreux projets. Nous l'avons développé parallèlement à cette thèse en vue de pouvoir y intégrer les différents projets à caractère géographique, en cours et à venir. La Figure 4.18 présente une vue d'ensemble de l'organisation de cette plate-forme.

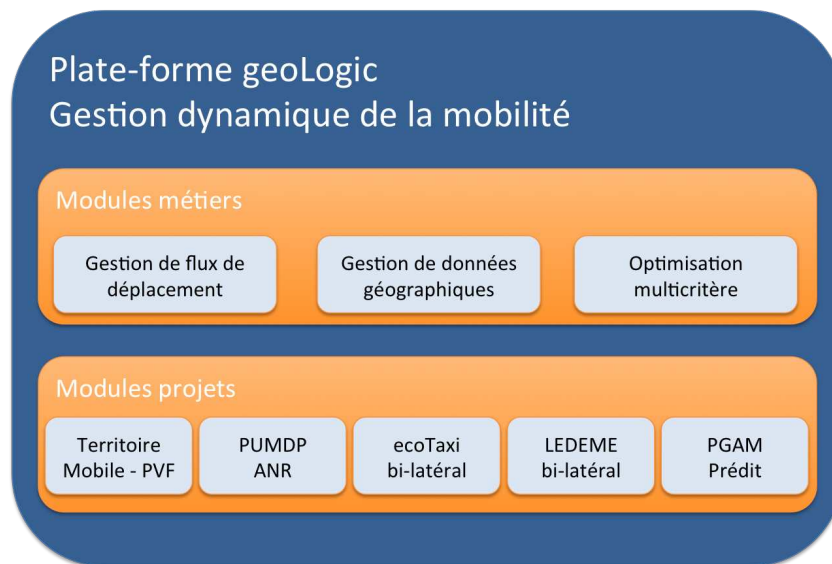


FIGURE 4.18 – Plate-forme geoLogic

Il s'agit d'une plate-forme de type SIG (Système d'Information Géographique) qui intègre, en plus d'un module de gestion des données géographiques, différents modules orientés métier. Parmi les éléments de base on retrouve le maillage des données géographiques, la prise en compte des flux de déplacements entre zones géographiques, entre mailles plus fines, ou encore suivant les réseaux routiers. Toutes les données intégrées à cette plate-forme peuvent être considérées comme statiques ou dynamiques, avec pour les données dynamiques la possibilité de définir le pas de temps. La Figure 4.19 illustre quelques vues issues de cet outil. De gauche à droite, cette figure montre tout d'abord la description vectorielle du territoire sur fond de carte. On distingue notamment des objets de type polygone (en bleu) qui localisent

les habitations, ou encore des objets de type ligne (polyline) pour représenter les axes routiers. Ensuite nous voyons le résultat obtenu après application d'une grille de maillage sur les données vectorielles. Chaque couleur représente une catégorie de sol, avec en bleu les habitations, en orange les commerces, etc. L'illustration suivante montre les déplacements sortant de la zone sélectionnée à un moment donné. L'orientation des flèches indique les destinations et leur épaisseur le nombre de personnes faisant ce déplacement. Enfin, la vue de droite représente sous forme de "vaisseaux sanguins" les axes de rue avec une intensité sur la mobilité observée. Par effet de transparence les axes les plus empruntés ressortent plus intensément.

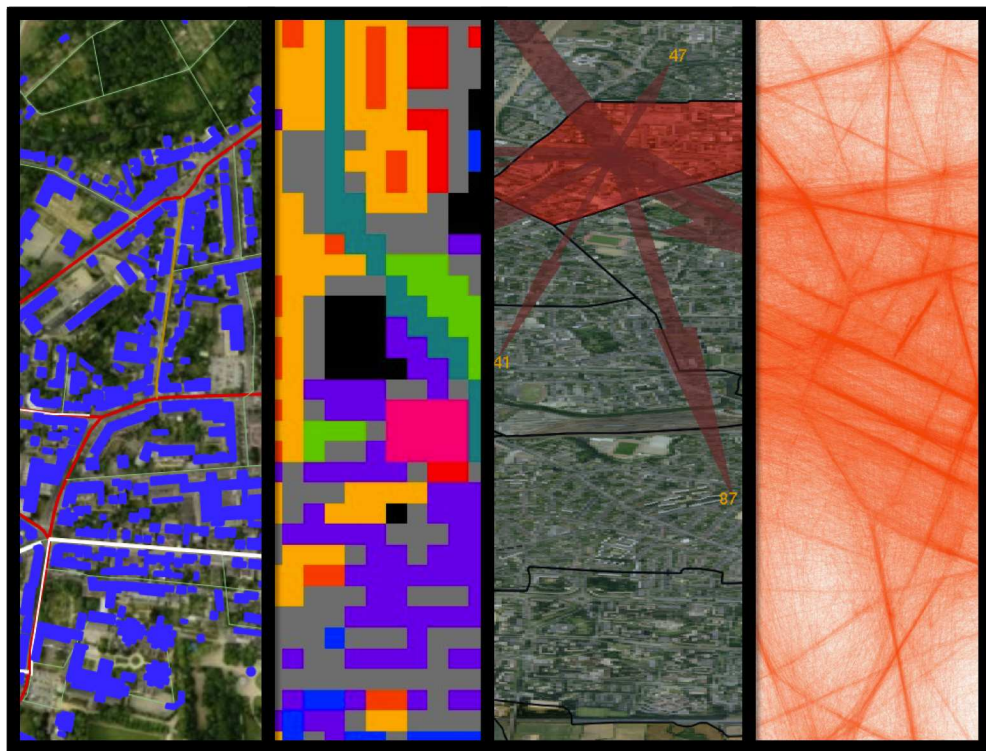


FIGURE 4.19 – Vues issues de la plate-forme geoLogic

La plate-forme proposée est réalisée en C++ standard, et s'appuie particulièrement sur la STL (Standard Template Library). Elle peut ainsi être déployée sur les OS existants tels que Linux, Mac OS et Windows. Pour permettre l'affichage des calques géographiques de façon particulièrement efficace, en termes de fluidité et de quantité de données, cette partie est déportée sur la carte graphique à travers l'utilisation de la bibliothèque OpenGL (Open Graphics Library).

Enfin, l'intégration de modules projets est facilitée grâce au recours à un canevas générique, la classe virtuelle *geoModule*, qui permet de renseigner de façon aisée les trois composants d'un module : la partie contrôle, la partie calcul et la partie visualisation.

4.6.3 Aide à la décision multiobjectif

Lorsqu'une décision doit être prise, en matière de création d'infrastructures ou de localisation de services, elle ne peut se baser uniquement sur un niveau de performance des solutions. En effet, la prise en compte des réalités du terrain, des orientations politiques visées, et de l'expertise du décideur, sont autant de facteurs agissant sur la décision finale. Dès lors, la mise en regard de l'espace de décision (implantation des solutions) et de l'espace des critères (évaluation des solutions) s'avère être essentielle. L'important dans le problème de l'auto-partage que nous optimisons, est de montrer au décideur comment une façon d'implanter les stations se situe en termes d'évaluation, et réciproquement comment un niveau de performance se situe en termes de solutions correspondantes.

Le choix que nous avons fait est donc de proposer au décideur une vue globale de toutes les solutions non-dominées, afin de lui laisser le soin de les analyser et d'en extraire la solution de son choix. Encore faut-il le guider au mieux dans cette tâche. Pour répondre à ce besoin nous avons réalisé un module spécifique aux problèmes d'optimisation multiobjectif à caractère géographique. Il permet à son utilisateur d'appliquer les différents algorithmes dédiés aux MOP présentés dans ce chapitre, tels que PLS [Paquete 2004], FLSMO [Moalic 2013a], NSGA-II [Deb 2002] et MEMO [Moalic 2013b], ainsi que des approches Tabou [Moalic 2012] ou encore des variantes de PLS. De plus il permet également de visualiser les solutions proposées simultanément dans les deux espaces, l'espace de décision qui est ici l'environnement cartographique, et l'espace des critères qui est l'évaluation des fonctions objectif.

L'outil permet cette mise en correspondance, en plus de la fonctionnalité permettant de filtrer les solutions. La Figure 4.20 propose un aperçu de l'interface de cet outil. La fenêtre principale permet de visualiser l'environnement cartographique, où sont affichés les calques souhaités tels que la carte, la description vectorielle du territoire (*shapefile*), la description maillée du territoire, les flux de déplacements, et l'implantation des stations pour les solutions calculées par l'algorithme d'optimisation. Il s'agit donc de l'espace de décision. La fenêtre du bas permet quant à elle d'afficher les solutions Pareto optimales dans l'espace des critères. Les solutions de l'espace à trois dimensions (trois fonctions objectif) sont projetées sur les axes considérés deux à deux ($f_1 - f_2$, $f_1 - f_3$, $f_2 - f_3$). Enfin, le panel de droite permet à l'utilisateur de choisir les calques à afficher, les différents paramètres pour l'optimisation, etc. Il s'agit du panel de contrôle.

Le nuage de points dans l'espace objectif permet de disposer de la vue d'ensemble des solutions non dominées trouvées par le module d'optimisation. Parmi toutes ces solutions, la solution sélectionnée est représentée en rouge. Ceci permet d'une part de déterminer où elle se situe dans l'espace objectif vis à vis de chaque critère, mais aussi d'identifier à quelle configuration physique correspond cette solution en référence à la localisation de stations. Dans la Figure 4.20 chaque disque rouge

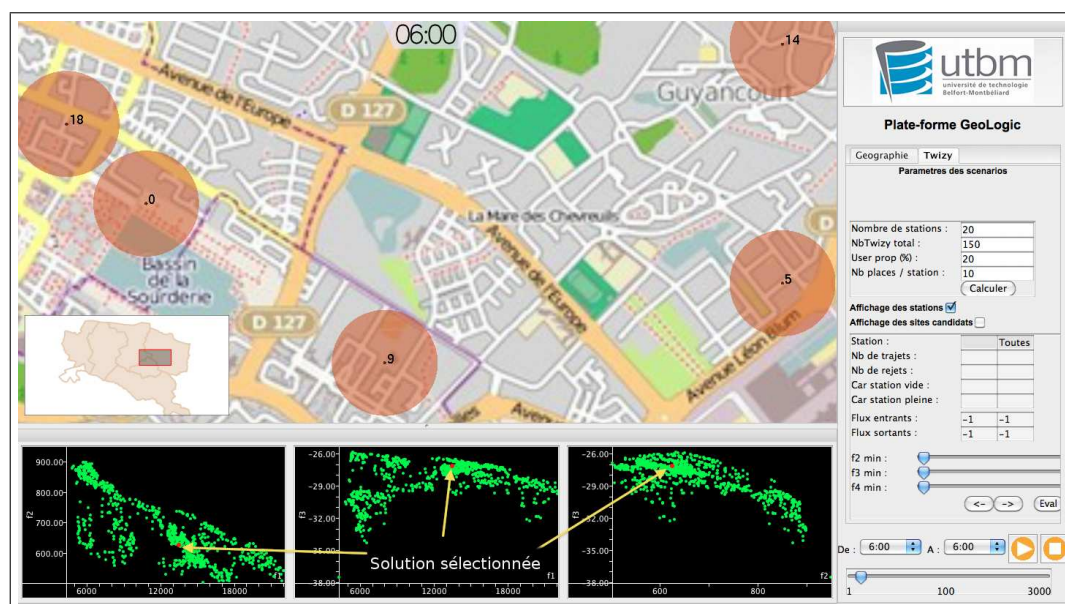


FIGURE 4.20 – Outil d’aide à la décision de geoLogic

correspond à l’implantation d’une station et les trois fenêtres du dessous montrent la distribution des solutions dans l’espace des trois critères f_1 , f_2 et f_3 pris deux à deux.

Nous l’avons indiqué, parmi toutes les solutions Pareto optimales il est vraisemblable que de nombreuses solutions ne soient pas acceptables : toutes les régions de l’espace des critères ne sont pas intéressantes pour le décideur. Pour l’aider dans son choix, nous avons mis en place des filtres qui lui permettent de cibler certaines parties de l’espace des critères. Cela peut être vu comme une approche ε -contrainte *a posteriori*, avec l’avantage de proposer au décideur une vue d’ensemble des solutions, et ainsi utiliser ces paramètres ε de façon interactive.

Sur la Figure 4.21 les filtres sont appliqués. Ils consistent donc à fixer des valeurs minimales pour f_1 , f_2 et f_3 dans l’espace des critères. Les solutions désactivées (exclues par les filtres) par l’utilisateur apparaissent en bleu et les solutions activées (qui respectent les filtres) apparaissent en vert. Ainsi le décideur peut se focaliser sur les solutions satisfaisant les valeurs seuils qu’il fixe, et identifier les emplacements correspondants sur l’espace géographique. La solution entourée d’un cercle blanc dans l’espace des critères correspond à l’implantation présentée sur la carte.

Réciproquement il est possible de choisir un site de l’espace de décision pour visualiser comment se répartissent les solutions ayant une station sur ce site dans l’espace objectif. La Figure 4.22 présente les solutions dans l’espace des critères ayant une station sur le site du cercle rouge. A gauche, on présente les solutions

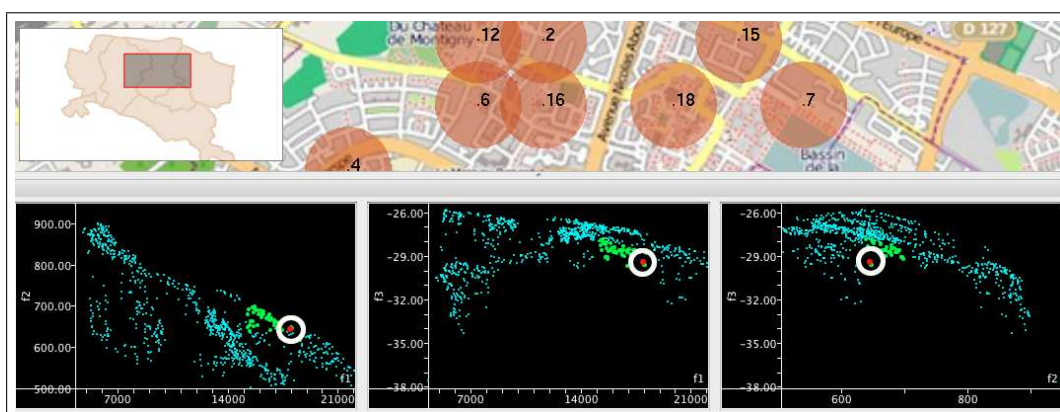


FIGURE 4.21 – Filtres sur l'espace des critères

dans l'espace f_1 (abscisse), f_2 (ordonnées); au centre dans l'espace f_1 (abscisse), f_3 (ordonnées) et à droite dans l'espace f_2 (abscisse), f_3 (ordonnées). L'intensité du violet de chaque site traduit le nombre de solutions ayant une station à cet emplacement. On voit notamment que le site sélectionné est utilisé dans beaucoup de solutions efficaces pour f_1 . En effet, dans la vue de gauche de l'espace des critères, les solutions sélectionnées (en vert) sont regroupées essentiellement vers la droite (bonne valeur de f_1). Mais on distingue également quelques bonnes configurations pour f_2 . Elles se retrouvent toujours dans la vue de gauche vers le haut (f_2 est sur l'axe des ordonnées). Par contre le fait de choisir ce site pour y mettre une station ne fournit pas de bons résultats pour f_3 .

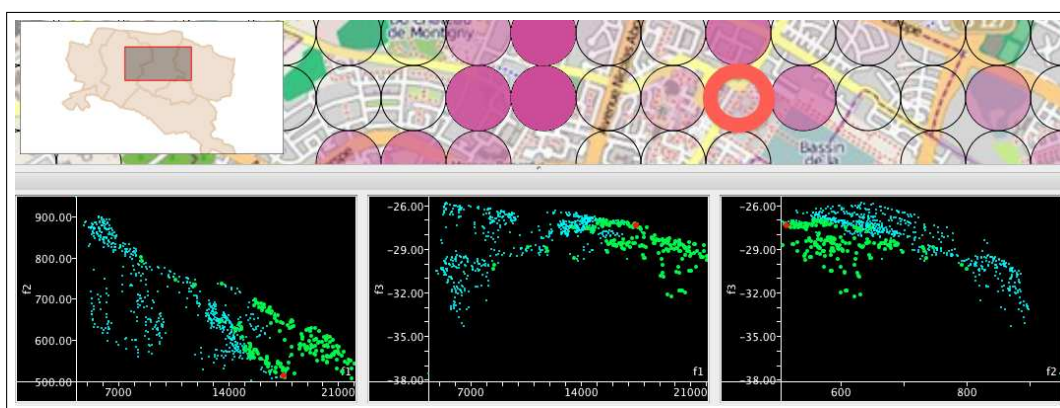


FIGURE 4.22 – Filtres sur l'espace de décision

Ainsi si un site intéresse particulièrement le décideur, il peut le sélectionner et étudier s'il appartient à des solutions dont l'évaluation peut lui convenir. A l'inverse, cette démarche peut permettre de justifier l'absence de station sur certains sites, en

montrant qu'il s'agit de sites n'appartenant à aucune solution satisfaisante.

Il s'agit bien ici d'un outil permettant de comprendre le problème traité et de guider la prise de décision quant à la solution à déployer.

4.6.4 Les projets s'appuyant sur la plate-forme geoLogic

Au cours de ces dernières années, parallèlement aux travaux présentés jusqu'à maintenant, nous avons mis en oeuvre la plate-forme geoLogic dans divers projets, toujours pour des problèmes d'aide à la décision en matière de transport. Ici l'aide à la décision ne s'entend pas forcément comme une phase d'un processus d'optimisation. On étend la définition d'aide à la décision aux outils permettant à un utilisateur de mieux comprendre son problème, où se situent les besoins, en vue de guider ses choix en matière de déploiement d'infrastructures sans forcément faire appel à un algorithme d'optimisation ensuite.

4.6.4.1 Le projet Territoire Mobile

Collaborations : SMTC Belfort (société mixte d'exploitation du réseau de bus), société Orange, Agence d'Urbanisme, collectivités territoriales du Territoire de Belfort (2006-2009)

Ce projet labellisé par le pôle de compétitivité Véhicules du Futur a été présenté dans le chapitre 2. C'est en effet dans le cadre de ce projet que nous avons pu expérimenter et mettre en oeuvre nos modèles de présence s'appuyant sur les téléphones portables. Ce projet, novateur quant à l'utilisation de la téléphonie mobile pour comprendre comment évolue la population sur le territoire au cours de la journée, était destiné à améliorer l'offre en transport en commun.

La mise en relation du fonctionnement du réseau d'alors, et de la localisation de la population issue de notre modèle, devait faire ressortir les éventuels dysfonctionnements du service. La Figure 4.23 en propose une illustration. Elle propose une vue d'ensemble du Territoire de Belfort avec la répartition de la population construite à partir de la téléphonie mobile. Le zoom sur le centre ville révèle la mise en correspondance entre la localisation des personnes et le réseau de bus. Les disques de chaque arrêt traduisent le nombre de montées ayant eu lieu dans les bus à ce moment de la journée. Cet outil a notamment permis de se rendre compte de zones ne se trouvant pas à proximité d'un arrêt de bus mais où de nombreuses personnes se rendaient. L'objet n'était donc pas d'optimiser le réseau de transport mais de permettre à l'exploitant d'analyser par lui-même ce qui pouvait être amélioré.

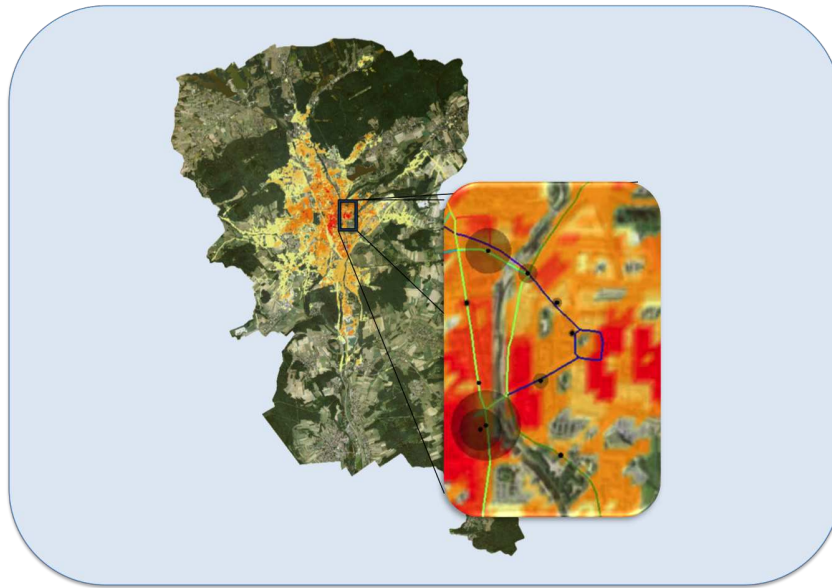


FIGURE 4.23 – Territoire Mobile : analyse d'un réseau de bus

4.6.4.2 Le projet EcoTaxi

Collaborations : Renault Technocentre, Taxis G7, Taxis Gescop, assureur MMA (2009-2010)

Ce projet mené avec le Technocentre de Renault (centre de recherche et de conception Renault) avait pour vocation de permettre une meilleure compréhension de la façon dont les taxis se déplacent à Paris, et par extension quels sont les besoins en mobilité. Ce travail s'est inscrit dans la politique actuelle du constructeur automobile visant à développer les véhicules électriques. L'objectif de ce projet, outre de révéler comment les taxis se déplacent tout au long de la journée, était de proposer un module d'optimisation pour la localisation et le dimensionnement de stations d'échange automatique de batteries pour des taxis électriques.

L'enjeu ici est de placer les stations de façon à ce qu'un maximum de taxis puissent être convertis en véhicules électriques sans rencontrer de problème d'autonomie et donc de panne.

Les données de mobilité utilisées dans ce cadre provenaient directement des sociétés de taxis. Des relevés de positions de plus de 5000 taxis ont été fait durant trois semaines à partir de leurs GPS embarqués. La Figure 4.24 montre sous forme de carte de chaleur où se situent les taxis ayant besoin de se recharger selon les hypothèses d'autonomie définies par l'utilisateur. Les zones rouges en particulier indiquent les lieux où les besoins sont les plus importants.

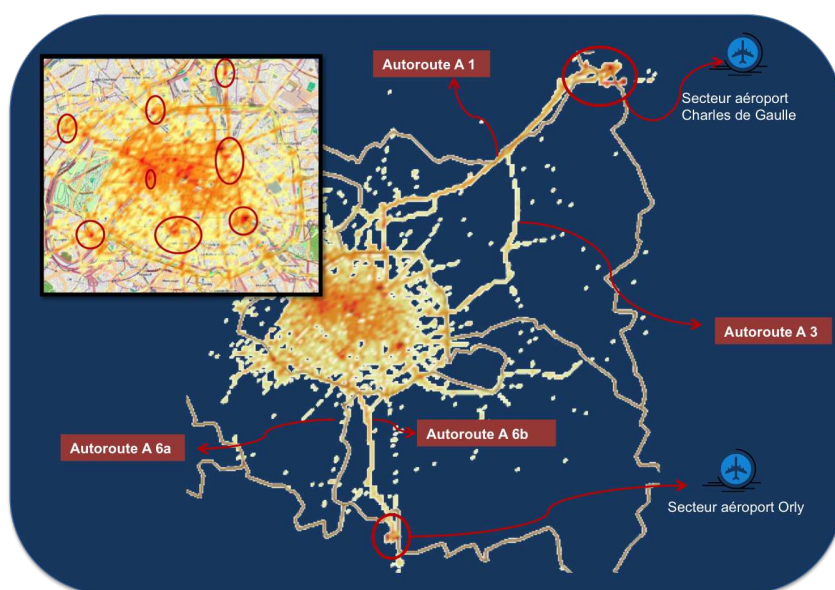


FIGURE 4.24 – ecoTaxi : optimisation d’infrastructures de stations d’échange automatique de batteries

L’approche utilisée pour réaliser l’optimisation est une approche constructive qui fonctionne avec un algorithme de descente. Le logiciel, en s’appuyant sur les relevés de trajectoires et sur les hypothèses de fonctionnement des véhicules électriques, propose successivement l’emplacement optimal pour placer la prochaine station en tenant compte des stations placées aux étapes précédentes et des taxis en besoin d’échange de batterie. Libre au décideur de suivre la recommandation de l’outil ou de la positionner ailleurs. En résultat, le logiciel fournit une analyse complète des performances estimées avec le réseau de stations d’échange proposé : nombre de taxis venant se recharger par plage horaire et par station, nombre de taxis pouvant passer en électrique sans rencontrer de problème, nombre de taxis en panne d’autonomie au cours de leurs courses, etc.

4.6.4.3 Le projet PUMDP

Collaborations : Bureau sociologique du transport Chronos, La FING, Club des Villes et Territoires Cyclables, cabinet d’avocats 11.100.34, Société Keolis Groupe, Communautés d’Agglomérations de Rennes et Angers, ville de Lorient (2010-2012).

Le projet ANR PUMDP (Partages, Usages et Modélisations de la Donnée Publique) était également labellisé par le pôle de compétitivité Véhicule du Futur. L’objectif de ce projet était de fournir des outils de diagnostic, d’évaluation et de planification urbaine. Le cas d’étude retenu concernait les modes de déplacements dits actifs, la marche et le vélo. Le logiciel développé dans le cadre de ce projet est un outil d’aide

à la décision qui ne propose pas directement de solution optimale à mettre en oeuvre mais qui présente le niveau d'adéquation entre les déplacements actifs (la demande) et les infrastructures disponibles tels que les aménagements cyclables (l'offre).

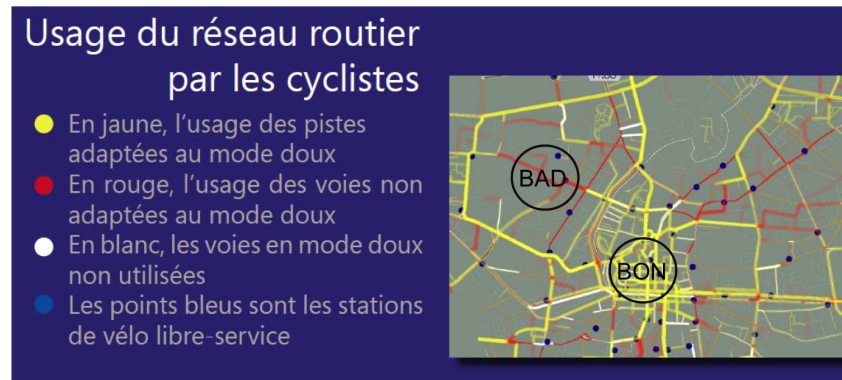


FIGURE 4.25 – PUMDP : adéquation entre offre et demande pour une mobilité à vélo

La Figure 4.25 montre ainsi en jaune, les réseaux adaptés aux vélos et empruntés, en blanc, les tronçons adaptés aux vélos mais peu empruntés, et surtout en rouge, les axes non aménagés qui sont empruntés par des vélos. Cet outil permet notamment d'attirer l'attention des collectivités concernées sur les axes qu'il faudrait aménager en priorité.

Les données utilisées pour déterminer la façon dont les cyclistes se déplacent sont de différentes natures. Elles vont d'enquêtes réalisées sur le terrain à des relevés GPS effectués par des volontaires, tout en passant par des comptages faits au niveau des stations de vélo en libre service (lieu et moment de prise d'un vélo, lieu et moment de dépose du vélo).

4.6.4.4 Les projets à venir

Outre les projets précédents, deux nouveaux projets s'inscrivent dans la continuité du travail réalisé durant cette thèse et s'appuieront sur la plate-forme geoLogic : SYSMO, un projet Investissement d'Avenir de l'ADEME (appel à projets 2011), et NORM-ATIS, un projet de l'ANR (appel à projet 2013). Leur démarrage est prévu pour janvier 2014.

Le projet Investissement d'Avenir traite de la mobilité et notamment de l'utilisation de la téléphonie mobile pour modéliser les déplacements. En ce sens il s'agit d'une suite directe du travail présenté dans le chapitre 2. Alors que dans le cas traité dans cette thèse les informations téléphoniques ne permettent que de construire des densités de présence dynamiques, avec ce nouveau projet les données téléphoniques permettront de construire des itinéraires. Ce niveau de finesse supplémentaire doit

nous inciter à faire évoluer les modèles présentés ici, afin de reconstituer des matrices de déplacements origines/destinations sur un territoire, et pour caractériser ces déplacements, en termes de vitesse par exemple.

Par ailleurs, ce projet comporte un volet consacré aux nouveaux usages de mobilité et à l'optimisation dans les transports. L'accent est mis sur la construction d'une offre de transport multimodale optimisée, appuyée sur l'articulation entre les différents modes disponibles tels que les transports en commun, l'auto-partage, les véhicules individuels, etc. Il s'agit donc non seulement d'une extension du travail présenté dans le chapitre 2 mais également du travail présenté dans ce chapitre.

Ce projet, labellisé par les pôles de compétitivité Systematic, Mov'eo et Advancity, va aboutir notamment à la mise au point d'un logiciel de planification et de gouvernance de la mobilité.

Le deuxième projet, financé par l'ANR et soutenu par les pôles de compétitivité Mov'eo et Véhicule du Futur, s'attache à établir des normes sur les données impliquées dans la modélisation de la mobilité et de l'information trafic. Il s'agit donc d'un travail de fond sur les façons de modéliser les déplacements, de manière à standardiser les données et les modèles entre les différents territoires, en France dans un premier temps, mais avec l'objectif clair d'établir de nouveaux standards européens. Ce projet s'inscrit également dans la continuité du travail réalisé dans cette thèse.

Notons que ces deux nouveaux projets porteront sur des volumes de données beaucoup plus importants que ceux que nous avons traités jusqu'alors, permettant une validation des modèles et algorithmes présentés dans cette thèse à une échelle bien plus large.

4.6.4.5 Un transfert technologique en cours

Les travaux que nous avons réalisés dans le cadre de cette thèse, et plus généralement autour de la plate-forme geoLogic, nous amènent actuellement à un transfert technologique vers une PME, la société Web Geo Services, qui apporte la structure et les compétences permettant de les valoriser. En effet, la téléphonie comme nouveau marqueur de mobilité ouvre certes des portes mais nécessite beaucoup de traitements appropriés pour leur donner du sens. Les modèles, qui incluent données de mobilité (téléphoniques, enquêtes ou autres) mais aussi les composantes géographiques et socio-économiques pour construire une information pertinente sur les déplacements de la population, constituent une plus value certaine pour les collectivités.

4.7 Synthèse

Dans ce chapitre nous avons présenté le processus complet permettant de traiter un problème d'optimisation multiobjectif réel (MOP), de la définition du problème, à la conception d'une heuristique permettant de le résoudre en passant par la modélisation du problème et des fonctions objectifs associées. Les problèmes combinatoires multiobjectifs rencontrés dans la littérature peuvent être abordés sous plusieurs angles, en fonction du besoin de précision, du temps de développement nécessaire ou encore du temps de calcul. Parmi les différentes approches permettant d'appréhender un MOP, les algorithmes mémétiques ont montré leur efficacité sur de nombreux problèmes. Ces approches hybrides, incorporant une recherche locale à l'intérieur d'un algorithme génétique, ont par ailleurs été souvent utilisées pour résoudre des problèmes mono-objectifs et fournissent encore aujourd'hui les meilleures performances sur de nombreux problèmes. L'étude de ces approches nous a permis de définir notre propre algorithme mémétique, MEMO (Memetic Evolutionary MultiObjective).

L'algorithme MEMO est une hybridation de NSGA-II, algorithme certainement le plus répandu parmi les approches génétiques pour résoudre un MOP, avec une descente simple du type Hill Climbing. Alors que NSGA-II est un algorithme multiobjectif, et donc directement prévu pour fonctionner sur ce type de problème, l'algorithme de descente simple est adapté aux problèmes mono-objectifs. Pour cause, un tel algorithme repose sur le fait qu'une transition vers un voisin n'est acceptable que si elle permet d'améliorer la fitness. Or comme nous l'avons montré, une telle valeur de fitness ne peut être déterminée en contexte multiobjectif. Les solutions d'un MOP ont une évaluation distincte pour chaque objectif, et très souvent ne sont pas comparables entre elles. Dans les seuls cas où tous les objectifs de s_1 sont soit meilleurs, soit moins bons, que ceux de s_2 , une relation d'ordre existe entre s_1 et s_2 . Mais dans le cas largement majoritaire où certains objectifs seulement de s_1 sont meilleurs que ceux de s_2 et les autres moins bons, il n'est plus possible de déterminer qui de s_1 et s_2 est la meilleure. Il n'est donc pas possible d'appliquer un algorithme de descente directement.

Nous avons proposé alors de procéder à une agrégation des fonctions objectifs pour rassembler les différents objectifs en une valeur de fitness unique lors de la recherche locale. La façon d'agréger les objectifs est déterminante dans le déroulement de l'algorithme. Il est possible de simplement faire la somme des objectifs, après s'être assuré de la normalisation des objectifs pour les ramener à un domaine de valeurs comparables. Ainsi on considère que tous les objectifs sont à favoriser de la même façon. Il est par ailleurs possible d'attribuer un poids à chaque objectif selon l'importance qu'on souhaite lui donner. On tend alors vers des approches de type lexicographique où l'on est capable d'ordonner prioritairement les objectifs. Dans notre approche nous proposons de générer une agrégation selon une pondération aléatoire, recalculée pour chaque recherche locale. Cette approche permet de ne pas

favoriser une direction de recherche (direction définie par la pondération), mais au contraire d'explorer au maximum l'espace de recherche.

Pour appliquer l'algorithme proposé, nous avons présenté le problème particulier de la localisation optimale des stations pour un service d'auto-partage. Alors que jusqu'à présent les approches rencontrées dans la littérature pour résoudre ce problème étaient des approches exactes, qui ne permettaient de résoudre qu'une version simplifiée du problème (maillage du territoire en maille d'un kilomètre carré), l'introduction d'une approche heuristique a permis d'offrir de nouvelles perspectives. Nous avons proposé une modélisation du problème sous forme de MOP qui permet de construire des configurations diverses du service, où chaque configuration est une façon de localiser n stations. Les solutions retournées par l'algorithme sont les solutions rencontrées appartenant à l'ensemble Pareto, non pas réel car cet ensemble n'est pas connu, mais l'ensemble Pareto composé de toutes les solutions non-dominées entre elles que peut trouver l'algorithme. La constitution de cet ensemble Pareto suppose la définition des fonctions objectifs du problème. Nous avons proposé et formalisé trois fonctions permettant de trouver les configurations qui (1) s'adressent au maximum de flux de déplacements, (2) permettent une auto-régulation du service maximisant l'équilibre entre flux entrant / flux sortant à une même plage horaire et (3) favorisent un fonctionnement tout au long de la journée et non uniquement aux heures de pointe.

L'application d'une heuristique à un problème donné nécessite souvent de définir certains paramètres. Il peut s'agir de paramètres numériques, mais cela peut également être des opérateurs à part entière. Afin d'obtenir le meilleur fonctionnement de l'algorithme MEMO nous avons procédé à une série de tests en faisant varier la taille de la population et l'opérateur de croisement de l'algorithme mémétique. Nous avons montré à travers ces tests qu'une population de 20 individus couplée à un croisement aléatoire fournissait les meilleurs résultats. Le niveau de performance d'un algorithme ne peut être défini qu'en le comparant à d'autres approches. C'est pourquoi nous avons proposé de recourir à trois approches, dont deux reconnues pour être performantes en multiobjectif, afin de confronter notre algorithme. Nous avons choisi une approche locale, PLS, et une approche génétique, NSGA-II. Une deuxième recherche locale, FLSMO, est un algorithme que nous avons proposé et également testé dans ce chapitre.

L'algorithme mémétique MEMO que nous proposons s'est révélé être de loin le plus performant. Plus performant que la recherche locale PLS sans grande surprise, mais également significativement plus performant que NSGA-II qui a été paramétré pour fournir les meilleurs résultats possibles. La comparaison entre les algorithmes a nécessité l'introduction des indicateurs de distance $I_{\epsilon+}$, de contribution I_C et surtout l'indicateur hypervolume I_H qui est le meilleur reflet de la qualité d'un ensemble de solutions. Ces indicateurs nous ont permis de quantifier le gain apporté par l'algorithme mémétique sur les autres approches testées.

L'ensemble des solutions Pareto optimales étant construit, nous nous sommes attachés ensuite à guider le décideur vers la solution à retenir, en fonction de ses besoins et de son expertise. Cette étape, essentielle dans la résolution d'un MOP, s'inscrit dans le domaine de l'aide à la décision. L'objet ici n'était pas de proposer un algorithme capable de déterminer la meilleure solution parmi l'ensemble des solutions non-dominées, mais plutôt de fournir au décideur des outils efficaces qui lui permettent de mieux comprendre son problème et de choisir par lui-même la solution la plus appropriée.

Pour répondre à ce besoin, et plus généralement aux problèmes posés lorsque l'on traite de mobilité, nous avons développé la plate-forme geoLogic. Elle permet, à travers la gestion d'un environnement cartographique et des panels de contrôle, de présenter simultanément les résultats issus des heuristiques dans les deux espaces : l'espace de décision (localisation des stations) et l'espace des critères (front de Pareto). L'ajout d'outils tels que des filtres sur les données permet au décideur de se focaliser sur les régions de l'espace de recherche qui l'intéressent.

Enfin, nous avons présenté quelques projets d'aide à la décision que nous avons développé sur la plate-forme geoLogic ainsi que des projets à venir.

Conclusion générale et perspectives

Ce mémoire de thèse traite de deux aspects essentiels dans le domaine du transport : l'identification des besoins en mobilité, à savoir la demande, et l'optimisation des infrastructures permettant de répondre à cette demande, à savoir l'offre. Ces deux volets sont abordés sous un regard novateur. La construction de la demande tout d'abord, où un modèle est proposé pour identifier comment au cours de la journée la population se répartie sur un territoire. Ce modèle original repose sur des données issues du fonctionnement d'un réseau cellulaire et sur une caractérisation spatio-temporelle du sursol. La construction de l'offre ensuite, où le cas d'étude retenu est la localisation de stations pour un service d'auto-partage. Une approche heuristique est proposée pour résoudre ce problème combinatoire, modélisé sous la forme d'un problème multiobjectif.

Principales contributions

La modélisation de la mobilité Nous avons dans une première partie de ce mémoire abordé le thème de la modélisation de la mobilité des personnes. Pour permettre une analyse pertinente de la mobilité, nous avons proposé une taxinomie des modèles existants en fonction de leur niveau de réalisme. Trois familles de modèle ont été identifiées : les modèles aléatoires, les modèles déterministes, et les modèles hybrides. Les modèles hybrides se sont révélés être les plus à même de modéliser finement la mobilité, suffisamment finement pour permettre des applications en ingénierie du transport. Ils reposent sur des données d'enquêtes ou de comptages routiers pour identifier la demande ; mais ils intègrent également des processus stochastiques permettant de refléter les fluctuations pouvant survenir d'un jour à l'autre. Ces modèles hybrides, tels qu'ils sont utilisés par les experts de la mobilité, sont eux-mêmes répartis en trois grandes classes : les modèles agrégés, les modèles désagrégés et les modèles activité-centré. Cette dernière classe qui s'appuie sur la notion d'activité, permet de rendre compte du chainage des déplacements. Les travaux les plus récents dans ce domaine ont montré que les modèles de transport du futur seront de cette nature, certes complexe à modéliser mais traduisant le plus justement les besoins en mobilité.

Pour compléter les approches existantes en matière de modélisation de la mobilité, nous avons vu que la téléphonie mobile pouvait être utilisée comme un marqueur de présence sur le territoire. Nous avons proposé un modèle de distribution spatio-temporel des mobiles basé sur les données téléphoniques et sur la description du

sursol. La prise en compte du sursol est l'atout majeur du modèle proposé par rapport aux approches qui existaient jusqu'alors. Les quelques rares travaux s'appuyant sur la téléphonie mobile pour marquer la présence des personnes sur un territoire ne permettent pas un niveau de précision au-delà de la cellule radio. L'association d'un vecteur de poids d'attraction spatio-temporel à un découpage fin du territoire, nous a permis d'intégrer la notion d'activité à notre modèle. A ce titre il s'inscrit dans les modèles d'avenir identifiés par les experts du transport.

La modélisation du problème de l'auto-partage Modéliser la mobilité n'a de sens que si cela permet d'améliorer le fonctionnement de différents services et d'aider à prendre des décisions en matière de déploiement d'infrastructures notamment. Dans la deuxième partie de ce mémoire nous nous sommes intéressés au service de l'auto-partage électrique. Pour choisir l'emplacement optimal de chaque station où seront empruntés et rendus les véhicules, il apparaît primordial d'identifier les critères intervenant sur le bon fonctionnement du service. Nous avons à cette fin identifié trois fonctions objectif, qui traduisent trois caractéristiques importantes. Tout d'abord les stations doivent être localisées de façon à satisfaire un maximum de demande de déplacements. Un deuxième objectif est de maximiser l'autonomie du service en limitant les situations de saturation, telles que des stations sans véhicule, ou au contraire sans place disponible pour déposer un véhicule. Enfin, un tel service doit permettre de satisfaire des demandes régulièrement réparties sur la journée, pas uniquement aux périodes de pointe. Cette façon novatrice d'appréhender ce problème a été expérimentée et validée dans un contexte réel.

MEMO, une métaheuristique multiobjectif hybride Enfin, notre troisième contribution concerne l'optimisation combinatoire appliquée aux problèmes multiobjectifs (MOP). Afin de résoudre des problèmes d'optimisation complexes, nous avons vu que les heuristiques apportaient des réponses particulièrement efficaces. En particulier, trois classes de méthodes ont été présentées : la première regroupe les approches par transformation du MOP en un problème mono-objectif, la deuxième concerne les approches non-Pareto, et enfin la troisième est composée des approches Pareto. Une quatrième classe, qui fournit souvent les meilleures performances, repose sur l'hybridation de composantes issues des trois classes principales. L'algorithme MEMO (Memetic Evolutionary MultiObjective) proposé dans ce mémoire fait partie de ces approches hybrides. Il repose sur l'algorithme génétique NSGA-II, qui fait référence dans le domaine de la résolution des MOP, dans lequel la mutation est remplacée par une recherche locale. L'algorithme mémétique ainsi proposé a été expérimenté sur le problème de localisation de stations pour le service d'auto-partage présenté précédemment. Il a montré de très bonnes performances en comparaison aux approches classiques NSGA-II, PLS et FLSMO.

Perspectives et travaux en cours

Chaque champ exploré dans ce mémoire laisse entrevoir de nombreuses perspectives pour améliorer les performances obtenues ou encore pour résoudre d'autres types de problème. Nous proposons quelques aspects qui nous semblent particulièrement intéressant à approfondir et que nous avons parfois déjà commencé à étudier. Ces perspectives sont présentées autour de trois axes correspondant aux trois contributions présentées.

La modélisation de la mobilité basée sur la téléphonie mobile La téléphonie mobile peut être vue comme un formidable réseau de capteurs, ne nécessitant aucun investissement d'infrastructure complémentaire à ce qui existe déjà. Les données générées et enregistrées à chaque communication contiennent une richesse d'informations sous exploitées à l'heure actuelle. Le défi qui s'inscrit dans la lignée des travaux que nous avons présentés concerne la sémantique que l'on sera capable d'attribuer à ces données.

Le modèle de répartition spatio-temporel des mobiles que nous avons proposé s'attache à répartir statistiquement un nombre de mobiles dans une cellule radio. Ce nombre fournit une information de présence, mais ne permet pas de suivre le cheminement d'un mobile individuellement. Pour cette raison, l'approche proposée s'apparente à ce qui se fait dans les modèles de mobilité de type agrégé. La distribution effectuée par le modèle sur le territoire en fonction des caractéristiques du sursol introduit cependant la notation d'activité associée. Il semble très prometteur de pousser plus loin ce concept d'activité mais il nécessite de nouveaux types de données issues des réseaux cellulaires. Depuis l'année dernière en France, de nouvelles données téléphoniques permettent d'identifier chaque mobile individuellement et anonymement. Cela ne remet pas en cause le modèle que nous proposons, au contraire cela va permettre de l'enrichir en intégrant plus finement la notion d'activité. Le fait de pouvoir identifier un chainage de déplacements doit en effet permettre de préciser la caractérisation du profil de mobilité de chaque mobile. Cette perspective fait partie intégrante des projets SYSMO (Investissement d'Avenir de l'ADEME) et NORM-ATIS (ANR) présentés en section 4.6.4.4.

La modélisation du problème de l'auto-partage Le problème de localisation de stations pour un service d'auto-partage électrique est abordé sous un nouvel angle dans ce mémoire. Comme nous l'avons décrit, outre le fait de positionner les stations pour répondre à la plus forte demande, un des points essentiels est d'éviter les situations de blocage rencontrées lorsqu'une station ne contient plus de véhicule ou au contraire lorsque plus aucune place n'est disponible. Pour répondre à cette problématique nous avons défini des fonctions objectif qui tendent à limiter ces phénomènes de blocage. Mais l'introduction d'un service de jockeyage pour

redéployer les véhicules entre stations bloquées reste indispensable et constitue une perspective intéressante à prendre en compte. Il s'agit d'introduire cet aspect dans la modélisation mathématique du problème en termes de coût de personnel et de gain sur l'utilisation des véhicules.

Evidemment les perspectives sont nombreuses. Parmi celles à envisager à court terme, nous pensons aussi à une généralisation du modèle afin de l'étendre au cas des vélos en libre service. En effet, de plus en plus de villes à travers le monde (la France est précurseur dans le domaine) envisagent ce service sans recourir à de tels modèles qui permettraient pourtant de rationaliser les investissements. Il s'agit là certainement de quelque chose à approfondir.

Une métaheuristique multiobjectif hybride avec gestion de l'incertitude

L'algorithme MEMO tel que nous l'avons appliqué pour localiser des stations d'auto-partage s'appuie sur une matrice fixe qui décrit la demande en service de mobilité au cours d'une journée moyenne. La constitution de cette matrice revêt une importance primordiale quant aux résultats que peut fournir l'algorithme d'optimisation. Or, dans la réalité, la demande est une donnée instable, qui suit des lois stochastiques. La matrice de mobilité utilisée décrit donc une estimation de la demande moyenne. La prise en compte de l'incertitude sur les données, mais aussi de la variabilité de l'utilisation du service d'un jour à l'autre, mériterait qu'une approche robuste soit envisagée. Nous ne l'avons pas abordé dans ce travail mais il s'agit là d'une perspective qui nous semble intéressante.

Des perspectives pour les métaheuristiques hybrides Enfin, parmi les nombreux champs qui s'ouvrent suite à ce travail, nous portons un intérêt tout particulier aux métaheuristiques hybrides. Elles ont été étudiées dans ce mémoire pour résoudre des problèmes multiobjectifs, mais elles constituent encore un formidable objet de recherche en contexte mono-objectif également. L'intégration d'une recherche locale au sein d'un algorithme génétique telle que nous l'avons présentée a montré les très bons niveaux de performance qu'il était possible d'atteindre. Nous avons appliqué des mécanismes similaires pour résoudre le problème académique de coloriage de graphe, un problème mono-objectif connu pour être très difficile. L'équilibre entre intensification et diversification fait partie des éléments critiques lorsque l'on a recours à des approches heuristiques. Afin d'en assurer une meilleure maîtrise, nous avons proposé un algorithme mémétique très particulier [Gondran 2013] qui s'appuie sur une population à deux individus seulement. Ce travail nous a permis de diminuer le nombre de couleurs nécessaires pour résoudre certaines instances DIMACS, par rapport aux approches les plus performantes présentées dans [Galinier 2013]. Cet algorithme laisse entrevoir de nouveaux axes de recherche que nous ne manquerons pas d'explorer.

Liste des publications

Conférences internationales avec comité de lecture

- [1] Laurent Moalic, Sid Lamrous, and Alexandre Caminada. A memetic algorithm for multiobjective problems. In *GECCO - Genetic and Evolutionary Computation Conference*, pages 93–94, Amsterdam, Netherlands, july 2013.
- [2] Laurent Moalic, Sid Lamrous, and Alexandre Caminada. A multiobjective memetic algorithm for solving the carsharing problem. In *WORLDCOMP'13 - The 2013 World Congress in Computer Science, Computer Engineering, and Applied Computing*, Las Vegas, USA, july 2013.
- [3] Alexandre Gondran and Laurent Moalic. Analysis of memetic approaches for graph coloring problem. In *26th European Conference on Operational Research*, Rome, Italy, june 2013.
- [4] Laurent Moalic, Alexandre Caminada, and Sid Lamrous. A fast local search approach for multiobjective problems. In *LION - Learning and Intelligent Optimization Conference*, Catania, Italy, january 2013.
- [5] Laurent Moalic, Sid Lamrous, and Alexandre Caminada. Multiobjective tabu algorithm for optimizing carsharing systems. In *International Conference on Metaheuristics and Nature Inspired Computing, META'12*, Sousse, Tunisia, october 2012.
- [6] K. Ait Ali, A. Gondran, A. Caminada, and L. Moalic. Vehicular radio connectivity in urban environment. In *Wireless Communication Systems (ISWCS), 2010 7th International Symposium on*, pages 305 –309, sept. 2010.
- [7] K. Ait Ali, M. Lalam, L. Moalic, and O. Baala. V-mbmm : Vehicular mask-based mobility model. In *Networks (ICN), 2010 Ninth International Conference on*, pages 243 –248, april 2010.
- [8] C. Joumaa, L. Moalic, S. Lamrous, and A. Caminada. Mobility models comparative study. In *Computers Industrial Engineering, 2009. CIE 2009. International Conference on*, pages 1798 –1803, july 2009.
- [9] C. Joumaa, L. Moalic, S. Lamrous, and A. Caminada. A simulation based mobility models comparative study. In *Proceedings of the 2nd International Conference on Simulation Tools and Techniques, Simutools'09*, pages 41 :1–41 :2, ICST, Brussels, Belgium, Belgium, 2009. ICST (Institute for Computer Sciences, Social-Informatics and Telecommunications Engineering).
- [10] Sid Lamrous, Alexandre Caminada, Laurent Moalic, and Chibli Joumaa. Cartographie de déplacements – de la réalité à la modélisation. In *4ème Symposium international Images, Multimédia, Applications graphiques et Environnements (IMAGE)*, Guelma, Algeria, november 2008.

Conférences et communications nationales

- [1] Laurent Moalic, Sid Lamrous, and Alexandre Caminada. Aide à la décision dans un contexte multi-objectifs. In *14e congrès annuel de la Société française de Recherche Opérationnelle et d'Aide à la Décision*, Troyes, France, feb. 2013.
- [2] Rabih Zakaria, Mohammad Dib, Alexandre Caminada, Laurent Moalic, and Omar Rifi. Cars relocation method for a carsharing system. In *14e congrès annuel de la Société française de Recherche Opérationnelle et d'Aide à la Décision*, Troyes, France, feb. 2013.
- [3] Sid Lamrous, Laurent Moalic, and Alexandre Caminada. Analyse des écarts entre l'offre d'infrastructures et les pratiques des cyclistes. In *Forum Chronos le vélo en continu*, Paris, France, octobre 2012.
- [4] Alexandre Caminada, Laurent Moalic, and Sid Lamrous. Connaître, évaluer, promouvoir, booster... les modes actifs : quelles données ? In *19ème Congrès du Club des Villes et Territoires Cyclabes*, Dijon, France, octobre 2011.
- [5] Alexandre Caminada, Laurent Moalic, and Sid Lamrous. Les enquêtes ménages déplacements. In *Atelier métier. 32e Rencontre des agences d'urbanisme*, Paris, France, octobre 2011.
- [6] Laurent Moalic, Sid Lamrous, and Alexandre Caminada. Analyse de sensibilité d'un algorithme glouton pour des services géo-localisés. In *12e congrès annuel de la Société française de Recherche Opérationnelle et d'Aide à la Décision*, Saint-Etienne, France, Mars 2011.
- [7] Sid Lamrous, Laurent Moalic, Chibli Joumaa, and Alexandre Caminada. Observation dynamique de la population à partir de la téléphonie mobile. In *Journées du CERTU sur la mobilité urbaine – Les recueils et les sources de données sur les déplacements*, Paris, France, Janvier 2010.
- [8] Laurent Moalic, Sid Lamrous, and Alexandre Caminada. Utilisation des réseaux cellulaires pour localiser une population. In *Journée thématique Géopositionnement et mobilités intelligentes*, Belfort, France, mars 2010.
- [9] Laurent Moalic, Alexandre Caminada, and Sid Lamrous. Une vision dynamique (spatiale et temporelle) de la mobilité urbaine en se basant sur le réseau de téléphonie mobile et les flux de déplacements dans le territoire. In *Déplacements et trafics : enquêtes, données, exploitations. FNAU - Club Transport et Mobilité*, Paris, France, février 2009.
- [10] Chibli Joumaa, Laurent Moalic, Sid Lamrous, and Alexandre Caminada. Mobility models comparative study. In *Journées non thématiques ResCom*, Strasbourg, France, octobre 2008.

Bibliographie

- [Ahas 2005] Rein Ahas et Ular Mark. *Location based services - new challenges for planning and public administration?* Futures, vol. 37, no. 6, pages 547 – 561, 2005. (Cité en page 28.)
- [Ahas 2010] Rein Ahas, Anto Aasa, Siiri Silm et Margus Tiru. *Daily rhythms of sub-urban commuters' movements in the Tallinn metropolitan area: Case study with mobile positioning data.* Transportation Research Part C: Emerging Technologies, vol. 18, no. 1, pages 45 – 54, 2010. (Cité en page 35.)
- [Ait Ali 2010a] K. Ait Ali, A. Gondran, A. Caminada et L. Moalic. *Vehicular radio connectivity in urban environment.* In Wireless Communication Systems (ISWCS), 2010 7th International Symposium on, pages 305 –309, sept. 2010. (Cité en page 41.)
- [Ait Ali 2010b] K. Ait Ali, M. Lalam, L. Moalic et O. Baala. *V-MBMM: Vehicular Mask-Based Mobility Model.* In Networks (ICN), 2010 Ninth International Conference on, pages 243 –248, avril 2010. (Cité en pages 21 et 41.)
- [Akyildiz 1996] I.F. Akyildiz, J.S.M. Ho et Yi-Bing Lin. *Movement-based location update and selective paging for PCS networks.* Networking, IEEE/ACM Transactions on, vol. 4, no. 4, pages 629–638, 1996. (Cité en page 12.)
- [Akyildiz 2000] Ian F. Akyildiz, Yi bing Lin, Wei ru Lai et Rong jaye Chen. *A new random walk model for PCS networks.* IEEE Journal of Selected Areas in Communications, pages 1254–1260, 2000. (Cité en page 12.)
- [Azarm 1994] K. Azarm et G. Schmidt. *Integrated mobile robot motion planning and execution in changing indoor environments.* In Intelligent Robots and Systems '94. 'Advanced Robotic Systems and the Real World', IROS '94. Proceedings of the IEEE/RSJ/GI International Conference on, volume 1, pages 298–305 vol.1, 1994. (Cité en page 19.)
- [Bar-Noy 1994] A. Bar-Noy, I. Kessler et M. Sidi. *Mobile users: to update or not to update?* In INFOCOM '94. Networking for Global Communications., 13th Proceedings IEEE, pages 570–576 vol.2, 1994. (Cité en pages ix, 13 et 14.)
- [Barth 2002] Matthew Barth et Susan Shaheen. *Shared-Use Vehicle Systems: Framework for Classifying Carsharing, Station Cars, and Combined Approaches.* Transportation Research Record: Journal of the Transportation Research Board, vol. 1791, no. -1, pages 105–112, 01 2002. (Cité en page 125.)
- [Basseur 2003] Matthieu Basseur, Franck Seynhaeve et El-Ghazali Talbi. *Adaptive mechanisms for multi-objective evolutionary algorithms.* In IMACS multi-conference, Computational Engineering in Systems Applications (CESA'03), pages 72–86, Lille, France, July 2003. (Cité en page 103.)

- [Bettstetter 2001] Christian Bettstetter. *Smooth is better than sharp: a random mobility model for simulation of wireless networks*. In Proceedings of the 4th ACM international workshop on Modeling, analysis and simulation of wireless and mobile systems, MSWIM '01, pages 19–27, New York, NY, USA, 2001. ACM. (Cité en page 16.)
- [Bettstetter 2003] C. Bettstetter, G. Resta et P. Santi. *The node distribution of the random waypoint mobility model for wireless ad hoc networks*. Mobile Computing, IEEE Transactions on, vol. 2, no. 3, pages 257–269, 2003. (Cité en page 15.)
- [Beume 2007] Nicola Beume, Boris Naujoks et Michael Emmerich. *SMS-EMOA: Multiobjective selection based on dominated hypervolume*. European Journal of Operational Research, vol. 181, no. 3, pages 1653 – 1669, 2007. (Cité en pages 101 et 109.)
- [Bittner 2005] S. Bittner, W.-U. Raffe et M. Scholz. *The area graph-based mobility model and its impact on data dissemination*. In Pervasive Computing and Communications Workshops, 2005. PerCom 2005 Workshops. Third IEEE International Conference on, pages 268–272, 2005. (Cité en page 26.)
- [Bonnell 2002] Patrick Bonnell. *Prévision de la demande de transport*. PhD thesis, Université Lumière-Lyon II, 2002. (Cité en page 21.)
- [Brockhoff 2008] Dimo Brockhoff, Tobias Friedrich et Frank Neumann. *Analyzing Hypervolume Indicator Based Algorithms*. In Günter Rudolph, Thomas Jansen, Simon Lucas, Carlo Poloni et Nicola Beume, éditeurs, Parallel Problem Solving from Nature – PPSN X, volume 5199 of *Lecture Notes in Computer Science*, pages 651–660. Springer Berlin Heidelberg, 2008. (Cité en page 101.)
- [Burulitisz 2004] A. Burulitisz, S. Imre et S. Szabo. *On the accuracy of mobility modelling in wireless networks*. In Communications, 2004 IEEE International Conference on, volume 4, pages 2302–2306 Vol.4, 2004. (Cité en page 13.)
- [Caceres 2007] N. Caceres, J.P. Wideberg et F.G. Benitez. *Deriving origin destination data from a mobile phone network*. Intelligent Transport Systems, IET, vol. 1, no. 1, pages 15 –26, march 2007. (Cité en pages 33 et 34.)
- [Caceres 2012] N. Caceres, L. M. Romero, F. G. Benitez et J. M. del Castillo. *Traffic Flow Estimation Models Using Cellular Phone Data*. Intelligent Transportation Systems, IEEE Transactions on, vol. 13, no. 3, pages 1430 –1441, sept. 2012. (Cité en page 33.)
- [Calabrese 2009] Francesco Calabrese. Wikicity: Real-time location-sensitive tools for the city, pages 390–413. IGI Global, Hershey, PA, USA, 2009. (Cité en page 30.)
- [Calabrese 2010] Francesco Calabrese, Francisco Pereira, Giusy Di Lorenzo, Liang Liu et Carlo Ratti. *The Geography of Taste: Analyzing Cell-Phone Mobility and Social Events*. In Patrik Floréen, Antonio Krüger et Mirjana Spasojevic,

- editeurs, Pervasive Computing, volume 6030 of *Lecture Notes in Computer Science*, pages 22–37. Springer Berlin / Heidelberg, 2010. (Cité en page 33.)
- [Calabrese 2011] F. Calabrese, G. Di Lorenzo, Liang Liu et C. Ratti. *Estimating Origin-Destination Flows Using Mobile Phone Location Data*. Pervasive Computing, IEEE, vol. 10, no. 4, pages 36–44, April 2011. (Cité en pages 33 et 34.)
- [Calabrese 2013] Francesco Calabrese, Mi Diao, Giusy Di Lorenzo, Joseph Ferreira Jr. et Carlo Ratti. *Understanding individual mobility patterns from urban sensing data: A mobile phone trace example*. Transportation Research Part C: Emerging Technologies, vol. 26, no. 0, pages 301 – 313, 2013. (Cité en page 33.)
- [Camp 2002] Tracy Camp, Jeff Boleng et Vanessa Davies. *A survey of mobility models for ad hoc network research*. Wireless Communications and Mobile Computing, vol. 2, no. 5, pages 483–502, 2002. (Cité en pages ix, 12 et 16.)
- [Campbell 2008] Scott W. Campbell et Yong Jin Park. *Social Implications of Mobile Telephony: The Rise of Personal Communication Society*. Sociology Compass, vol. 2, no. 2, pages 371–387, 2008. (Cité en page 40.)
- [Cano 2004a] J.-C. Cano et P. Manzoni. *Group mobility impact over TCP and CBR traffic in mobile ad hoc networks*. In Parallel, Distributed and Network-Based Processing, 2004. Proceedings. 12th Euromicro Conference on, pages 382–389, 2004. (Cité en page 15.)
- [Cano 2004b] J.-C. Cano, P. Manzoni et M. Sanchez. *Evaluating the impact of group mobility on the performance of mobile ad hoc networks*. In Communications, 2004 IEEE International Conference on, volume 7, pages 4039–4043 Vol.7, 2004. (Cité en page 15.)
- [Chariete 2013] Abderrahim Chariete, Oumaya Baala et Alexandre Caminada. *The impact of ground-surfaces on the mobile traffic for LTE cellular networks deployment*. In Wireless Days, Valencia, Spain, november 2013. (Cité en page 41.)
- [Coello Coello 1998] CarlosA. Coello Coello. *Using the Min-Max Method to Solve Multiobjective Optimization Problems with Genetic Algorithms*. In Helder Coelho, editeur, Progress in Artificial Intelligence, IBERAMIA 98, volume 1484 of *Lecture Notes in Computer Science*, pages 303–313. Springer Berlin Heidelberg, 1998. (Cité en page 91.)
- [Coello 2004] C.A.C. Coello et G.B. Lamont. Applications of multi-objective evolutionary algorithms, volume 1. World Scientific Publishing Company Incorporated, 2004. (Cité en page 112.)
- [Coello 2007] Carlos A. Coello Coello, Gary B Lamont et David A. Van Veldhuizen. Evolutionary algorithms for solving multi-objective problems : Second edition. Springer Science+Business Media, LLC, Boston, MA, 2007. (Cité en pages 80 et 115.)

- [Cohon 1978] J. L. Cohon. Multiobjective programming and planning. Academic Press, New York [etc.], 1978. (Cité en page 158.)
- [Collette 2002] Yann Collette et Patrick Siarry. Optimisation multiobjectif. Ed. Techniques Ingénieur, 2002. (Cité en page 108.)
- [Couronne 2011] T. Couronne, A.-M. Olteanu et Z. Smoreda. *Urban Mobility: Velocity and Uncertainty in Mobile Phone Data*. In Privacy, security, risk and trust (passat), 2011 ieee third international conference on and 2011 ieee third international conference on social computing (socialcom), pages 1425 –1430, oct. 2011. (Cité en page 35.)
- [de Almeida Correia 2012] Goncalo Homem de Almeida Correia et Antonio Pais Antunes. *Optimization approach to depot location and trip selection in one-way carsharing systems*. Transportation Research Part E: Logistics and Transportation Review, vol. 48, no. 1, pages 233 – 247, 2012. (Cité en page 126.)
- [Deb 2001] Kalyanmoy Deb et Tushar Goel. *Controlled Elitist Non-dominated Sorting Genetic Algorithms for Better Convergence*. In Proceedings of the First International Conference on Evolutionary Multi-Criterion Optimization, EMO '01, pages 67–81, London, UK, UK, 2001. Springer-Verlag. (Cité en page 87.)
- [Deb 2002] K. Deb, A. Pratap, S. Agarwal et T. Meyarivan. *A fast and elitist multiobjective genetic algorithm: NSGA-2*. Evolutionary Computation, IEEE Transactions on, vol. 6, no. 2, pages 182 –197, apr 2002. (Cité en pages 93, 94, 109, 115, 117 et 161.)
- [Deb 2003] Kalyanmoy Deb, Manikanth Mohan et Shikhar Mishra. *Towards a Quick Computation of Well-Spread Pareto-Optimal Solutions*. In CarlosM. Fonseca, PeterJ. Fleming, Eckart Zitzler, Lothar Thiele et Kalyanmoy Deb, éditeurs, Evolutionary Multi-Criterion Optimization, volume 2632 of *Lecture Notes in Computer Science*, pages 222–236. Springer Berlin Heidelberg, 2003. (Cité en page 96.)
- [Deb 2005] Kalyanmoy Deb, Manikanth Mohan et Shikhar Mishra. *Evaluating the epsilon-Domination Based Multi-Objective Evolutionary Algorithm for a Quick Computation of Pareto-Optimal Solutions*. Evolutionary Computation, vol. 13, no. 4, pages 501–525, 2013/10/07 2005. (Cité en page 94.)
- [Dowek 2007] Gilles Dowek. Les métamorphoses du calcul : une étonnante histoire de mathématiques. Le Pommier 2007 (18-Saint-Amand-Montrond, Paris, 2007. (Cité en page 83.)
- [Duan 2011] Zhengyu Duan, Liang Liu et Shang Wang. *MobilePulse: Dynamic profiling of land use pattern and OD matrix estimation from 10 million individual cell phone records in Shanghai*. In Geoinformatics, 2011 19th International Conference on, pages 1 –6, june 2011. (Cité en page 33.)
- [Erbaş 2001] F. Erbaş, J. Steuer, D. Eggesieker, K. Kyamakya et K. Jobinann. *A regular path recognition method and prediction of user movements in wireless*

- networks*. In Vehicular Technology Conference, 2001. VTC 2001 Fall. IEEE VTS 54th, volume 4, pages 2672–2676 vol.4, 2001. (Cit  en pages ix et 19.)
- [Erbas 2002] F. Erbas, K. Kyamakya, J. Steuer et K. Jobmann. *On the user profiles and the prediction of user movements in wireless networks*. In Personal, Indoor and Mobile Radio Communications, 2002. The 13th IEEE International Symposium on, volume 5, pages 2282–2286 vol.5, 2002. (Cit  en page 19.)
- [Fonseca 1993] Carlos M. Fonseca et Peter J. Fleming. *Genetic Algorithms for Multiobjective Optimization: Formulation, Discussion and Generalization*, 1993. (Cit  en page 93.)
- [Francesco 2006] CALABRESE Francesco et RATTI Carlo. *Real time Rome*. NET-COM, vol. 20, no. 3-4, pages 247–257, 2006. (Cit  en page 30.)
- [Frias-Martinez 2012] Vanessa Frias-Martinez, Cristina Soguero et Enrique Frias-Martinez. *Estimation of urban commuting patterns using cellphone network data*. In Proceedings of the ACM SIGKDD International Workshop on Urban Computing, UrbComp '12, pages 9–16, New York, NY, USA, 2012. ACM. (Cit  en page 33.)
- [Galinier 2013] Philippe Galinier, Jean-Philippe Hamiez, Jin-Kao Hao et Daniel Porumbel. *Recent Advances in Graph Vertex Coloring*. In Ivan Zelinka, Vaclav Snasel et Ajith Abraham,  diteurs, Handbook of Optimization, volume 38 of *Intelligent Systems Reference Library*, pages 505–528. Springer Berlin Heidelberg, 2013. (Cit  en pages 112 et 176.)
- [Gandibleux 1997] Xavier Gandibleux, Nazik Mezdaoui et Arnaud Fr eville. *A Tabu Search Procedure to Solve MultiObjective Combinatorial Optimization Problems*. In Rafael Caballero, Francisco Ruiz et Ralph Steuer,  diteurs, Advances in Multiple Objective and Goal Programming, volume 455 of *Lecture Notes in Economics and Mathematical Systems*, pages 291–300. Springer Berlin Heidelberg, 1997. (Cit  en page 90.)
- [Garey 1979] Michael R. Garey et David S. Johnson. *Computers and intractability; a guide to the theory of np-completeness*. W. H. Freeman & Co., New York, NY, USA, 1979. (Cit  en page 83.)
- [Goldberg 1989] David E. Goldberg. *Genetic algorithms in search, optimization and machine learning*. Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA, 1st  dition, 1989. (Cit  en pages 93 et 94.)
- [Gondran 2013] Alexandre Gondran et Laurent Moalic. *Analysis of Memetic Approaches for Graph Coloring Problem*. In 26th European Conference on Operational Research, Rome, Italy, june 2013. (Cit  en page 176.)
- [Gonzalez 2008] Marta C. Gonzalez, Cesar A. Hidalgo et Albert-Laszlo Barabasi. *Understanding individual human mobility patterns*. *Nature*, vol. 453, no. 7196, pages 779–782, 06 2008. (Cit  en page 33.)
- [Hansen 1997] Michael Pilegaard Hansen. *Tabu Search for Multiobjective Optimization: MOTS*. In MCDM'97. Springer-Verlag, 1997. (Cit  en page 90.)

- [Hansen 2001] Nikolaus Hansen et Andreas Ostermeier. *Completely Derandomized Self-Adaptation in Evolution Strategies*. *Evol. Comput.*, vol. 9, no. 2, pages 159–195, Juin 2001. (Cité en page 101.)
- [Ho 1995] Joseph S. M. Ho et Ian F. Akyildiz. *Mobile user location update and paging under delay constraints*. *Wirel. Netw.*, vol. 1, no. 4, pages 413–425, Décembre 1995. (Cité en page 12.)
- [Holleczek 2013] Thomas Holleczech, Liang Yu, Joseph K Lee, Oliver Senn, Kristian Kloeckl, Carlo Ratti et Patrick Jaillet. *Digital breadcrumbs: Detecting urban mobility patterns and transport mode choices from cellphone networks*. arXiv preprint arXiv:1308.6705, 2013. (Cité en page 32.)
- [Hong 1999] Xiaoyan Hong, Mario Gerla, Guangyu Pei et Ching-Chuan Chiang. *A group mobility model for ad hoc wireless networks*. In *Proceedings of the 2nd ACM international workshop on Modeling, analysis and simulation of wireless and mobile systems, MSWiM '99*, pages 53–60, New York, NY, USA, 1999. ACM. (Cité en page 15.)
- [Hoteit 2013] Sahar Hoteit, Stefano Secci, Stanislav Sobolevsky, Guy Pujolle et Carlo Ratti. *Estimating Real Human Trajectories through Mobile Phone Data*. In *Mobile Data Management (MDM), 2013 IEEE 14th International Conference on*, volume 2, pages 148–153. IEEE, 2013. (Cité en page 33.)
- [Igel 2007] Christian Igel, Nikolaus Hansen et Stefan Roth. *Covariance Matrix Adaptation for Multi-objective Optimization*. *Evolutionary Computation*, vol. 15, no. 1, pages 1–28, 2013/10/09 2007. (Cité en page 101.)
- [Ishibuchi 1998] H. Ishibuchi et T. Murata. *A multi-objective genetic local search algorithm and its application to flowshop scheduling*. *Systems, Man, and Cybernetics, Part C: Applications and Reviews, IEEE Transactions on*, vol. 28, no. 3, pages 392–403, 1998. (Cité en pages 89 et 102.)
- [Jardosh 2005] A.P. Jardosh, E.M. Belding-Royer, K.C. Almeroth et S. Suri. *Real-world environment models for mobile network evaluation*. *Selected Areas in Communications, IEEE Journal on*, vol. 23, no. 3, pages 622–632, 2005. (Cité en pages ix, 27 et 28.)
- [Jaszkiewicz 2002] Andrzej Jaszkiewicz. *Genetic local search for multi-objective combinatorial optimization*. *European Journal of Operational Research*, vol. 137, no. 1, pages 50 – 71, 2002. (Cité en pages 87, 102 et 112.)
- [Jeannic 2011] Thomas Le Jeannic, Philippe Roussel et Dominique Francois. *LamobilitéésFrançais,panoramaissudel'enquêtenationale transports et déplacements 2008*. *La Revue du Service de l'Observation et des Statistiques (SOeS) du Commissariat Général au Développement Durable (CGDD)*, 2011. (Cité en page 70.)
- [Jeon 2000] Wha Sook Jeon et Dong Geun Jeong. *Performance of improved probabilistic location update scheme for cellular mobile networks*. *Vehicular Technology, IEEE Transactions on*, vol. 49, no. 6, pages 2164–2173, 2000. (Cité en page 12.)

- [Joumaa 2007] C. Joumaa, A. Caminada et S. Lamrous. *Mask Based Mobility Model A new mobility model with smooth trajectories*. In Modeling and Optimization in Mobile, Ad Hoc and Wireless Networks and Workshops, 2007. WiOpt 2007. 5th International Symposium on, pages 1–6, 2007. (Cité en page 20.)
- [Joumaa 2009a] C. Joumaa, L. Moalic, S. Lamrous et A. Caminada. *Mobility models comparative study*. In Computers Industrial Engineering, 2009. CIE 2009. International Conference on, pages 1798–1803, july 2009. (Cité en page 41.)
- [Joumaa 2009b] C. Joumaa, L. Moalic, S. Lamrous et A. Caminada. *A simulation based mobility models comparative study*. In Proceedings of the 2nd International Conference on Simulation Tools and Techniques, Simutools'09, pages 41:1–41:2, ICST, Brussels, Belgium, Belgium, 2009. ICST (Institute for Computer Sciences, Social-Informatics and Telecommunications Engineering). (Cité en page 41.)
- [Kammann 2003] J. Kammann, M. Angermann et B. Lami. *A new mobility model based on maps*. In Vehicular Technology Conference, 2003. VTC 2003-Fall. 2003 IEEE 58th, volume 5, pages 3045–3049 Vol.5, 2003. (Cité en pages ix, 19 et 20.)
- [Kang 2013] Chaogui Kang, Stanislav Sobolevsky, Yu Liu et Carlo Ratti. *Exploring Human Movements in Singapore: A Comparative Analysis Based on Mobile Phone and Taxicab Usages*. UrbComp'13, Chicago, Illinois, USA., August 11-14 2013. (Cité en page 31.)
- [Kim 1998] Tai Suk Kim, Min Young Chung, Dan Keun Sung et M. Sengoku. *Mobility modeling and traffic analysis in three-dimensional indoor environments*. Vehicular Technology, IEEE Transactions on, vol. 47, no. 2, pages 546–557, 1998. (Cité en page 26.)
- [Kim 2000] Tai Suk Kim, Jae Kyun Kwon et Dan Keun Sung. *Mobility modeling and traffic analysis in three-dimensional high-rise building environments*. Vehicular Technology, IEEE Transactions on, vol. 49, no. 5, pages 1633–1640, 2000. (Cité en page 26.)
- [Knowles 1999] J. Knowles et D. Corne. *The Pareto archived evolution strategy: a new baseline algorithm for Pareto multiobjective optimisation*. In Evolutionary Computation, 1999. CEC 99. Proceedings of the 1999 Congress on, volume 1, pages 3 vol. (xxxvii+2348), 1999. (Cité en page 99.)
- [Knowles 2000] J.D. Knowles et D.W. Corne. *M-PAES: a memetic algorithm for multiobjective optimization*. In Evolutionary Computation, 2000. Proceedings of the 2000 Congress on, volume 1, pages 325–332 vol.1, 2000. (Cité en pages 87 et 114.)
- [Lam 1997] D. Lam, D.C. Cox et J. Widom. *Teletraffic modeling for personal communications services*. Communications Magazine, IEEE, vol. 35, no. 2, pages 79–87, 1997. (Cité en pages 17 et 18.)
- [Lamrous 2008] Sid Lamrous, Alexandre Caminada, Laurent Moalic et Chibli Joumaa. *Cartographie de déplacements – De la réalité à la modélisation*. In

- 4ème Symposium international Images, Multimédia, Applications graphiques et Environnements (IMAGE), Guelma, Algeria, november 2008. (Cité en pages 47 et 50.)
- [Lee 1961] C. Y. Lee. *An Algorithm for Path Connections and Its Applications*. Electronic Computers, IRE Transactions on, vol. EC-10, no. 3, pages 346–365, 1961. (Cité en page 19.)
- [Liang 2003] Ben Liang et Zygmunt J. Haas. *Predictive distance-based mobility management for multidimensional PCS networks*. IEEE/ACM Trans. Netw., vol. 11, no. 5, pages 718–732, Octobre 2003. (Cité en page 16.)
- [Liefoghe 2010] Arnaud Liefoghe, Laetitia Jourdan et El-Ghazali Talbi. *Metaheuristics and cooperative approaches for the Bi-objective Ring Star Problem*. Computers & Operations Research, vol. 37, no. 6, pages 1033 – 1044, 2010. (Cité en page 97.)
- [Liefoghe 2012] Arnaud Liefoghe, Jérémie Humeau, Salma Mesmoudi, Laetitia Jourdan et El-Ghazali Talbi. *On dominance-based multiobjective local search: design, implementation and experimental analysis on scheduling and traveling salesman problems*. Journal of Heuristics, vol. 18, pages 317–352, 2012. 10.1007/s10732-011-9181-3. (Cité en pages 98, 114 et 123.)
- [Liu 1998] Tong Liu, P. Bahl et I. Chlamtac. *Mobility modeling, location tracking, and trajectory prediction in wireless ATM networks*. Selected Areas in Communications, IEEE Journal on, vol. 16, no. 6, pages 922–936, 1998. (Cité en page 18.)
- [Lizee 2001] G. Lizée et A.O. Fapojuwo. *Highway traffic models for wireless networks*. In Vehicular Technology Conference, 2001. VTC 2001 Spring. IEEE VTS 53rd, volume 4, pages 2766–2770 vol.4, 2001. (Cité en page 17.)
- [M.-H. 2000] MASSOT M.-H. *PRAXITELE : Un concept, un service, une expérimentation*, 2000. (Cité en page 128.)
- [Markoulidakis 1997] J.G. Markoulidakis, G.L. Lyberopoulos, D.F. Tsirkas et E.D. Sykas. *Mobility modeling in third-generation mobile telecommunications systems*. Personal Communications, IEEE, vol. 4, no. 4, pages 41–56, 1997. (Cité en pages 17 et 18.)
- [Melab 2013] Nouredine Melab, Thé Van Luong, Karima Boufaras et El-Ghazali Talbi. *ParadisEO-MO-GPU: a Framework for Parallel GPU-based Local Search Metaheuristics*. In ACM GECCO’2013, Amsterdam, Pays-Bas, Juillet 2013. (Cité en page 103.)
- [Meunier 2000] Hervé Meunier, El-Ghazali Talbi et Philippe Reininger. *A Multiobjective Genetic Algorithm for Radio Network Optimization*. In In Proceedings of the 2000 Congress on Evolutionary Computation, pages 317–324. IEEE Press, 2000. (Cité en pages 104 et 135.)
- [Moalic 2012] Laurent Moalic, Sid Lamrous et Alexandre Caminada. *Multiobjective Tabu Algorithm for Optimizing Carsharing Systems*. In International Confer-

- ence on Metaheuristics and Nature Inspired Computing, META'12, Sousse, Tunisia, october 2012. (Cit  en page 161.)
- [Moalic 2013a] Laurent Moalic, Alexandre Caminada et Sid Lamrous. *A Fast Local Search Approach For Multiobjective problems*. In LION - Learning and Intelligent Optimization Conference, Catania, Italy, january 2013. (Cit  en pages 146 et 161.)
- [Moalic 2013b] Laurent Moalic, Sid Lamrous et Alexandre Caminada. *A memetic algorithm for multiobjective problems*. In GECCO - Genetic and Evolutionary Computation Conference, pages 93–94, Amsterdam, Netherlands, july 2013. (Cit  en page 161.)
- [Molina 2009] Juli n Molina, Luis V. Santana, Alfredo G. Hern ndez-D az, Carlos A. Coello Coello et Rafael Caballero. *g-dominance: Reference point based dominance for multiobjective metaheuristics*. European Journal of Operational Research, vol. 197, no. 2, pages 685 – 692, 2009. (Cit  en page 94.)
- [Moscato 1989] Pablo Moscato. *On evolution, search, optimization, genetic algorithms and martial arts: Towards memetic algorithms*. Caltech concurrent computation program, C3P Report, vol. 826, page 1989, 1989. (Cit  en page 112.)
- [Nguyen-Luong 2000a] D. Nguyen-Luong. *Modeles de pr visions de trafic aux Etats-Unis : Application   l laboration des plans de transports r gionaux*. Rapport technique, IAURIF, Janv. 2000. (Cit  en pages 22 et 23.)
- [Nguyen-Luong 2000b] D. Nguyen-Luong. *Recherche sur le choix modal en milieu urbain*. Rapport technique, IAURIF, Juin 2000. (Cit  en page 23.)
- [Nguyen-Luong 2008] D. Nguyen-Luong. *Projet SIMAURIF – Perfectionnement et valorisation - PREDIT R f : n 06MTE053*. Rapport technique, IAURIF, Juil. 2008. (Cit  en page 24.)
- [Pan 2006] Changxuan Pan, Jiangang Lu, Shan Di et Bin Ran. *Cellular-Based Data-Extracting Method for Trip Distribution*. Transportation Research Record: Journal of the Transportation Research Board, vol. 1945, no. -1, pages 33–39, 01 2006. (Cit  en page 33.)
- [Paquete 2003] Luis Paquete et Thomas St tzle. *A Two-Phase Local Search for the Biobjective Traveling Salesman Problem*. In CarlosM. Fonseca, PeterJ. Fleming, Eckart Zitzler, Lothar Thiele et Kalyanmoy Deb,  diteurs, Evolutionary Multi-Criterion Optimization, volume 2632 of *Lecture Notes in Computer Science*, pages 479–493. Springer Berlin Heidelberg, 2003. (Cit  en page 90.)
- [Paquete 2004] Luis Paquete, Marco Chiarandini et Thomas St tzle. *Pareto Local Optimum Sets in the Biobjective Traveling Salesman Problem: An Experimental Study*. In Xavier Gandibleux, Marc Sevaux, Kenneth S rensen, Vincent T'kindt, G. Fandel et W. Trockel,  diteurs, Metaheuristics for Multiobjective Optimisation, volume 535 of *Lecture Notes in Economics and Mathematical Systems*, pages 177–199. Springer Berlin Heidelberg, 2004. (Cit  en pages 98, 146 et 161.)

- [Pareto 1896] Vilfredo Pareto. Cours d'Economie Politique. Droz, Genève, 1896. (Cité en page 85.)
- [Perera 2002] Ranjit Perera, Andreas Eisenblätter, Erik Fledderus, Carmelita Görg, Michael Scheutzow et Stefan Verwijmeren. *Pixel Oriented Mobility Modelling for UMTS Network Simulations*. presented at IST Mobile and Wireless Telecommunications, 2002. (Cité en page 14.)
- [Phithakkitnukoon 2012] Santi Phithakkitnukoon, Zbigniew Smoreda et Patrick Olivier. *Socio-Geography of Human Mobility: A Study Using Longitudinal Mobile Phone Data*. PLoS ONE, vol. 7, no. 6, page e39253, 06 2012. (Cité en page 35.)
- [PREDIT 1999] PREDIT. *La modélisation dans les transports terrestres, présidé par Gilles Kahn, directeur scientifique INRIA*. Rapport du groupe de travail, PREDIT, Juin 1999. (Cité en page 10.)
- [PREDIT 2007] PREDIT. *Miro – Rapport Final – PREDIT Groupe opérationnel n°1 Mobilité territoires et développement durable*. Rapport technique, Laboratoire d'Informatique de l'Université de Franche-Comté, Laboratoire Théoriser et Modéliser pour Aménager, Laboratoire PACTE, Laboratoire Image et Ville, GEODES, Laboratoire CEDETE, Septembre 2007. (Cité en pages 24 et 25.)
- [Pulselli 2006] R M Pulselli, C Ratti et E Tiezzi. *City out of Chaos: Social Patterns and Organization in Urban Systems*. International Journal of Ecodynamics, vol. 1, no. 2, pages 125–134, 2006. (Cité en page 29.)
- [Pulselli 2008] R M Pulselli, P Romano, S Magaudda et Tiezzi E. *Monitoring human mobility in urban systems: a new technique based on cell-phone activity*. In 13th International Conference on Urban Planning in Information and Knowledge Society (REAL CORP), pages 629–635. CORP – Competence Center of Urban and Regional Planning, May 2008. (Cité en page 35.)
- [RATP 2010] RATP. *Métro Grand Paris – Étude de prévision de trafic*. Rapport technique, Département du Développement et de l'Action Territoriale, Oct. 2010. (Cité en pages 23 et 24.)
- [Ratti 2006] C Ratti, R M Pulselli, S Williams et D Frenchman. *Mobile Landscapes: using location data from cell phones for urban analysis*. Environment and Planning B: Planning and Design, vol. 33, no. 5, pages 727–748, 2006. (Cité en page 28.)
- [Ratti 2007] Carlo Ratti, Andres Sevtsuk, Sonya Huang et Rudolf Pailer. *Mobile Landscapes: Graz in Real Time*. In Georg Gartner, William Cartwright, Michael P. Peterson, William Cartwright, Georg Gartner, Liqiu Meng et Michael P. Peterson, éditeurs, Location Based Services and TeleCartography, Lecture Notes in Geoinformation and Cartography, pages 433–444. Springer Berlin Heidelberg, 2007. (Cité en page 31.)

- [Reades 2007] J. Reades, F. Calabrese, A. Sevtsuk et C. Ratti. *Cellular Census: Explorations in Urban Data Collection*. Pervasive Computing, IEEE, vol. 6, no. 3, pages 30–38, july-sept. 2007. (Cité en page 33.)
- [Rheingold 2002] H. Rheingold. *Smart mobs: the next social revolution*. Perseus Pub., 2002. (Cité en page 40.)
- [Roy 1985] B. Roy. *Méthodologie multicritère d'aide à la décision*. Economica, Paris, 1985. (Cité en page 159.)
- [Schaffer 1985] J. David Schaffer. *Multiple Objective Optimization with Vector Evaluated Genetic Algorithms*. In Proceedings of the 1st International Conference on Genetic Algorithms, pages 93–100, Hillsdale, NJ, USA, 1985. L. Erlbaum Associates Inc. (Cité en pages 92 et 109.)
- [Scourias 1999] John Scourias et Thomas Kunz. *An activity-based mobility model and location management simulation framework*. In Proceedings of the 2nd ACM international workshop on Modeling, analysis and simulation of wireless and mobile systems, MSWiM '99, pages 61–68, New York, NY, USA, 1999. ACM. (Cité en page 16.)
- [Sevtsuk 2010] Andres Sevtsuk. *Does Urban Mobility Have a Daily Routine? Learning from the Aggregate Data of Mobile Networks*. Journal of urban technology, vol. 17, no. 1, page 41, 2010. (Cité en page 35.)
- [Shaheen 2008] Susan A. Shaheen et Adam P Cohen. *Worldwide Carsharing Growth: An International Comparison*, 2008. (Cité en page 125.)
- [Silverman 1986] B. W. Silverman. *Density estimation for statistics and data analysis*. Chapman & Hall, London, 1986. (Cité en page 96.)
- [Song 2010] Chaoming Song, Zehui Qu, Nicholas Blumm et Albert-László Barabási. *Limits of Predictability in Human Mobility*. Science, vol. 327, no. 5968, pages 1018–1021, 2010. (Cité en page 35.)
- [Srinivas 1994] N. Srinivas et Kalyanmoy Deb. *Multiobjective Optimization Using Nondominated Sorting in Genetic Algorithms*. Evolutionary Computation, vol. 2, pages 221–248, 1994. (Cité en pages 93 et 94.)
- [Steenbruggen 2011] John Steenbruggen, Maria Borzacchiello, Peter Nijkamp et Henk Scholten. *Mobile phone data from GSM networks for traffic parameter and urban spatial pattern assessment: a review of applications and opportunities*. GeoJournal, pages 1–21, May 2011. 10.1007/s10708-011-9413-y. (Cité en pages 28 et 33.)
- [Systems 1997] Raytheon Systems et P. B. Farradyne. *FINAL EVALUATION REPORT For The CAPITAL-ITS Operational Test and Demonstration Program*, May 1997. (Cité en pages iii, 9 et 29.)
- [Tabbane 2002] S. Tabbane. *Ingénierie des réseaux cellulaires*. Collection Réseaux et télécommunications. Hermès science publications, 2002. (Cité en page 45.)
- [Talbi 2001] El-Ghazali Talbi, Malek Rahoual, Mohamed Hakim Mabed et Clarisse Dhaenens. *A Hybrid Evolutionary Approach for Multicriteria Optimization*

- Problems: Application to the Flow Shop.* In EMO, pages 416–428, 2001. (Cité en pages 103 et 112.)
- [Talbi 2009] El-Ghazali Talbi. *Metaheuristics: from design to implementation.* Wiley, 2009. (Cité en page 80.)
- [Talbi 2013] El-Ghazali Talbi. *Hybrid metaheuristics.* Springer, 2013. (Cité en pages 102 et 113.)
- [Tian 2002] J. Tian, J. Hähner, C. Becker, I. Stepanov et K. Rothermel. *Graph-based mobility model for mobile ad hoc network simulation.* In Simulation Symposium, 2002. Proceedings. 35th Annual, pages 337–344, 2002. (Cité en pages ix, 26 et 27.)
- [Tsai 1999] I-Fei Tsai et Rong-Hong Jan. *The lookahead strategy for distance-based location tracking in wireless cellular networks.* SIGMOBILE Mob. Comput. Commun. Rev., vol. 3, no. 4, pages 27–38, Octobre 1999. (Cité en page 14.)
- [Tuan 2003] Chiu-Ching Tuan et Chen-Chau Yang. *A compact normal walk model for PCS networks.* SIGMOBILE Mob. Comput. Commun. Rev., vol. 7, no. 4, pages 20–31, Octobre 2003. (Cité en pages ix, 14 et 15.)
- [Ulungu 1999] E.L. Ulungu, J. Teghem, P.H. Fortemps et D. Tuyttens. *MOSA method: a tool for solving multiobjective combinatorial optimization problems.* Journal of Multi-Criteria Decision Analysis, vol. 8, no. 4, pages 221–236, 1999. (Cité en page 90.)
- [While 2006] L. While, P. Hingston, L. Barone et S. Huband. *A faster algorithm for calculating hypervolume.* Evolutionary Computation, IEEE Transactions on, vol. 10, no. 1, pages 29–38, 2006. (Cité en pages xi, 106 et 108.)
- [Wu 2002] Chien-Hsing Wu, Huang-Pao Lin et Leu-Shing Lan. *A new analytic framework for dynamic mobility management of PCS networks.* Mobile Computing, IEEE Transactions on, vol. 1, no. 3, pages 208–220, 2002. (Cité en page 13.)
- [Yang 1994] Xiaofeng Yang et Mitsuo Gen. *Evolution program for bicriteria transportation problem.* Computers & Industrial Engineering, vol. 27, no. 1–4, pages 481 – 484, 1994. <ce:title>16th Annual Conference on Computers and Industrial Engineering</ce:title>. (Cité en page 90.)
- [Zitzler 1999] E. Zitzler et L. Thiele. *Multiobjective evolutionary algorithms: a comparative case study and the strength Pareto approach.* Evolutionary Computation, IEEE Transactions on, vol. 3, no. 4, pages 257 –271, nov 1999. (Cité en pages 94, 100, 104, 105, 135 et 136.)
- [Zitzler 2001] Eckart Zitzler, Marco Laumanns et Lothar Thiele. *SPEA2: Improving the Strength Pareto Evolutionary Algorithm.* TIK-Report 103, 2001. (Cité en pages 94, 95, 109 et 115.)
- [Zitzler 2003] E. Zitzler, L. Thiele, M. Laumanns, C.M. Fonseca et V.G. da Fonseca. *Performance assessment of multiobjective optimizers: an analysis and review.* Evolutionary Computation, IEEE Transactions on, vol. 7, no. 2, pages 117 – 132, april 2003. (Cité en pages 105 et 135.)

- [Zitzler 2004] Eckart Zitzler et Simon Künzli. *Indicator-based selection in multiobjective search*. In in Proc. 8th International Conference on Parallel Problem Solving from Nature (PPSN VIII, pages 832–842. Springer, 2004. (Cité en pages 93, 96 et 100.)
- [Zitzler 2007] Eckart Zitzler, Dima Brockhoff et Lothar Thiele. *The Hypervolume Indicator Revisited: On the Design of Pareto-compliant Indicators Via Weighted Integration*. In Shigeru Obayashi, Kalyanmoy Deb, Carlo Poloni, Tomoyuki Hiroyasu et Tadahiko Murata, éditeurs, Evolutionary Multi-Criterion Optimization, volume 4403 of *Lecture Notes in Computer Science*, pages 862–876. Springer Berlin Heidelberg, 2007. (Cité en page 101.)

