



Contributions au processus d'Aide Multicritère à la Décision : Méthodes, Outils et Applications

Patrick Meyer

► To cite this version:

Patrick Meyer. Contributions au processus d'Aide Multicritère à la Décision : Méthodes, Outils et Applications. Recherche opérationnelle [math.OC]. Dauphine recherche en Management, 2013. tel-00959830

HAL Id: tel-00959830

<https://theses.hal.science/tel-00959830>

Submitted on 17 Mar 2014

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

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

Contributions au processus d'Aide Multicritère à la Décision : Méthodes, Outils et Applications

Mémoire présenté et soutenu publiquement
le 7 novembre 2013 en vue de l'obtention de
l'Habilitation à Diriger des Recherches en Informatique

par

Patrick Meyer

Composition du jury

Rapporteurs : Alberto Colorni, Professeur, Politecnico di Milano
Michel Grabisch, Professeur, Université Paris I Panthéon-Sorbonne
Vincent Mousseau, Professeur, École Centrale Paris
Examineurs : Raymond Bisdorff, Professeur, Université du Luxembourg
Denis Bouyssou, Directeur de Recherche CNRS, Université Paris Dauphine
Marc Pirlot, Professeur, Université de Mons
Coordinateur : Alexis Tsoukiàs, Professeur, Université Paris Dauphine

Contributions au processus d'Aide Multicritère à la Décision : Méthodes, Outils et Applications

Mémoire présenté et soutenu publiquement
le 7 novembre 2013 en vue de l'obtention de
l'Habilitation à Diriger des Recherches en Informatique

par

Patrick Meyer

Composition du jury

Rapporteurs : Alberto Colorni, Professeur, Polytecnico di Milano
Michel Grabisch, Professeur, Université Paris I Panthéon-Sorbonne
Vincent Mousseau, Professeur, École Centrale Paris
Examineurs : Raymond Bisdorff, Professeur, Université du Luxembourg
Denis Bouyssou, Directeur de Recherche CNRS, Université Paris Dauphine
Marc Pirlot, Professeur, Université de Mons
Coordinateur : Alexis Tsoukiàs, Professeur, Université Paris Dauphine

'Sir, What is the secret of your success ?' a reporter asked a bank president.

'Two words.'

'And, sir, what are they ?'

'Good decisions.'

'And how do you make good decisions ?'

'One word.'

'And sir, what is that ?'

'Experience.'

'And how do you get Experience ?'

'Two words.'

'And, sir, what are they ?'

'Bad decisions.'

Unknown author

Remerciements

J'aimerais tout d'abord remercier Alberto Colorni, Michel Grabisch et Vincent Mousseau d'avoir consenti de jouer le rôle de rapporteur de ce travail. Merci également à Raymond Bisdorff, Denis Bouyssou et Marc Pirlot pour leur participation au jury, ainsi qu'à Alexis Tsoukiàs pour avoir accepté de coordonner cette habilitation.

Ce travail résume cinq années de recherche, une période qui a été ponctuée par de nombreuses rencontres stimulantes et créatives.

J'ai notamment eu la chance de collaborer avec de nombreuses personnes sur des sujets divers et variés : Sébastien Bigaret, Raymond Bisdorff, Olivier Cailloux, Vanea Chirprianov, Dhouha Kbaier, Yannick Le Bras, Stéphane Lallich, Philippe Lenca, Vincent Mousseau, Pritesh Narayan, Alexandru-Liviu Olteanu, Marc Pirlot, Grégory Ponthière, Gérald-Reparate Retali, Jacques Simonin, Thomas Veneziano. Chacune de ces personnes a apporté à sa façon une pierre à l'édifice qu'est ce mémoire, et je les en remercie chaleureusement.

Sans vouloir faire de jaloux, j'aimerais remercier plus spécifiquement Sébastien Bigaret, sans l'aide de qui, une grande partie des outils informatiques qui sont décrits dans ce mémoire n'existeraient probablement pas sous leur forme actuelle, et n'auraient pas le succès qu'ils rencontrent aujourd'hui. Merci Sébastien de supporter l'informaticien en herbe que je suis, et merci pour toutes tes initiatives ainsi que pour ton regard percutant, incisif et constructif sur mes travaux.

La grande majorité de ces travaux a été réalisée à Télécom Bretagne dans l'UMR CNRS Lab-STICC. J'aimerais remercier Gilles Coppin, Yvon Kermarrec et Philippe Lenca pour leur accueil, ainsi que les conseils qu'ils ont pu me donner tout au long de ces cinq années pour m'aider à orienter ma carrière et mes travaux.

Merci aussi à Sophie et à Noé, grâce à qui je quitte tous les soirs mon bureau avec la perspective d'une deuxième partie de journée riche en événements et en surprises, et surtout, qui me rappellent régulièrement les choses essentielles de la vie.

Table des matières

Remerciements	vii
Table des figures	xiii
Liste des tableaux	xv
Préface	xvii
1 Contexte et objectifs	1
1.1 L'Aide Multicritère à la Décision : brève introduction et concepts essentiels	1
1.2 Motivations et objectifs scientifiques	5
2 Aspects algorithmiques et méthodologiques de l'AMCD	7
2.1 Expressivité empirique de modèles d'agrégation	8
2.1.1 Cadre expérimental	9
2.1.2 Résultats et observations	11
2.1.3 Conclusions	12
2.2 Formalisation du concept de classification non supervisée et méthode de résolution	13
2.2.1 Définitions et objectifs pour une classification non supervisée	14
2.2.2 Méthode de résolution	15
2.2.3 Illustration et validation empirique	17
2.2.4 Conclusions	19
2.3 Aide au tri en présence de critères et de décideurs multiples	20
2.3.1 Electre Tri et méthodes d'élicitation des préférences connexes	21
2.3.2 Inférence des profils des catégories à partir d'exemples d'affectation .	23
2.3.3 Processus d'aide au tri pour un groupe de décideurs	23
2.3.4 Exemple illustratif et validation empirique	24
2.3.5 Conclusions et perspectives	27
2.4 Elicitation stable des poids d'une relation de surclassement	28
2.4.1 Au sujet de la stabilité de la relation de surclassement	29

2.4.2	Algorithme de détermination stable des poids	31
2.4.3	Validation empirique et processus d'élicitation	32
2.4.4	Conclusions et perspectives	34
3	Outils de soutien au processus d'AMCD	37
3.1	Le standard XMCDA	38
3.1.1	Présentation	39
3.1.2	Illustration	40
3.1.3	Conclusion	43
3.2	Les services-web XMCDA	44
3.2.1	Motivations	44
3.2.2	Détails techniques et accès	46
3.2.3	Conclusion	47
3.3	La plateforme diviz	48
3.3.1	Description et utilisation	48
3.3.2	Exemple d'utilisation	51
3.3.3	Conclusions et perspectives	55
3.4	Modélisation du processus d'AMCD	56
3.4.1	Aperçu du modèle	57
3.4.2	Une instance du processus d'AMCD	61
3.4.3	Perspectives	65
4	Applications de l'AMCD	67
4.1	Au sujet de la construction d'indicateurs de bien-être en économie	68
4.1.1	Cadre expérimental	70
4.1.2	Résultats	73
4.1.3	Conclusions et perspectives	76
4.2	Construction d'une échelle de risque territorial prenant en compte plusieurs critères et experts	76
4.2.1	Méthodologie multicritère et multi-décideurs pour la construction d'une échelle de risque territorial	77
4.2.2	Exemple illustratif	79
4.2.3	Conclusion	83
5	Conclusion et recherches futures	85
	Bibliographie	91

Annexe A : Curriculum vitæ et activités de recherche	99
A.1 Curriculum vitæ	99
A.1.1 État civil	99
A.1.2 Formation	100
A.1.3 Fonctions assurées	100
A.1.4 Compétences	101
A.2 Thèmes de recherche	101
A.2.1 Contributions au processus d'AMCD	101
A.2.2 Fouille de données	102
A.3 Responsabilités scientifiques et animation de la recherche	103
A.3.1 Responsabilités	103
A.3.2 Encadrement de thèses de doctorat	103
A.3.3 Encadrement de stages de Master Recherche	104
A.4 Collaborations récentes nationales et internationales	104
A.5 Participation à des projets	104
A.6 Affiliations	105
A.7 Travaux de rédacteur	105
A.8 Participation à des comités de lecture	106
A.9 Membre de comités de programme de conférences	106
A.10 Membre de comités d'organisation de conférences scientifiques	106
A.11 Séjours de recherche	107
Annexe B : Publications	109
B.1 Articles dans des revues internationales avec comité de lecture	109
B.2 Articles dans des ouvrages collectifs avec comité de lecture	110
B.3 Articles dans des revues françaises avec comité de lecture	111
B.4 Rédacteur invité	111
B.5 Articles dans des actes de conférences avec comité de lecture	111
B.6 Articles dans des bulletins de sociétés scientifiques	113
B.7 Communications à des séminaires et des conférences sans actes	113

Table des figures

2.1	Structures potentielles du résultat d'une classification non supervisée en AMCD.	15
2.2	Résultats de la classification non supervisée.	18
2.3	Les restrictions sur les poids imposées par le choix des profils communs. Un arc d'un critère j à un critère j' signifie que le choix de ces profils implique que ce décideur devra choisir un poids pour j strictement plus élevé que pour j'	26
3.1	Le tableau de performance de Thierry au format HTML.	44
3.2	Architecture des services-web XMCD.	46
3.3	Une session de travail typique dans diviz.	49
3.4	Un workflow représentant les données d'entrée (à gauche), et un module de somme pondérée (au centre) qui est combiné à un module de représentation graphique du résultat (à droite).	50
3.5	Le workflow pour le problème de sélection d'une voiture.	52
3.6	La partie "ACUTA" du workflow.	52
3.7	Fonctions de valeur marginales.	52
3.8	Diagramme à barres des scores des alternatives.	53
3.9	La partie PROMETHEE du workflow.	53
3.10	Diagramme à barres des flots nets.	54
3.11	Comparaison des deux méthodes.	54
3.12	Graphe XY représentant les deux classements (abscisse : rangs par la méthode PROMETHEE, ordonnée : rangs par ACUTA).	55
3.13	La méthode de construction du modèle du processus d'AMCD.	58
3.14	Le processus général (gauche) et le sous-processus <i>Decision aid</i> (droite). . .	58
3.15	L'activité <i>Construct the evaluation model</i>	60
3.16	Graphes en étoile de deux alternatives obtenus par le service-web <code>plotStarGraphPerformanceTable</code>	62
3.17	La première étape de la construction progressive du workflow diviz supportant le processus d'AMCD.	63
3.18	Représentations graphiques des fonctions de valeur pour les critères g_1 (coût) et g_2 (accélération), obtenues à l'aide du service <code>plotValueFunctions</code> . . .	63

3.19	La deuxième étape de la construction progressive du workflow diviz.	63
3.20	Début du classement des voitures.	64
3.21	La troisième étape de la construction progressive du workflow diviz.	64
4.1	Courbes d'indifférence pour les 3 individus fictifs.	71
4.2	Le workflow diviz implémentant le processus de l'exemple illustratif.	82
4.3	Diagramme à barres des poids consensuels et diagramme en étoile de la zone ZoneC (resp. ZoneB), affectée à la catégorie "risque bas" (resp. "risque élevé") d'après la règle d'affectation ELECTRE TRI.	82
4.4	Affectation automatique des 6 zones réelles à l'aide du modèle ELECTRE TRI.	83

Liste des tableaux

2.1	Moyennes des résultats (écarts types entre parenthèses).	19
2.2	Procédures d'inference pour ELECTRE TRI.	22
2.3	Quelques exemples d'affectation fournis par les quatre décideurs.	24
2.4	Performances de quelques exemples d'affectations.	25
2.5	Profils inférés et veto.	25
2.6	Des jeux de poids et des seuils de majorité pour chaque décideur, compatibles avec les profils élicités et les exemples d'affectation.	25
2.7	Poids et seuil de majorité λ consensuels.	26
2.8	Pourcentage de problèmes résolus en moins de 90 minutes.	27
2.9	Tableau de performance, relation de surclassement $\tilde{S}^w$ et stabilité associée.	31
2.10	Accroissements de stabilité et temps de calculs.	33
2.11	Résultats pour la reconstruction itérative, en utilisant l'heuristique de sélection VSF.	34
3.1	Tableau de performance pour le problème de choix de Thierry.	40
3.2	Les informations intra- et inter-critères.	54
3.3	Correspondance entre les activités du processus d'AMCD et les services-web XMCDa pour le problème de Thierry.	64
4.1	Classement international, Human Development Index, 2007	69
4.2	Description des préférences des trois individus A , B et C	70
4.3	Construction des sociétés fictives	72
4.4	Fonction de valeur "en s"	73
4.5	Niveaux de k -additivity minimaux pour les 13 répondants analysés.	74
4.6	Les valeurs des indices de Shapley pour les 13 répondants analysés.	74
4.7	Nombre d'interactions strictement positives et négatives pour les 13 répondants.	75
4.8	Performances des 7 pays sur les 5 critères.	75
4.9	Rangs des 7 pays pour chacun des répondants	75
4.10	Partie des exemples de zones données par les experts.	80

4.11	Limites des catégories inférées. b_1 sépare les catégories “risque élevé” et “risque moyen”, alors que b_2 sépare les catégories “risque moyen” et “risque bas”.	80
4.12	Des vecteurs de poids pour chaque expert, déterminés par le programme ICL.	80
4.13	Limites des catégories inférées lors du second essai. b_1 sépare les catégories “risque élevé” et “risque moyen”, alors que b_2 sépare les catégories “risque moyen” et “risque bas”.	81
4.14	Des vecteurs de poids pour chaque expert, déterminés par le programme ICLV.	81
4.15	Un vecteur de poids consensuel respectant 96 des exemples d’affectation.	81
4.16	Zones réelles à évaluer en fonction de leur risque, et l’affectation par ELECTRE TRI.	82

Préface

Le but de ce mémoire est de présenter une synthèse des résultats que j’ai obtenus depuis la soutenance de ma thèse de doctorat en 2007 (Meyer, 2007) et de proposer des pistes de recherches à court et à moyen terme. Ces réalisations passées concernent principalement le domaine de l’Aide Multicritère à la Décision (AMCD). Mon objectif ici est d’éviter de faire un inventaire à la Prévert, et de démontrer une cohérence entre les différents travaux que j’ai pu effectuer ces 5 dernières années.

Le premier chapitre me permet d’introduire le domaine de recherche en AMCD, avant de passer à une présentation des objectifs scientifiques qui ont guidé mes travaux. J’y évoque notamment mon souci récurrent de proposer des solutions d’aide à la décision opérationnelles, mon désir de validation empirique de mes propositions, ainsi que mon envie de développer des outils de soutien du processus d’AMCD performants et pérennes. Ces objectifs sont perceptibles tout au long des résultats présentés dans ce mémoire.

Le deuxième chapitre présente mes contributions algorithmiques et méthodologiques aux 2 principaux courants de pensée de l’AMCD. Au delà du désir de répondre à des besoins issus d’applications réelles, le fil rouge de ces travaux est la validation expérimentale des algorithmes proposés. Je pense en effet que celle-ci doit aller au-delà de caractérisations mathématiques et de confrontations à quelques cas réels, et doit rejoindre ce qui se fait dans des domaines comme l’intelligence artificielle ou la fouille de données, où les algorithmes sont confrontés à des grandes quantités de données générées artificiellement. Je montre notamment dans ce chapitre comment ces expérimentations peuvent aider à mieux utiliser les algorithmes proposés en pratique, en plus de leur donner une validité accrue.

L’amélioration des outils informatiques de soutien du processus d’AMCD est une autre de mes préoccupations. Le troisième chapitre traite de cette thématique, et présente les travaux que j’ai réalisés dans le contexte du Decision Deck Consortium. J’y présente notamment l’écosystème diviz qui réunit dans un même environnement un grand nombre d’algorithmes et de procédures d’AMCD. L’outil diviz est utilisable aussi bien en recherche et en enseignement que lors de la résolution de problèmes de décision réels. Son fort potentiel de dissémination de résultats a récemment été démontré, et, comme en témoignent les nombreux cours d’AMCD qui l’ont adopté, diviz apparaît comme un excellent outil de formation.

Je porte également beaucoup d’intérêt à la diffusion des résultats de mes recherches dans d’autres domaines scientifiques qui pourraient bénéficier d’une prise en compte des préférences humaines dans leurs modèles mathématiques. Le quatrième chapitre présente ainsi quelques exemples de telles applications, et je m’arrête plus particulièrement à l’étude d’indicateurs de bien-être en économie et la construction d’échelles de risque territorial.

Pour terminer, en annexe le lecteur trouvera d’abord un curriculum vitæ et un descriptif de mes activités de recherche, avant de pouvoir consulter la liste de mes publications scientifiques et de mes présentations à des colloques et séminaires.

En dehors de mes travaux en AMCD, j'ai récemment participé à la définition de la robustesse de mesures de qualité en fouille de données, plus particulièrement dans un contexte de règles d'associations (Le Bras *et al.*, 2010a,b, 2011). Par ailleurs, toujours dans ce domaine, j'ai contribué à proposer une méthode de classification non supervisée (Bisdorff *et al.*, 2011). Ces travaux ne sont cependant pas détaillés dans ce mémoire.

PATRICK MEYER

Quilhouarn le 17 septembre 2013

Chapitre 1

Contexte et objectifs

Sommaire

1.1 L'Aide Multicritère à la Décision : brève introduction et concepts essentiels	1
1.2 Motivations et objectifs scientifiques	5

Ce chapitre a comme objectif principal de délimiter le cadre de nos recherches et de situer notre discours dans le contexte scientifique de l'Aide Multicritère à la Décision. A cette fin, une première section présente ce domaine de manière succincte et introduit les concepts que nous utiliserons et auxquels nous ferons référence par la suite. Ensuite, une deuxième section présente notre vision de la recherche dans ce domaine, les points qui nous ont motivés à travailler sur les sujets présentés dans ce mémoire, ainsi que nos objectifs scientifiques.

1.1 L'Aide Multicritère à la Décision : brève introduction et concepts essentiels

Aider à prendre des décisions de manière plus éclairée, tenir compte de facteurs humains, organisationnels et sociaux dans les problèmes de décision, formaliser les préférences des différents acteurs et être capable de faire face à plusieurs points de vue dans le processus décisionnel, gérer les contraintes et exigences de décideurs multiples, . . . sont des sujets qui font tous partie des préoccupations des chercheurs du domaine de l'Aide Multicritère à la Décision (AMCD).

L'AMCD est une branche de la recherche opérationnelle, et a comme objectif d'aider un ou plusieurs décideurs à préparer et à prendre des décisions, lorsque plusieurs points de vue, le plus souvent conflictuels, doivent être pris en compte. L'objectif de cette aide n'étant pas de forcer une prise de décision à tout prix ou de se substituer au(x) décideur(s), l'AMCD peut simplement se résumer à une structuration du problème de décision en vue de sa meilleure compréhension, ou aller jusqu'à l'élaboration d'une recommandation de décision.

Les problèmes de décision pour lesquels une telle aide multicritère peut être envisagée concernent les décisions complexes du quotidien, comme par exemple les décisions managériales dans des organisations, la gestion du risque dans l'industrie, la sélection de

fournisseurs, le tri de réponses à des appels d’offres, les décisions de triage dans les urgences hospitalières, le design de systèmes complexes, le choix de stratégies d’investissement, ...

L’idée que les mathématiques doivent être utilisées pour servir la connaissance dans son sens le plus large est très répandue parmi les scientifiques. Il est ainsi très tentant de modéliser des situations de la vie réelle par des principes logiques forts en vue de les décrire, de les expliquer, ou même de prédire des événements futures qui pourraient en découler. Cependant, lorsqu’il s’agit de représenter des décisions, les modèles mathématiques se heurtent souvent à l’irrationalité des humains.

Cette nécessaire intervention de l’humain dans la prise de décision représente ainsi une des difficultés majeures que rencontre un chercheur tentant de fournir des outils de recommandation au décideur. L’AMCD peut donc être vue comme un processus préparatoire à la prise de décision tentant de garantir un certain niveau d’objectivité.

Ce processus d’AMCD nécessite en général au moins 2 acteurs : d’une part le décideur (qui peut dans le cas le plus général être un groupe de décideurs), qui portera la responsabilité de la prise de décision et de ses conséquences, et qui est régi par un certain système de valeurs qu’il s’agira d’explicitier ; et d’autre part l’analyste qui a comme tâche de faciliter la prise de décision en guidant le décideur à travers le processus d’aide à la décision. Ce processus est défini par toute une série de tâches qui vont de la définition et la structuration du problème à la construction d’une recommandation de décision, en passant par l’utilisation d’un algorithme d’agrégation. La complexité de ce processus dépend donc notamment du contexte du problème de décision et des acteurs qui ont une influence sur la prise de décision.

Les méthodes classiques de la recherche opérationnelle ont tendance à formuler un problème de décision en des termes analytiques (contraintes, fonction objectif, ...) et proposent une résolution en recherchant une solution optimale. Le décideur intervient dans ce cas dans la délimitation et la formulation du problème ainsi que lors de la validation de la solution. En AMCD par contre, l’accent est mis sur le rôle central du décideur, et une attention toute particulière est portée à l’inclusion d’une représentation de son système de valeur dans les modèles mathématiques. Il intervient ainsi à toutes les étapes du processus d’aide à la décision et la recommandation qui en résulte est en général assez fidèle à ses objectifs initiaux.

D’autre part, les problèmes liés à la prise de décision sont souvent dus à la nature multidimensionnelle des options de décision. En effet, même les problèmes de décision d’apparence simple auxquels nous sommes régulièrement confrontés sont souvent soumis à de multiples dimensions de préférences (coût, qualité, écologie, ...), qui de plus ont tendance à être conflictuelles (un coût peu élevé n’est en général pas compatible avec une qualité de fabrication élevée). Lorsqu’il s’agit alors de choisir la “meilleure” option, l’AMCD propose plutôt de rechercher un bon compromis, qui prend en compte les préférences du décideur, afin de faciliter l’adoption de cette recommandation par le décideur.

D’un point de vue historique, les origines de l’AMCD remontent au moins jusqu’au 18^e siècle, où le Marquis de Condorcet a été le premier à appliquer systématiquement des théories mathématiques aux sciences sociales. En 1785 il écrit son ouvrage “Essai sur l’application de l’analyse à la probabilité des décisions rendues à la pluralité des voix”, qui traite de prise de décision en présence de plusieurs votants. Cependant les bases de l’AMCD moderne ont été posées au milieu du 20^e siècle avec la théorie des préférences révélées de [Samuelson \(1938\)](#), les débuts de la théorie des jeux par [von Neumann et Morgenstern \(1944\)](#), l’émergence de la théorie du choix social par [Arrow \(1951\)](#) et les recherches sur

les aspects psychologiques et mathématiques des décisions par Luce et Raiffa (1957) et Fishburn (1970).

Dans les années 1950, un pas important est fait par Simon vers une fondation pragmatique des théories de la décision avec son concept de rationalité limitée (Simon, 1957). Il stipule que dans une décision de la vie réelle, différents facteurs limitent la capacité d'un décideur à faire un choix complètement rationnel. Par conséquent, le décideur est soumis à une rationalité limitée et il choisira une option qui prendra en compte à la fois ses limitations sur la connaissance et celles sur ses capacités cognitives. D'un point de vue mathématique, cette décision ne sera pas forcément optimale, mais aura tendance à satisfaire le décideur.

A la fin des années 1960, les premières méthodes de résolution de problèmes de décision multidimensionnels commencent à apparaître. En 1968 Roy crée la branche des méthodes de surclassement (Roy, 1968) alors qu'en 1976, Keeney et Raiffa élargissent la théorie de l'utilité (ou de la valeur) au cas multidimensionnel (Keeney et Raiffa, 1976). Ces deux courants de pensée vont générer deux conceptions méthodologiques différentes pour l'AMCD : l'école européenne et l'école anglo-saxonne.

D'un point de vue méthodologique, les outils mathématiques générés par les deux courants diffèrent fortement. D'un côté, l'école européenne s'est développée autour du concept de relation de surclassement, où la recommandation de décision est construite à partir de comparaisons par paires des différentes options de décision. D'un autre côté, l'école anglo-saxonne se base sur la notion d'utilité ou de valeur dans la théorie de l'utilité ou de la valeur multiattribut (MAVT) afin d'obtenir, par agrégation, une comparabilité totale des alternatives.

Notons X l'ensemble fini des n alternatives de décision. Elles représentent les options possibles au sujet desquelles un décideur doit prendre une décision. La définition de cet ensemble est en général assez délicate, et doit permettre au décideur de mieux appréhender le problème de décision, en clarifiant ses contours. Vincke (1989) observe que " X ne s'impose généralement pas comme une réalité objective facile à cerner". Il en découle que les choix de modélisation pour X ont une influence non négligeable sur l'ensemble des étapes du processus d'AMCD.

Roy (1985) constate que l'objectif d'un processus d'AMCD est de résoudre une des trois problématiques suivantes : déterminer une alternative pouvant être considérée comme la meilleure (choix), affecter les alternatives à des catégories prédéfinies (tri), classer les alternatives de la meilleure à la plus mauvaise (classement). Ces problématiques de décision se distinguent principalement par le fait que les jugements se font de manière absolue (tri) ou relative (choix et classement).

Dans le cadre multidimensionnel dans lequel se situent ces travaux, les conséquences de la mise en oeuvre d'une alternative sont multiples. Elles traduisent la diversité des caractéristiques qui interviennent dans la comparaison des alternatives. Un critère est principalement une fonction qui sert à résumer les évaluations d'une alternative sur diverses conséquences relativement homogènes qui se rattachent à un même point de vue général. Roy et Bouyssou (1993) définissent ainsi un *critère* comme une fonction à valeur réelle définie sur l'ensemble X des alternatives et qui permet de comparer deux alternatives x et y de sorte que

$$g(x) \geq g(y) \Rightarrow \begin{array}{l} \text{"}x \text{ est au moins aussi bon que } y \\ \text{relativement au point de vue considéré"} \end{array}.$$

La valeur numérique prise par une alternative sur un critère est appelée sa performance ou son évaluation, et le tableau regroupant toutes les évaluations des alternatives sur les critères est appelé *tableau de performance*.

La notion de critère sous-entend que la comparaison qualitative des alternatives se fait par rapport au point de vue du décideur. Il est donc essentiel de prendre en compte son système de valeurs et ses objectifs à travers ce qu'on appelle la modélisation de ses *préférences*. Pour ce faire, on fait en général appel à des relations binaires afin de représenter la manière avec laquelle deux alternatives se comparent. La théorie classique basée sur le concept d'utilité ou de valeur (Fishburn, 1970) distingue deux relations différentes pour cette comparaison, l'indifférence et la préférences, toutes deux étant considérées comme transitives. L'école européenne d'AMCD a enrichi cette dichotomie par l'introduction d'une relation d'incomparabilité (impossibilité pour le décideur de trancher) et de préférence faible (hésitation entre l'indifférence et la préférence) (Roy et Vanderpooten, 1996), aucune des relations n'étant supposée transitive.

L'information préférentielle qu'un analyste peut espérer obtenir du décideur peut en général prendre deux formes. D'un côté, il s'agit d'information qu'on qualifiera de *directe* lorsqu'il s'exprime par exemple sur des différences d'évaluation non significatives sur un critère, des coalitions de critères suffisantes pour emporter une comparaison entre deux alternatives ou des taux de substitution entre deux critères. D'un autre côté on parle d'information *indirecte* lorsque que le décideur fait référence aux résultats attendus du modèle d'aide à la décision en comparant deux alternatives entre elles ou en donnant des exemples d'affectation d'alternatives à des catégories.

Un processus d'AMCD vise en général à fournir une recommandation au décideur. Dans ce contexte multidimensionnel, cela sous-entend qu'une procédure d'agrégation doit être mise en oeuvre afin de synthétiser les préférences exprimées par le décideur sur chacun des critères. Mousseau (2003) définit une telle *procédure d'agrégation multicritère* par une règle ou un procédé qui permet d'établir, sur la base des performances des alternatives sur les critères et de paramètres préférentiels, un ou plusieurs systèmes relationnels permettant de comparer les alternatives de X . L'exploitation de ce système relationnel permet alors de répondre à la problématique adoptée par une recommandation de décision.

Les valeurs des *paramètres préférentiels* sont quant à eux déterminées via un processus d'*élicitation des préférences*. Ce dernier consiste en une interaction entre l'analyste et le décideur et amène ce dernier à exprimer une information sur ses préférences dans le cadre de la procédure d'agrégation qui a été choisie pour les modéliser. Cette information se traduit en un ensemble de valeurs plausibles pour les paramètres préférentiels de la procédure d'agrégation (Mousseau, 2003).

Les deux principales écoles méthodologiques pour l'AMCD fournissent différentes procédures d'agrégation multicritère. Du côté de l'école européenne, le principe général des méthodes de *surclassement* est de comparer les alternatives par le biais d'une relation majoritaire qui se base sur des différences d'évaluations qui sont comparées à des seuils de discrimination (indifférence, préférence, veto) pour chacun des critères. Une agrégation de ces comparaisons locales est ensuite effectuée, dans laquelle chacun des critères a un *pouvoir de vote* qui s'exprime par un coefficient d'importance. L'opérationnalisation de ces comparaisons par paires peut se faire de différentes façons (Roy et Bouyssou, 1993; Brans et Mareschal, 2002; Bisdorff *et al.*, 2008) et mène au concept de relation de surclassement. Cette relation n'étant en général pas transitive, une deuxième étape d'exploitation est nécessaire afin d'aboutir à une recommandation de décision.

Du côté de l'école anglo-saxonne, le but est de construire une représentation numérique

des préférences du décideur sur X à travers une fonction de valeur $U : X \rightarrow \mathbb{R}$ telle que

$$x \text{ est au moins aussi bon que } y \iff U(x) \geq U(y), \quad \forall x, y \in X.$$

Une nouvelle fois, la définition de U peut prendre différentes formes, comme par exemple une somme pondérée, l'intégrale de Choquet ([Grabisch, 1996](#)) ou la fonction de valeur additive ([von Winterfeldt et Edwards, 1986](#), chapitre 8). La relation de préférence induite par cette fonction de valeur est nécessairement un préordre, et son exploitation en vue d'une recommandation est assez aisée.

Cette brève introduction au domaine de l'Aide Multicritère à la Décision montre bien la diversité des sujets de recherche qui peuvent être traités. Sans être exhaustif, on peut donc y trouver par exemple la proposition de nouvelles procédures d'agrégation multicritère, leur caractérisation mathématique, l'étude du comportement du décideur en vue de proposer la procédure d'agrégation appropriée, l'élaboration de nouvelles procédures d'élicitation des préférences, la recherche de recommandations robustes, la formalisation du processus d'AMCD et le développement d'outils informatiques qui lui viennent en soutien, l'application des méthodes d'AMCD à d'autres champs scientifiques ou à des problèmes décisionnels de la vie réelle ...

1.2 Motivations et objectifs scientifiques de nos recherches

Dans cette section, nous commençons par présenter des observations qui ont initié plusieurs de nos travaux. Nous présentons ensuite les objectifs scientifiques qui en découlent et qui structurent notre recherche, ainsi que ce mémoire.

Les aspects liés à la mise en œuvre pratique de résultats de recherche nous paraissent essentiels afin que ceux-ci puissent avoir un intérêt pour des applications réelles, et pour faciliter la tâche de l'analyste dans le processus d'aide à la décision. En conséquence, nous avons en général veillé à ce que les résultats plus théoriques de nos recherches aient aussi une déclinaison dans un outil informatique de soutien au processus d'AMCD.

De nombreux travaux existent sur les caractérisations mathématiques des procédures d'agrégation multicritère ([Krantz et al., 1971](#); [Bouyssou et Marchant, 2007a](#), entre autres). Il existe aussi de nombreuses procédures qui argumentent leur validité par des justifications méthodologiques ou par plusieurs applications réussies à des cas réel. Ces deux voies sont intéressantes et essentielles pour comprendre les situations dans lesquelles il est préférable d'appliquer une procédure d'agrégation donnée. Cependant, comme cela se fait classiquement en intelligence artificielle ou en fouille de données, nous pensons que des tests empiriques de procédures, à l'aide d'un grand nombre d'instances de problèmes générées artificiellement, doivent donner des indications précieuses sur leur validité et leur utilisation en pratique.

Les deux principaux courants méthodologiques de l'AMCD ont tendance à être opposés. Nous préférons ne pas souscrire à cette vision et estimons que chacune des écoles a ses avantages et inconvénients. De manière générale, ce sont la situation de décision réelle combinée au comportement du décideur qui doivent guider l'analyste dans le choix de la méthodologie à appliquer.

Le processus d'AMCD, en tant que processus métier de l'analyste, a à ce jour été très peu modélisé et formalisé. Des travaux comme ceux de [Tsoukiàs \(2007\)](#) et de [Franco et Montibeller \(2010\)](#) sont des tentatives intéressantes. La formalisation du processus d'AMCD peut cependant être poussée plus loin en utilisant des concepts et méthodologies

issus de l'ingénierie dirigée par les modèles, qui vise à créer et à exploiter des connaissances et des activités d'un domaine métier.

Jusqu'il y a peu, le paysage des outils informatiques pour l'AMCD était composé d'une diversité de logiciels totalement indépendants, ayant tous leurs spécificités d'interaction et souvent incapables de partager des données entre eux. Dans d'autres domaines scientifiques, comme par exemple les statistiques ou la fouille de données, il existe des plateformes qui permettent de tester des méthodes d'analyse différentes dans un cadre commun (le système statistique R du [R Development Core Team \(2005\)](#) et le logiciel Weka de [Hall *et al.* \(2009\)](#)). Beaucoup de chercheurs du domaine sont convaincus qu'un tel outil, appliqué à l'AMCD, permettrait clairement d'augmenter le rayon de dissémination de leurs résultats, et contribuerait à une adoption plus large des techniques d'aide à la décision.

En partant de ces cinq observations, voici les objectifs qui ont guidé et qui guident encore actuellement nos travaux de recherche :

1. Proposer des procédures d'élicitation des préférences ou d'agrégation multicritère répondant à des besoins issus de réflexions méthodologiques ou d'applications réelles et dont l'intérêt est garanti minimalement par une validation empirique, en les confrontant à un grand nombre de situations de décision générées artificiellement ;
2. Fournir des recommandations d'utilisation des procédures proposées à partir de leur validation empirique ;
3. Proposer un écosystème cohérent d'outils informatiques de soutien au processus d'AMCD, englobant des techniques issues des deux écoles méthodologiques, et, proposer un cadre formel pour la pratique du processus d'AMCD ;
4. Appliquer les concepts et outils issus de l'AMCD à des problèmes réels et promouvoir leur utilisation dans d'autres domaines scientifiques.

Ces objectifs nous mènent à présenter nos travaux dans ce mémoire à travers trois axes thématiques :

1. Au chapitre 2, nos contributions aux aspects algorithmiques et méthodologiques de l'AMCD sont présentés, via une étude empirique de l'*expressivité* de différentes procédures d'agrégation multicritère, la définition et l'étude de la problématique de la classification non supervisée en AMCD, la proposition d'une aide au tri multicritère en présence de décideurs multiples et l'élicitation stable de paramètres préférentiels dans un contexte de surclassement. Tous ces résultats sont soumis à des validations empiriques en les confrontant à un grand nombre de problèmes de décision générés artificiellement, afin notamment d'en extraire des recommandations d'utilisation ;
2. Ensuite, au chapitre 3, nous nous focalisons sur des outils de soutien au processus d'AMCD, en détaillant l'initiative *diviz* ainsi qu'une première tentative de modélisation du processus métier de l'analyste, à l'aide de l'ingénierie dirigée par les modèles ;
3. Pour terminer, au chapitre 4 nous présentons des applications de l'AMCD à d'autres domaines scientifiques, en particulier l'amélioration d'indicateurs de bien-être en économie et la construction d'échelles de risque territorial.

Chapitre 2

Aspects algorithmiques et méthodologiques de l'AMCD

Sommaire

2.1	Expressivité empirique de modèles d'agrégation	8
2.1.1	Cadre expérimental	9
2.1.2	Résultats et observations	11
2.1.3	Conclusions	12
2.2	Formalisation du concept de classification non supervisée et méthode de résolution	13
2.2.1	Définitions et objectifs pour une classification non supervisée	14
2.2.2	Méthode de résolution	15
2.2.3	Illustration et validation empirique	17
2.2.4	Conclusions	19
2.3	Aide au tri en présence de critères et de décideurs multiples	20
2.3.1	Electre Tri et méthodes d'élicitation des préférences connexes	21
2.3.2	Inférence des profils des catégories à partir d'exemples d'affectation	23
2.3.3	Processus d'aide au tri pour un groupe de décideurs	23
2.3.4	Exemple illustratif et validation empirique	24
2.3.5	Conclusions et perspectives	27
2.4	Elicitation stable des poids d'une relation de surclassement	28
2.4.1	Au sujet de la stabilité de la relation de surclassement	29
2.4.2	Algorithme de détermination stable des poids	31
2.4.3	Validation empirique et processus d'élicitation	32
2.4.4	Conclusions et perspectives	34

Ce chapitre présente quatre apports algorithmiques et méthodologiques au domaine de l'AMCD. Ces travaux très différents au premier abord se rejoignent sur un point très important à nos yeux : la confrontation des algorithmes à des données générées artificiellement.

En premier lieu, cette étape sert à valider de manière empirique les procédures. Elle peut paraître tout à fait évidente et essentielle, mais n'est pas forcément naturelle dans les travaux d'AMCD. En fouille de données ou en intelligence artificielle par contre, toutes les méthodes proposées sont confrontées soit à des problèmes classiques (benchmarks) ou à des données artificielles.

Ensuite, en analysant leur comportement face à ces problèmes de décision artificiels, cette étape sert à fournir des recommandations sur l'utilisation des procédures dans des situations réelles. Ainsi, pour des procédures d'élicitation des préférences, ces simulations de comportement peuvent notamment fournir des indications sur le volume d'information à obtenir d'un décideur, avant d'arriver à un modèle "stable" et "représentatif" des préférences du décideur.

Cette validation empirique représente ainsi le fil rouge de ces travaux et servira à confirmer l'intérêt des algorithmes, ainsi qu'à en déduire des recommandations de bonne pratique pour la résolution de problèmes de décision réels.

Tout d'abord nous présentons en section 2.1 des résultats sur l'expressivité de différents modèles d'agrégation multicritères, c'est à dire leur capacité à représenter des préférences de décideurs (Pirlot *et al.*, 2010; Meyer et Pirlot, 2012). Ensuite, la section 2.2 s'intéresse plus particulièrement à la problématique de la classification non supervisée en AMCD, dont l'objectif est de regrouper des alternatives de décision que le décideur considère comme indifférentes, et de séparer celles qui ne le sont pas à ses yeux (Meyer et Olteanu, 2013). Les deux sections suivantes traitent d'élicitation des préférences. Tout d'abord, la section 2.3 détaille une méthode d'aide au tri d'alternatives dans un contexte où plusieurs décideurs interviennent dans le processus d'évaluation (Cailloux *et al.*, 2012). Ensuite, nous présentons des travaux sur l'élicitation stable de paramètres préférentiels (Bisdorff *et al.*, 2009, 2013) en section 2.4.

2.1 Expressivité empirique de modèles d'agrégation

Dans le contexte de l'école anglo-saxonne pour l'AMCD, la fonction de valeur additive occupe une position dominante dans les modèles utilisés pour représenter les préférences d'un décideur sur des alternatives décrites par plusieurs critères (Keeney et Raiffa, 1976; Bouyssou et Pirlot, 2004). Ce modèle implique l'indépendance préférentielle (Keeney et Raiffa, 1976, page 32), ce qui ne signifie cependant pas qu'il peut être utilisé pour représenter toutes les préférences qui satisfont cette condition.

Un exemple simple pour montrer que cette condition d'indépendance peut être facilement violée est celui d'un commerçant qui veut louer un nouveau local commercial. Les locaux sont évalués par rapport à leur surface, le montant de la location et le quartier dans lequel ils sont situés. Le commerçant estime que dans un quartier commercialement attractif, il préfère le local qui est grand et cher par rapport à celui qui est petit et économique, alors que dans un quartier peu intéressant pour ses affaires, il préfère le local petit et économique à celui qui est grand et cher. Cette situation parfaitement plausible ne satisfait pas la condition d'indépendance préférentielle et ne peut pas être représenté par un modèle de fonction de valeur additive.

Dans les années 1980, le modèle basé sur l'intégrale de Choquet (Schmeidler, 1986) a émergé pour représenter des préférences dans lesquelles un certain type d'interaction entre les critères doit être pris en compte (Grabisch, 1996).

Le type d'interaction qui peut être modélisé par l'intégrale de Choquet s'illustre facilement via l'exemple du classement d'étudiants de gestion de Grabisch et Labreuche (2004). Ces étudiants sont évalués par leurs performances passées en mathématiques, statistiques et langues. Étant donné que les notes de mathématiques et de statistiques sont corrélées, le jury qui doit classer les étudiants estime que pour ceux qui sont bons en mathématiques, il préférera ceux qui sont bons en langues à ceux qui sont bons en statistiques. Cependant

dans le cas opposé, pour les étudiants mauvais en mathématiques, le jury préférera ceux qui sont bons en statistiques à ceux qui sont bons en langues. Dans cette situation, il existe une interaction (négative) entre les critères mathématiques et statistiques et l'indépendance préférentielle est clairement violée.

Il est cependant intéressant de noter que l'interaction entre les critères est une propriété qui ne doit pas être identifiée avec la violation de l'indépendance préférentielle : en effet, des préférences représentables par une intégrale de Choquet ne satisfont pas nécessairement la condition d'indépendance préférentielle, mais comme le montre Wakker (1989, page 111) elles satisfont une forme d'indépendance plus faible comme l'indépendance comonotone ainsi que la séparabilité faible (Bouyssou et Pirlot, 2004).

L'objectif de ces travaux, menés en collaboration avec Marc Pirlot¹, est d'étudier la capacité de différentes procédures d'agrégation à représenter les préférences de décideurs, dans le cadre de la théorie de la valeur multiattribut (MAVT).

Nous appelons cette capacité de représentation l'*expressivité* d'un modèle. Ainsi, un modèle très expressif aura tendance à être approprié pour représenter les préférences d'un décideur pris au hasard, alors qu'un modèle peu expressif aura de fortes chances de n'être utilisable que dans des situations très spécifiques.

Ces travaux tentent en particulier de répondre aux questions suivantes :

1. Quelle est l'expressivité de l'intégrale de Choquet comparée à celles du modèle additif général et de la somme pondérée ?
2. Quelle est l'expressivité du modèle additif général comparée à celle de modèles linéaires par morceaux ?
3. Comment se comporte l'intégrale de Choquet quand elle est confrontée à des préférences non indépendantes ?

Nous supposons dans cette étude que l'intégrale de Choquet est calculée directement sur le vecteur des performances des alternatives, sans passer par un recodage de ces performances par des fonctions de valeur. Le modèle additif général au contraire utilise ces fonctions de valeur potentiellement non linéaires pour représenter les préférences du décideur.

Afin de déterminer l'expressivité de ces modèles et de répondre aux questions posées, nous effectuons une série d'expériences qui sont décrites brièvement dans la section 2.1.1. Les résultats et conclusions que nous avons obtenus sont ensuite présentés en section 2.1.2. Le lecteur intéressé pourra trouver plus de détails, ainsi que des tableaux résumant les résultats, dans les deux articles (Pirlot *et al.*, 2010; Meyer et Pirlot, 2012).

2.1.1 Cadre expérimental

Soit $(a_1, \dots, a_m)$ le vecteur de performances de l'alternative a de X . Supposons sans perte de généralité que ces performances ont leurs valeurs dans l'intervalle $[0, 1]$. Les modèles qui sont testés dans ce travail sont :

1. La *somme pondérée* pour laquelle le score de l'alternative a est $\sum_{i=1}^m w_i a_i$, où w_i sont des poids qui doivent être déterminés.

1. Marc Pirlot, Université de Mons, Belgique

2. La *fonction de valeur additive* pour laquelle le score de l'alternative a est $\sum_{i=1}^m u_i(a_i)$, où les u_i sont les fonctions de valeur marginales de $[0, 1]$ vers lui-même qui doivent être déterminées. Deux cas particuliers de ce modèle général sont aussi étudiés, dans lesquels les fonctions marginales considérées sont des fonctions linéaires par morceaux. Dans la première variante, u_i a deux morceaux linéaires correspondant à une division de l'intervalle $[0, 1]$ en deux parties égales. La seconde variante quant à elle divise l'intervalle $[0, 1]$ en trois parties égales pour considérer des fonctions à 3 morceaux linéaires.
3. L'*intégrale de Choquet* dans laquelle le score de l'alternative a est calculé en utilisant des capacités 2-additives et 3-additives. Dans le cas de capacités 2-additives, le score de a s'obtient par la formule : $\sum_{i=1}^m m_i a_i + \sum_{i,j,i < j} m_{ij}(a_i \wedge a_j)$, où m_i (resp. m_{ij}) sont les poids associés aux critères (resp. aux paires de critères) et $a_i \wedge a_j$ indique le minimum de a_i et de a_j . Ces poids sont évidemment régis par un ensemble de contraintes qui sont par exemple détaillées dans [Grabisch et al. \(2008\)](#). La formule pour des capacités 3-additives est similaire et contient un terme additionnel de poids m_{ijk} associé aux triplets de critères.

En vue d'étudier les questions mentionnées précédemment, les deux cadres expérimentaux suivants sont mis en place. Tout d'abord, pour tester l'expressivité des différents modèles, pour un nombre n d'alternatives et un nombre m de critères donnés, nous générons aléatoirement n vecteurs de performances à m composantes. Ces performances sont générées aléatoirement à partir d'une distribution uniforme sur l'intervalle $[0, 1]$. Nous supposons aussi, sans perte de généralité, que le décideur préfère les performances élevées aux performances basses pour tous les critères. Un ordre total est ensuite généré pour les n alternatives à partir d'une distribution uniforme de tous les ordres totaux. Des routines de programmation linéaire sont alors utilisées pour vérifier que cet ordre est représentable par le biais des différents modèles retenus. Dans chacun de ces modèles, l'alternative a est préférée à l'alternative b si son score est plus élevé que celui de b .

Ensuite, pour tester le comportement de l'intégrale de Choquet quand elle est confrontée à des préférences non indépendantes, des données et des ordres appropriés sont générés d'après un cadre similaire au précédent. Nous générons d'abord aléatoirement les performances de deux alternatives a et b . Ensuite nous tirons au hasard un indice de critère q dans $\{1, \dots, m\}$ et générons aléatoirement deux évaluations x_q et y_q . Nous construisons ensuite les quatre vecteurs de performances (x_q, a_{-q}) , (x_q, b_{-q}) , (y_q, a_{-q}) et (y_q, b_{-q}) , où la notation (x_q, a_{-q}) représente le profil de l'alternative a pour laquelle l'évaluation a_q a été remplacée par l'évaluation x_q . Nous générons ensuite les $n - 4$ alternatives comme dans le cadre précédent. Les ordres sont générés de façon à ce que (x_q, a_{-q}) soit classé avant (x_q, b_{-q}) et (y_q, a_{-q}) soit classé après (y_q, b_{-q}) , ce qui garantit que cet essai viole la condition d'indépendance préférentielle. En pratique nous générons un ordre aléatoire sur les n alternatives et nous le retenons si la condition précédente sur les quatre profils est respectée ; si non, l'ordre est rejeté.

Toutes ces expériences sont effectuées dans le système statistique R ([R Development Core Team, 2005](#)). Les paramètres pour les deux premiers modèles sont déterminés en utilisant un solveur qui peut être appelé à partir de R ([lpSolve](#)). La structure des contraintes pour les fonctions de valeur linéaires par morceaux sont inspirées de la technique d'agrégation-désagrégation UTA de [Jacquet-Lagrèze et Siskos \(1982\)](#), alors que la formulation pour tester l'existence d'une fonction de valeur générale est de [Greco et al. \(2008\)](#). Pour vérifier l'existence d'une représentation des préférences à travers une intégrale de Choquet

2-additive ou 3-additive, la librairie Kappalab de [Grabisch et al. \(2008\)](#) est utilisée, via la fonction “lin.prog.capa.ident”. Dans tous les modèles, la routine d’optimisation tente de maximiser une variable d’écart δ qui représente la différence minimale entre deux alternatives qui se suivent dans le classement aléatoire.

Nous explorons les cas où $n = 4, 5, 6, 8, 10$ et pour chaque valeur de n un nombre de critères égal à $m = 3, 4, 5, 6, 8$. Pour des petites valeurs de n (4 à 6), nous générons entre 50 et 100 instances de n alternatives et nous essayons de représenter systématiquement tous les ordres possibles (pour le deuxième cadre expérimental, nous ne considérons évidemment que les ordres qui violent la condition d’indépendance préférentielle). Pour $n = 8$ et $n = 10$, nous générons 50 instances et ne considérons qu’un échantillon de 3000 ordres.

2.1.2 Résultats et observations

Nous ne présentons ici que les observations principales qui peuvent être déduites des résultats des expériences que nous avons menées. Pour des résultats détaillés, le lecteur intéressé peut se référer à [Meyer et Pirlot \(2012\)](#).

Tout d’abord, en ce qui concerne l’expressivité, nous observons que :

1. Pour tous les modèles, la proportion d’instances qui peuvent être représentées croît avec le nombre de critères (pour n fixé) et décroît avec le nombre d’alternatives (pour m fixé).
2. Le modèle additif est plus expressif que l’intégrale de Choquet 3-additive ; cette dernière étant légèrement plus expressive que le modèle 2-additif ; qui lui-même est plus expressif que la somme pondérée. Ces différences sont d’autant plus marquées que n est grand.
3. Les différences d’expressivité entre le modèle additif général et ses variantes linéaires par morceaux croissent avec n . Pour des valeurs de n inférieures à 6 cette différence reste petite. Pour une valeur fixée de n cette différence croît avec le nombre de critères.

Il est également intéressant de comparer l’expressivité des différents modèles. A cette fin, pour chaque modèle, nous calculons l’amélioration qu’il apporte par rapport aux autres modèles en terme de représentativité des ordres. Nous en déduisons les observations suivantes :

1. Dans très peu de cas (moins d’1% des instances que nous avons générées), les modèles à base d’intégrale de Choquet 2-additive ou 3-additive sont capables de représenter un ordre que le modèle additif général ne peut pas modéliser.
2. L’avantage d’une intégrale de Choquet 3-additive par rapport à une intégrale de Choquet 2-additive n’est pas flagrant jusqu’à $n = 8$.
3. L’intégrale de Choquet est capable de représenter un grand nombre d’ordres que la somme pondérée ne peut pas modéliser.

Il peut paraître assez surprenant que l’intégrale de Choquet n’est capable de représenter que quelques instances qui ne sont pas modélisables par un modèle additif général. La question que l’on peut alors se poser concerne la proportion de préférences qui violent l’indépendance que l’intégrale de Choquet peut représenter. Le deuxième cadre expérimental que nous avons décrit à la section 2.1.1 nous permet de faire les observations suivantes à ce sujet :

1. Le pourcentage d'ordres qui violent la condition d'indépendance, mais représentables par une intégrale de Choquet 2-additive ou 3-additive, est plus élevé que ce que la première expérience ne laissait penser (en moyenne une dizaine de pourcents).
2. La proportion de préférences représentables par ces deux modèles augmente avec le nombre de critères (pour n fixé) et décroît avec le nombre d'alternatives (pour m fixé).
3. Le passage d'un modèle 2-additif à un modèle 3-additif ne résulte pas en une augmentation flagrante de l'expressivité.

Ces observations confirment un certain nombre d'hypothèses "naturelles", mais sont également accompagnées de nouvelles questions. En particulier, on peut se poser des questions sur la génération de nos données, et plus spécifiquement sur les données qui violent l'indépendance préférentielle. Il est clair que l'approche choisie pour le deuxième cadre expérimental n'est pas garante d'un échantillonnage uniforme des préférences qui ne sont pas représentables par une fonction de valeur additive générale. Nous en concluons donc que l'intégrale de Choquet est capable de modéliser certaines préférences que la fonction de valeur additive générale ne peut pas représenter, mais nous préférons ne pas nous avancer, dans l'état actuel de nos recherches, sur une estimation numérique de cette proportion.

L'intégrale de Choquet peut aussi être utilisée à la place de la somme dans le modèle de fonctions de valeur additive (Grabisch et Labreuche, 2004). Ceci aurait évidemment comme conséquence d'accroître l'expressivité de ce dernier, pour laquelle nous donnons une estimation dans Meyer et Pirlot (2012).

2.1.3 Conclusions

A la vue de ces résultats, il pourrait être tentant de recommander de restreindre l'utilisation de l'intégrale de Choquet à des cas où les fonctions de valeur ont été élicitées et où il y a des indices forts d'interactions entre des critères.

Cependant, nous pensons qu'il faut être un peu moins restrictif en se rappelant que les modèles les plus expressifs ne sont pas forcément le meilleur choix dans un contexte d'apprentissage des préférences, car ils nécessitent la spécification d'un très grand nombre de paramètres. En effet, si les informations préférentielles fournies par le décideur sont pauvres, l'utilisation d'un modèle moins expressif, comme l'intégrale de Choquet 2-additive, peut être conseillée, car elle devrait fournir un degré d'indétermination des paramètres plus bas qu'avec une fonction de valeur additive générale.

Notons également que de bonnes alternatives au modèle additif général sont les modèles linéaires par morceaux, qui nécessitent la détermination de moins de paramètres, tout en ayant une bonne expressivité.

*

Dans la section suivante, nous présentons des travaux sur la classification non supervisée en AMCD, une problématique de décision assez peu étudiée à ce jour, et qui peut se rapprocher des problématiques de tri et de rangement, suivant la perspective adoptée. Une nouvelle fois, nous soulignerons des aspects de validation empirique.

2.2 Formalisation du concept de classification non supervisée et méthode de résolution

En analyse des données, il est classique de rassembler des objets décrits par une série de variables afin de mieux comprendre la structure du problème sous-jacent. Cette action de regrouper est très naturelle pour l'humain, car elle constitue notre façon d'appréhender un nouveau concept, en tentant de le rapprocher d'autres éléments de notre champs de connaissance (Anderberg, 1973; Everitt *et al.*, 2001). En analyse des données, la classification supervisée et la classification non supervisée sont deux approches générales pour résoudre ces problèmes de regroupement d'objets.

La classification supervisée se base sur des informations concernant les groupes d'objets, appelés des classes. Dans ce cas, le but est de tenter de faire correspondre un certain modèle aux données et de le valider ensuite. La classification non supervisée quant à elle ne bénéficie pas de cette information a priori sur la structure des données, et a comme objectif de la découvrir, en utilisant des mesures de similarité.

Un des atouts principaux de l'AMCD est de tenter d'inclure des préférences dans la comparaison des alternatives et par conséquent, l'information disponible est plus riche qu'en analyse des données. De manière générale on distingue entre trois types de relations entre ces alternatives : l'indifférence, la préférence stricte et l'incomparabilité (Vincke, 1989). La problématique du tri en AMCD peut d'une certaine manière se rapprocher de la classification supervisée en analyse des données, étant donné que de l'information sur les classes (ou catégories) est à fournir a priori. La problématique du classement en AMCD quant à elle se rapproche sous certains aspects de la classification non supervisée en analyse des données, puisque en général aucune information n'est donnée a priori sur la structure du problème.

Néanmoins, il reste que la problématique de la classification non supervisée en AMCD a jusqu'à présent été très peu abordée et formalisée. Cailloux *et al.* (2007) proposent une taxonomie de méthodes existantes, qui les sépare en deux catégories, suivant qu'elles prennent en compte de l'information préférentielle ou non.

Il convient aussi de citer Bisdorff (2002a) qui procède à une classification non supervisée des critères, et non des alternatives. Par ailleurs, De Smet et Guzman (2004) ont étendu l'algorithme classique des K-Means au contexte de l'AMCD. Une extension de ces travaux est proposée par la suite par De Smet et Eppe (2009) afin de générer des relations de préférence entre les classes générées. Cet algorithme des K-Means est aussi à la base des travaux de Figueira *et al.* (2004) et Baroudi et Safia (2010).

Dans Nemery et De Smet (2005) et Fernandez *et al.* (2010) les auteurs proposent des approches qui déterminent des classes totalement ordonnées, alors que plus récemment, dans Rocha *et al.* (2012) les classes peuvent être partiellement ordonnées.

Ce bref tour d'horizon constitue le point de départ de ce travail que nous avons mené en collaboration avec Alexandru Olteanu², lors de l'encadrement de sa thèse de doctorat. Dans la suite, nous présenterons dans la section 2.2.1 notre définition de la problématique de classification non supervisée en AMCD qui se base sur des relations de préférence entre les alternatives. Ensuite, dans la section 2.2.2 nous présenterons brièvement notre approche pour résoudre cette problématique, avant de passer à une illustration et la validation empirique en section 2.2.3. Le lecteur intéressé pourra se référer à Meyer et Olteanu (2013) pour plus de détails sur ces travaux.

2. Alexandru Olteanu, Université du Luxembourg, Luxembourg

2.2.1 Définitions et objectifs pour une classification non supervisée

La classification non supervisée en analyse des données a comme but de regrouper des objets qui se ressemblent et de séparer ceux qui ne se ressemblent pas. Dans le cadre de l'AMCD nous pouvons rajouter de l'information préférentielle à ces comparaisons d'objets afin de prendre en compte les préférences d'un décideur.

Il existe de nombreuses méthodes pour comparer des alternatives décrites sur un certain nombre de critères, et qui prennent en compte différents types d'informations préférentielles. Cependant, d'un point de vue général, ces procédures fournissent trois types de comparaisons : l'indifférence, la préférence stricte et l'incomparabilité (notées I, P et R).

La relation d'indifférence I est réflexive et symétrique, l'incomparabilité R est irreflexive et symétrique et la préférence stricte P est asymétrique. La méthode de construction de ces relation n'est pas l'objectif de ces travaux, et nous supposons donc qu'elles sont données (il est ainsi possible de les dériver d'une relation de surclassement, par exemple).

Étant donné qu'entre deux alternatives de X une et une seule de ces relations peut exister, nous pouvons adapter la définition de base de l'analyse des données au contexte de l'AMCD.

Définition 1. *La **classification non supervisée non relationnelle** en AMCD est le processus qui regroupe les alternatives qui sont indifférentes entre elles et qui sépare celles qui ne le sont pas.*

Il est intéressant de noter que dans ce cas les relations entre alternatives sont divisées en deux : d'un côté celle qui rassemble les alternatives (l'indifférence), et de l'autre côté celles qui les séparent (préférence stricte et incomparabilité). En analyse des données, chacun de ces ensembles ne contient qu'une relation, et elles sont complémentaires (similarité et dissimilarité). En AMCD, le "complémentaire" de l'indifférence est soit de la préférence stricte dans un sens ou dans l'autre ou de l'incomparabilité.

Suivant l'objectif de la classification, on peut être tentés d'exploiter la richesse des relations qui séparent les alternatives pour définir d'autres résultats.

Définition 2. *Nous appelons **classification non supervisée relationnelle** en AMCD le processus qui regroupe les alternatives qui sont indifférentes entre elles, tout en séparant les classes qui sont strictement préférées ou incomparables entre elles.*

Dans ce cas, on ne cherche pas que des ensembles d'alternatives clairement indifférentes, mais aussi des classes qui se comparent entre-elles via une relation de préférence stricte ou d'incomparabilité. Il est aussi possible d'être plus discriminant au sujet de ces relations entre les classes, et de rechercher un **tournoi partiel** ou un **tournoi complet** entre les classes.

Un résultat encore plus structuré est celui de la relation d'ordre entre les classes, où les ensembles d'alternatives indifférentes sont ordonnées du meilleur au plus mauvais.

Définition 3. *On appelle **classification non supervisée ordonnée** en AMCD le processus qui regroupe les alternatives qui sont indifférentes entre elles, tout en séparant les classes strictement préférées ou incomparables entre elles, de manière à créer un ordre entre les classes.*

Il est à nouveau possible de distinguer dans ce cas entre deux structures pouvant apparaître entre les classes : un **ordre partiel strict** et un **ordre total strict**.

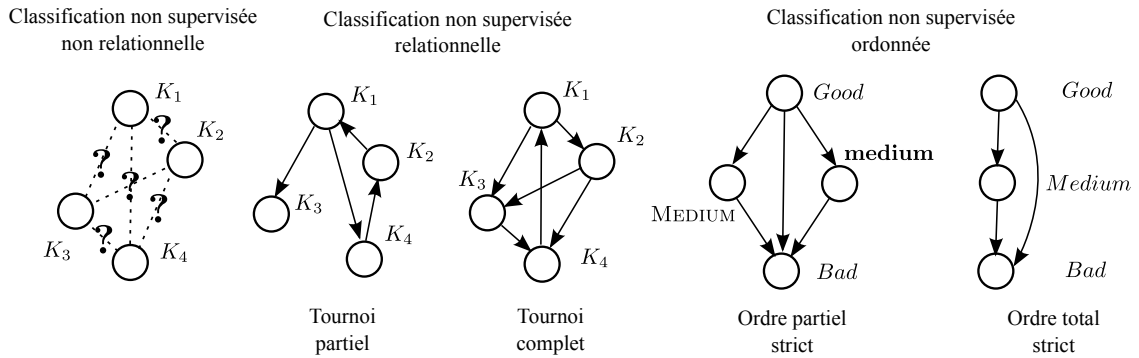


FIGURE 2.1 – Structures potentielles du résultat d’une classification non supervisée en AMCD.

La figure 2.1 résume ces différentes structures qui peuvent émerger d’une classification non supervisée en AMCD.

Afin de rendre ces différentes définitions opérationnelles, nous construisons pour chacune d’elles une fonction objectif (à maximiser). Pour une définition donnée, cette fonction permet de qualifier un résultat de classification. La valeur maximale de cette fonction est alors atteinte lorsque la partition est “idéale” pour la définition choisie.

De manière générale, le but de la classification non supervisée est de déterminer une partition $K = \{K_l, K_m, K_n, \dots\}$ de l’ensemble des alternatives X ayant certaines propriétés suivant la définition choisie.

Pour la définition la moins restrictive, c’est à dire celle de la classification non supervisée non relationnelle, le résultat idéal serait d’avoir une partition pour laquelle les alternatives à l’intérieur d’une classe seraient toutes indifférentes entre elles, et aucune relation d’indifférence n’apparaîsse entre des alternatives appartenant à des classes différentes.

Soit $S_O(A, B)$ la fonction qui compte le nombre de relations $O \in \{I, P, R\}$ qui apparaissent entre deux ensembles d’alternatives A et B de X . La fonction objectif pour la classification non supervisée non relationnelle, à maximiser, est donc donnée par :

$$f_{NR}(K) := \sum_l S_I(K_l, K_l) + \sum_{l < m} (S_P(K_l, K_m) + S_P(K_m, K_l) + S_R(K_l, K_m)). \quad (2.1)$$

Le premier terme de f_{NR} mesure à quel point la relation d’indifférence est présente à l’intérieur des classes, alors que le second terme évalue l’existence des autres relations entre les classes. L’optimum est atteint lorsque $f_{NR}(K) = \frac{|X| \cdot (|X| - 1)}{2}$, et dans ce cas K est une partition idéale par rapport à l’objectif de la classification non supervisée non relationnelle.

Les autres fonctions objectifs se définissent de manière similaire. Le lecteur intéressé pourra se référer à [Meyer et Olteanu \(2013\)](#) pour de plus amples détails.

2.2.2 Méthode de résolution

Afin de déterminer une solution au problème de la classification non supervisée non relationnelle, nous pourrions énumérer toutes les partitions de l’ensemble des alternatives X

et sélectionner la partition qui maximise f_{NR} . La complexité d'une telle approche est clairement exponentielle, étant donné que le nombre de partitions d'un ensemble est donné par le nombre de Bell (pour 10 alternatives, le nombre de partitions est d'environ 10^5). Ainsi, même dans le contexte de l'AMCD où le nombre d'alternatives est en général restreint, cette énumération exhaustive n'est pas raisonnable.

L'approche que nous proposons, qui se base sur nos travaux [Bisdorff et al. \(2011\)](#) en fouille de données, est composée de deux étapes : tout d'abord nous tentons de trouver une partition des alternatives en ne considérant que la relation d'indifférence ; ensuite nous raffinons ce résultat en déplaçant des alternatives entre les éléments de la partition afin de se rapprocher du résultat optimal.

Dans la première étape, nous tentons de créer des ensembles, à l'intérieur desquels les alternatives sont majoritairement indifférentes entre elles, et entre lesquels il existe peu d'indifférences. Cela réduit la complexité du problème original à la recherche d'ensembles denses dans le graphe représentant la relation d'indifférence, noté $G(X, I)$.

En vue de trouver ces régions denses, nous recherchons les *cœurs* des classes, que nous définissons comme des ensembles d'alternatives qui sont toutes indifférentes entre-elles et qui se comparent de manière consistante aux autres alternatives de X . Ceci signifie que nous aimerions que ces cœurs contiennent des alternatives qui sont soit indifférentes aux mêmes alternatives de X (ces dernières seront rajoutées à la classe définie par le cœur par la suite), soit non indifférentes aux mêmes alternatives de X (ces dernières seront placées dans d'autres classes). De cette définition découle le fait que les cœurs que nous recherchons sont des cliques dans $G(X, I)$.

Idéalement, soit un cœur Y ne présentera que des relations d'indifférences vers l'extérieur, soit il n'en présentera aucune. La fonction pour mesurer la qualité d'un cœur est donc définie comme :

$$f_C(Y) := \sum_{x \in X} |S_I(x, Y) - S_P(x, Y) - S_P(Y, x) - S_R(x, Y)|. \quad (2.2)$$

où $S_O(x, Y)$ compte le nombre de relations $O \in \{I, P, R\}$ de x vers les éléments de Y . Une valeur élevée pour $f_C(Y)$ indique que Y est en moyenne plutôt indifférent ou plutôt pas indifférent aux alternatives à l'extérieur de Y , ce qui traduit la propriété que nous recherchons pour un cœur.

Nous recherchons donc les cliques de $G(X, I)$ qui sont les maxima locaux de f_C , afin de déterminer le nombre de régions denses du graphe, et par conséquent le nombre de classes (que nous notons k).

Ceci peut être fait d'une manière exacte en utilisant l'algorithme de [Bron et Kerbosch \(1973\)](#), amélioration de [Koch \(2001\)](#), et qui énumère de manière efficace toutes les cliques maximales d'un graphe. Cette procédure est exponentielle en complexité à cause du fait qu'un graphe contient un nombre exponentiel de cliques maximales ([Moon et Moser, 1965](#)). Cependant, comme les problèmes en AMCD ne sont généralement pas très grands, cette approche pour rechercher les cœurs est bien adaptée ($|X| < 100$). Dans le cas où le nombre d'alternatives devient plus grand, nous proposons d'appliquer une meta-heuristique telle que celle que nous avons proposé dans [Bisdorff et al. \(2011\)](#). Nous considérerons ici que le problème est raisonnable en taille, et nous notons $C = \{C_1, \dots, C_k\}$ l'ensemble des cœurs déterminés.

Chaque classe est alors construite en lui ajoutant d'abord les éléments de son cœur, et ensuite en ajoutant les alternatives restantes à la classe la plus appropriée d'après

l'heuristique gloutonne suivante :

$$\begin{aligned} 1) & K_l = C_l, \forall l \in 1..k, \\ 2) & K_l = K_l \cup \{x\} : l = \arg \max_m (S_l(x, C_m)), \forall x \in X - \bigcup_i C_i, \end{aligned} \quad (2.3)$$

Cette première partition devrait déjà être assez proche de l'optimum de f_{NR} (comme nous le montrons empiriquement dans la section 2.2.3). Un raffinement de ce premier partitionnement est effectué à l'étape suivante, en fonction de l'objectif de la classification non supervisée recherché.

L'idée de cette deuxième étape est de déplacer des alternatives entre les classes en vue d'augmenter la valeur de la fonction objectif. A cette fin nous définissons des heuristiques pour chacun des objectifs de classification qui permettent d'évaluer si le déplacement entraîne une augmentation de la fonction objectif ou non. Nous invitons le lecteur intéressé à consulter [Meyer et Olteanu \(2013\)](#) pour une argumentation détaillée sur ces heuristiques.

Toute meta-heuristique de recherche locale peut alors être utilisée, comme par exemple la recherche tabou ([Glover, 1989, 1990](#)) ou le recuit simulé ([Kirkpatrick et al., 1983](#)) pour trouver la classification optimale. De manière générale, ce type d'algorithme commence avec une partition initiale K (trouvée à l'étape précédente), et génère ensuite son voisinage (l'ensemble des déplacements d'alternatives à d'autres classes). Ces déplacements sont évalués grâce aux heuristiques précédemment mentionnées, en fonction de l'objectif de classification retenu. Le meilleur voisin est alors choisi, et la mécanique itérative de la meta-heuristique se met en marche, suivant sa spécificité.

2.2.3 Illustration et validation empirique

L'algorithme est ici illustré sur des données issues de [Bouyssou et al. \(2000\)](#) concernant le choix d'une voiture par un étudiant. Pour rappel, le problème contient 14 alternatives (voitures) définies sur 5 critères : un critère de coût, deux critères de performance, et deux critères de sécurité.

Les alternatives sont comparées au moyen d'une relation de surclassement évaluée en utilisant la règle d'ELECTRE III ([Roy, 1978](#)). Les poids des critères sont égaux, et les seuils d'indifférence et de préférence sont fixés à 10% et à 20% de la longueur de l'intervalle des valeurs de chaque critère. Des seuils de veto ont été définis pour les critères de coût et d'accélération à 60% de la longueur de l'intervalle des valeurs. Une coupe de la relation évaluée a été effectuée à 0.5 afin d'en extraire les relations d'indifférence, de préférence stricte et d'incomparabilité, comme présenté dans [Vincke \(1992\)](#).

La figure 2.2 présente les sorties de l'algorithme pour des classifications non supervisées ordonnées où les classes sont ordonnées partiellement et totalement. Sur la gauche sont détaillées les relations entre les alternatives. Ces dernières ont été réordonnées afin de tenir compte du résultat de la classification. Les graphes sur la droite représentent les classes et les relations entre celles-ci. Afin de faciliter la représentation, les classes du haut sont considérées comme préférées à celles du bas et des classes au même niveau sont considérées comme incomparables. Les confiances en dessous de chaque tableau correspondent à f_{OPS} (ordre partiel strict) et à f_{OCS} (ordre complet strict).

Pour l'ordre partiel strict l'algorithme trouve cinq classes. Les classes $\{a8, a9, a14\}$ et $\{a6\}$ sont incomparables entre elles. Les cases grisées dans le tableau indiquent des

Ordre partiel strict

<i>K</i>	a3	a7	a11	a5	a10	a12	a8	a9	a14	a6	a1	a2	a4	a13
a3		<i>I</i>	<i>I</i>	<i>P</i>	<i>I</i>	<i>I</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a7			<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>R</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a11				<i>P</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a5					<i>I</i>	<i>I</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>I</i>	<i>I</i>	<i>P</i>
a10						<i>I</i>	<i>I</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>I</i>	<i>I</i>	<i>I</i>	<i>I</i>
a12							<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>P</i>
a8								<i>I</i>	<i>I</i>	<i>R</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a9									<i>I</i>	<i>R</i>	<i>P</i>			<i>I</i>
a14										<i>R</i>	<i>P</i>	<i>R</i>	<i>R</i>	<i>P</i>
a6											<i>P</i>	<i>I</i>	<i>I</i>	<i>I</i>
a1												<i>I</i>	<i>I</i>	<i>I</i>
a2									<i>P</i>				<i>I</i>	<i>I</i>
a4									<i>P</i>					<i>I</i>
a13														

Confiance du résultat : 72.52%

Ordre total strict

<i>K</i>	a3	a7	a8	a11	a12	a1	a2	a4	a5	a6	a10	a13	a9	a14
a3		<i>I</i>	<i>I</i>	<i>I</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>
a7			<i>I</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>R</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a8						<i>P</i>	<i>P</i>	<i>P</i>	<i>I</i>	<i>R</i>	<i>I</i>	<i>P</i>	<i>I</i>	<i>I</i>
a11					<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a12						<i>P</i>	<i>I</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>
a1							<i>I</i>	<i>I</i>			<i>I</i>	<i>I</i>		
a2								<i>I</i>	<i>I</i>	<i>I</i>	<i>I</i>	<i>I</i>	<i>P</i>	<i>R</i>
a4									<i>I</i>	<i>I</i>	<i>I</i>	<i>I</i>	<i>P</i>	<i>R</i>
a5									<i>I</i>	<i>I</i>	<i>P</i>	<i>P</i>	<i>P</i>	<i>P</i>
a6											<i>I</i>	<i>R</i>	<i>R</i>	<i>R</i>
a10										<i>P</i>		<i>I</i>	<i>I</i>	<i>P</i>
a13													<i>I</i>	
a9									<i>P</i>					<i>I</i>
a14									<i>P</i>			<i>P</i>		

Confiance du résultat : 71.42%

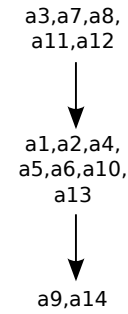
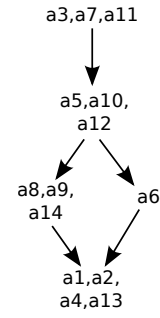


FIGURE 2.2 – Résultats de la classification non supervisée.

Objectif	Fitness (% de l'optimum)	
NR / étape 1	96.00	
NR	99.94	(0.08)
TP	99.82	(0.33)
TC	99.80	(0.29)
OPS	76.00	(17.38)
OCS	98.51	(2.10)

TABLE 2.1 – Moyennes des résultats (écarts types entre parenthèses).

incohérences avec le résultat de la classification (par exemple *a5* et *a8* sont indifférentes alors qu'elles n'appartiennent pas à la même classe).

Pour le second résultat, la confiance est évidemment moins élevée que pour le premier calcul, car l'algorithme tente d'écarter l'incomparabilité entre les classes.

En vue de valider notre algorithme de manière empirique, nous générons des jeux de données aléatoires. Comme l'objectif est de montrer que la méthode approxime bien la solution optimale, ces données ne contiennent que 10 alternatives, afin de nous permettre d'utiliser une méthode exacte comme algorithme de contrôle. D'autre part, nous générons directement les relations de préférences et ne passons pas par une procédure d'agrégation particulière comme dans l'exemple illustratif.

800 problèmes sont générés en faisant varier plusieurs paramètres décrits ci-dessous :

- Nombre de classes de 2 à 10 ;
- Taille des classes (différentes combinaisons de petit, moyen et grand) ;
- Pourcentage d'indifférences dans les classes (60% ou 80%) ;
- Pourcentage d'indifférences entre les classes (5% ou 15%) ;
- Relation entre les classes ;
- Perturbations sur la structure créée (5% ou 15%).

Pour les 800 problèmes nous avons d'abord utilisé la méthode d'énumération de toutes les partitions afin de déterminer le résultat optimal pour chaque objectif. Ensuite nous avons appliqué notre algorithme 100 fois sur chaque problème et pour chaque objectif. Chaque essai de résolution a été arrêté après 1 seconde.

Le tableau 2.1 présente les résultats de ces expériences. La première colonne indique l'objectif de classification non supervisée recherchée. La seconde colonne donne le ratio entre la valeur de la fonction objectif déterminée par notre méthode et celle du résultat optimal.

Les résultats indiquent bien la difficulté croissante des différents objectifs de classification. Pour des objectifs simples, on trouve des résultats proches de la méthode exhaustive. La première ligne indique la qualité de la solution après l'étape 1 de notre méthode, ce qui tend à prouver que cette étape fournit déjà de bons résultats.

2.2.4 Conclusions

La validation empirique de notre méthode de classification non supervisée montre l'intérêt de cette approche. Une nouvelle fois nous estimons que cette étape de validation est essentielle, car elle permet de donner des indications sur le comportement de l'algorithme proposé dans des situations réelles.

Justement en parlant de situations réelles, le lecteur attentif aura remarqué que nous n'avons pas encore abordé de domaines applicatifs pour la classification non supervisée en AMCD. Tout d'abord il y a évidemment l'analyse exploratoire, qui permet de mieux comprendre la structure du problème sous-jacent, en regroupant les alternatives indifférentes dans des classes. On pourrait également imaginer des variations de la classification ordonnée stricte pour résoudre des problèmes de tri, lorsque la structure des classes n'est pas tout à fait claire, et qu'elle devrait être découverte.

Il est possible d'étendre ces travaux à des cas où le nombre d'alternatives est plus élevé. En effet, en utilisant la meta-heuristique que nous avons proposée dans [Bisdorff et al. \(2011\)](#), il est possible de déterminer les cœurs de manière approchée. Cependant la deuxième étape de notre méthode nécessitera des améliorations, étant donné que le nombre d'opérations de déplacements est potentiellement quadratique.

*

La section suivante présente une contribution à l'élicitation des préférences pour la problématique du tri dans un contexte où plusieurs décideurs interviennent. Le paradigme utilisé est celui des relations de surclassement, et comme dans les deux précédents travaux, une attention particulière est portée à la validation empirique de la procédure proposée. Par ailleurs, pour des raisons opérationnelles, nous soulignons comment se positionne cette procédure dans l'ensemble des outils de soutien à l'élicitation de préférences d'un groupe de décideurs.

2.3 Aide au tri en présence de critères et de décideurs multiples

Ce travail se situe dans le contexte de la problématique du tri multicritère, dans laquelle chaque alternative de X doit être affectée à une catégorie prédéfinie. Les k classes disponibles sont par ailleurs ordonnées préférentiellement $C_1 \ll C_2 \ll \dots \ll C_h \ll \dots \ll C_k$ (C_1 étant la plus mauvaise catégorie, et C_k la meilleure).

Par rapport au travail de la section précédente, les catégories et leur nombre sont donnés à l'avance et l'affectation des alternatives ne se fait pas en les comparant entre-elles, mais par des comparaisons à des normes données a priori.

De nombreuses approches ont été proposées afin de résoudre des problèmes de tri multicritère ([Perny, 1998](#); [Greco et al., 2002](#)). Dans ce travail, nous nous intéressons à une variante de la méthode ELECTRE TRI ([Figueira et al., 2005](#); [Mousseau et al., 2000](#); [Roy, 1991](#)) qui respecte l'axiomatique de [Bouyssou et Marchant \(2007b,c\)](#).

Cette variante affecte les alternatives aux catégories en utilisant les performances des alternatives ainsi que de l'information préférentielle de trois types : les profils qui délimitent les catégories, des poids qui spécifient l'importance de chacun des critères, et des seuils de veto. Afin d'éliciter ces préférences, nous supposons que les décideurs sont capables de fournir des exemples d'affectation, i.e. des alternatives fictives ou réelles associées à la catégorie à laquelle les décideurs les affecteraient. Ces exemples d'affectation peuvent correspondre à des décisions passées ou être liées à des alternatives que les décideurs connaissent bien.

Bon nombre de travaux sur l'élicitation des préférences dans le contexte de l'AMCD se focalisent sur la représentation des préférences d'un décideur unique. Dans ce travail,

que nous avons effectué en collaboration avec Vincent Mousseau³ et Olivier Cailloux⁴, nous nous sommes intéressés à une procédure d'élicitation pour des groupes de décideurs qui permet à chacun d'eux de fournir de l'information préférentielle individuelle en vue de construire un modèle de tri multicritère représentant les préférences du groupe.

Cette procédure utilise des programmes linéaires afin de déterminer, à partir d'exemples d'affectations fournis par les décideurs, des profils communs, partagés par tous les décideurs. Les poids des critères pouvant par contre varier d'un décideur à un autre, cette procédure peut être vue comme une première étape vers un modèle consensuel.

La suite de la section se structure de la façon suivante. Nous présentons d'abord dans la section 2.3.1 le contexte de ce travail ainsi que les procédures d'élicitation connexes, avant de passer à la section 2.3.2 à une présentation succincte des outils que nous proposons. Les formulations des programmes linéaires qui permettent de déterminer les paramètres préférentiels ne sont pas détaillés ici et le lecteur intéressé pourra se référer à [Cailloux et al. \(2012\)](#) pour plus d'informations à ce sujet. En section 2.3.3 nous passons ensuite à la proposition d'un processus d'élicitation incluant ces nouveaux algorithmes, avant de détailler leur application à un exemple illustratif et leur validation empirique dans la section 2.3.4.

2.3.1 Electre Tri et méthodes d'élicitation des préférences connexes

Par rapport à la méthode classique d'ELECTRE TRI, la procédure utilisée dans le cadre de ces travaux ([Bouyssou et Marchant, 2007b,c](#)) utilise uniquement la règle d'affectation pessimiste, sans seuils d'indifférence ou de préférence, et avec un veto binaire, qui invalide un surclassement, peu importe la situation de concordance.

Dans la suite nous rappelons brièvement le principe de cette méthode de tri multicritère. Soit un ensemble de profils $B = \{b_0, \dots, b_k\}$ et un ensemble fini de critères $\{g_j, j \in J\}$. $g_j(a)$ dénote la performance de l'alternative a sur le critère g_j . Chaque catégorie C_h est définie par les performances de son profil inférieur b_{h-1} et de son profil supérieur b_h , avec $b_{h-1}, b_h \in B$. Nous supposons qu'une performance plus élevée représente une performance meilleure et que les performances des profils sont non-décroissantes, i.e. $\forall j \in J, 1 \leq h \leq k : g_j(b_{h-1}) \leq g_j(b_h)$.

En vue de trier les alternatives, la méthodologie d'ELECTRE TRI définit une relation de surclassement S sur l'ensemble des alternatives de manière à ce qu'une alternative a surclasse une alternative b si et seulement si a est considérée comme au moins aussi bonne que b . La règle d'affectation pessimiste affecte une alternative a à la catégorie la plus élevée C_h telle que l'alternative surclasse le profil inférieur b_{h-1} de cette catégorie.

Une alternative a surclasse un profil b_{h-1} si et seulement si la coalition de critères en faveur de l'assertion " a est au moins aussi bon que b_{h-1} " forme une majorité (pondérée) et qu'aucun critère ne s'oppose fortement à cette assertion.

La coalition de critères en faveur de aSb_{h-1} , $\forall a \in X, 1 \leq h \leq k$, forme une majorité si et seulement si

$$\sum_{j \in J} w_j C_j(a, b_{h-1}) \geq \lambda, \quad (2.4)$$

où w_j est le poids du critère g_j (avec $\sum_{j \in J} w_j = 1$), $C_j(a, b_{h-1}) \in \{0, 1\}$, et $C_j(a, b_{h-1}) = 1 \Leftrightarrow g_j(a) \geq g_j(b_{h-1})$, 0 sinon. Le résultat de la somme des poids des critères en faveur

3. Vincent Mousseau, Ecole centrale Paris, France

4. Olivier Cailloux, Ecole centrale Paris, France

Article	input	output
MS98 Mousseau et Słowiński (1998)	i	$\mathcal{P}, W$
MFN01 Mousseau <i>et al.</i> (2001)	$i, \mathcal{P}$	W
NM00 The et Mousseau (2002)	i, W	$\mathcal{P}$
DMFC02 Dias <i>et al.</i> (2002)	i	modèle robuste $(\mathcal{P}, W)$, affectations robustes
MDFGC03 Mousseau <i>et al.</i> (2003)	i^*	exemples à enlever pour rétablir la consistance
MDF06 Mousseau <i>et al.</i> (2006)	i^*	exemples à relâcher pour rétablir la consistance
DDM07 Damart <i>et al.</i> (2007)	$g, \mathcal{P}$	modèle progressif collectif (W)
CMM13 Cailloux <i>et al.</i> (2012)	g	modèle collectif $(\mathcal{P})$

TABLE 2.2 – Procédures d’inference pour ELECTRE TRI.

du surclassement est comparé à un seuil de majorité $\lambda \in [0.5, 1]$ défini avec le décideur en même temps que les poids. Si la coalition n’est pas suffisante, l’alternative ne surclasse pas le profil b_{h-1} et sera par conséquent affectée à une catégorie en dessous de C_h .

Même en cas de coalition suffisamment forte, il est possible qu’un critère s’oppose au surclassement en levant un veto. Ceci arrive lorsque $g_j(a) < v_j^{h-1}$, où le seuil de veto v_j^{h-1} représente la performance en dessous de laquelle le décideur interdit à l’alternative a de surclasser le profil b_{h-1} , et n’est par conséquent pas autorisé à être affecté à la catégorie c_h . En résumé, l’alternative a surclasse le profil b_{h-1} (et est affecté par conséquent au moins à la catégorie C_h) si et seulement si $\sum_{j \in J} w_j C_j(a, b_{h-1}) \geq \lambda$ et $\forall j \in J : g_j(a) \geq v_j^{h-1}$.

Dans une situation où il n’y a qu’un seul décideur, les poids, le seuil de majorité, les seuils de veto et les profils des catégories peuvent être soit directement donnés par le décideur, ou peuvent être inférés à partir d’exemples d’affectation. La situation se complique lorsque plusieurs décideurs sont impliqués. Nous supposons que l’ordre des catégories, les critères à utiliser et les performances des alternatives sont consensuels. Cependant, il n’y a aucune raison de penser que tous les décideurs sont d’accord a priori sur l’importance des critères, le seuil de majorité ou les limites des catégories. Dans ce travail nous montrons donc comment trouver un consensus sur ces paramètres préférentiels, à partir d’exemples d’affectation non nécessairement consensuels fournis par plusieurs décideurs.

De nombreux travaux ont été effectués sur l’éllicitation de paramètres de la procédure ELECTRE TRI pour un décideur isolé à partir d’exemples d’affectation qu’il a fourni (Mousseau et Słowiński, 1998; Mousseau *et al.*, 2001; The et Mousseau, 2002; Dias *et al.*, 2002; Dias et Clímaco, 1999, 2000), ainsi que sur la détection et la résolution d’inconsistances dans ces affectations (Mousseau *et al.*, 2003, 2006).

Damart *et al.* (2007) proposent quant à eux une méthode qui permet de gérer une multiplicité de décideurs en proposant un processus itératif de construction des poids individuels et collectifs.

Nous résumons dans le tableau 2.2 ces différentes procédures. Pour chacun des articles répertoriés, la deuxième colonne indique l’entrée attendue, alors que la dernière colonne résume la sortie produite. i désigne des exemples d’affectation d’un décideur isolé, i^* désigne les exemples d’affectation potentiellement entachés d’inconsistances d’un décideur isolé, g désigne les exemples d’affectation d’un groupe de décideurs, $\mathcal{P}$ est un ensemble de profils, W est un vecteur de poids. La dernière ligne représente notre contribution.

2.3.2 Inférence des profils des catégories à partir d'exemples d'affectation

Dans ce travail, nous proposons trois outils permettant d'aider à la recherche du consensus en vue de la construction des profils des catégories. Ces outils utilisent de la programmation linéaire en nombres entiers et sont décrits ci-après :

- **ICL**, ou “Infer Category Limits”, détermine, si possible, les profils des catégories à partir d'exemples d'affectation des décideurs en utilisant des poids et un seuil de majorité individuels, sans l'intervention de vetos ;
- **ICLV**, ou “Infer Category Limits with Vetoes”, est une généralisation d'ICL, qui détermine, si possible, les profils des catégories à partir d'exemples d'affectation des décideurs en utilisant des poids et un seuil de majorité individuels, et en faisant intervenir des vetos, si nécessaire ;
- **CWR**, ou “Compute Weights Restrictions”, mesure la latitude que les décideurs ont pour fixer l'ordre de leurs poids, étant donné que les profils consensuels ont été déterminés.

Nous renvoyons le lecteur intéressé à [Cailloux *et al.* \(2012\)](#) pour des formulations détaillées de ces programmes.

2.3.3 Processus d'aide au tri pour un groupe de décideurs

Les outils que nous avons développés ne constituent évidemment qu'une partie de la boîte à outil dont dispose un analyste pour aider un groupe de décideur dans cette situation de tri. Les autres procédures qu'il peut utiliser incluent les résultats cités à la section 2.3.1.

Afin de rendre les outils proposés opérationnels, nous présentons dans cette section un processus d'AMCD visant à obtenir un modèle consensuel de tri pour un groupe de décideur, qui se base notamment sur nos résultats. D'autre part, ce processus permet d'illustrer où ces algorithmes interviennent, et comment ils se combinent avec d'autres outils existants.

Le processus proposé comporte deux parties. Dans une première étape, il s'agit d'obtenir des profils de catégories, déterminés à partir d'exemples d'affectation et de poids des critères individuels. Ensuite, dans une deuxième étape, en partant de ces profils consensuels, le but est de trouver des poids et un seuil de majorité consensuels afin d'aboutir à un modèle de tri du groupe de décideurs.

La première partie du processus est détaillée ci-après :

- Obtenir des exemples d'affectation individuels des décideurs ;
- Rechercher des profils communs avec des poids individuels à l'aide du programme ICL ;
- Si aucune solution ne peut être trouvée, autoriser des vetos en utilisant le programme ICLV. Une autre possibilité consiste à utiliser MDFGC03 (voir tableau 2.2) afin de suggérer aux décideurs de changer certains exemples d'affectation en vue d'éliminer l'inconsistance ; le groupe étant alors considéré comme un individu ;
- Lorsqu'une solution est trouvée, calculer les latitudes des décideurs concernant la fixation des poids à l'aide de l'algorithme CWR. Ceci permet d'exclure une solution trop restrictive ou trop inéquitable ;
- Pour terminer, demander aux décideurs si les choix des profils, ainsi que les contraintes sur leurs poids sont satisfaisants. Si non, relancer les programmes ICL ou ICLV avec des contraintes supplémentaires.

décideur 1	décideur 2	décideur 3	décideur 4
Prj 1 → Average	Prj 31 → Good	Prj 61 → Bad	Prj 91 → Bad
Prj 2 → Good	Prj 32 → Bad	Prj 62 → Good	Prj 92 → Bad
Prj 3 → Good	Prj 33 → Bad	Prj 63 → Bad	Prj 93 → Bad

TABLE 2.3 – Quelques exemples d’affectation fournis par les quatre décideurs.

Lorsqu’une solution consensuelle a été trouvée pour les profils des catégories, des poids et un seuil de majorité consensuels doivent être déterminés. Il est évidemment possible qu’à ce stade, le choix des profils implique qu’aucun consensus sur les poids ne puisse être trouvé qui respecte tous les exemples d’affectation. Dans ce cas, DDM07 (voir tableau 2.2) peut être utilisé pour créer progressivement un modèle des préférences du groupe en partant de profils consensuels et d’exemples d’affectation.

Il faut également noter que nous supposons dans ce processus que les décideurs sont d’accord de discuter et d’adapter leurs exemples d’affectation, afin d’obtenir un modèle qui satisfasse tout le groupe. Dans le cas contraire, la procédure ne peut pas s’appliquer.

2.3.4 Exemple illustratif et validation empirique

Afin d’illustrer ce travail, nous imaginons un scénario fictif. Un comité a la responsabilité de déterminer quels projets de recherche doivent être financés parmi une liste de soumissions. Ce comité d’évaluation désire mettre en place une procédure systématique afin d’affecter les projets à trois catégories : celle des bons projets qui méritent directement un financement (catégorie *Good*) ; celle des bons projets qui pourront être financés si un budget supplémentaire peut être trouvé (catégorie *Average*) ; et celle des mauvais projets et qui n’obtiendront pas de financement (catégorie *Bad*). Les quatre membres du comité se mettent d’accord pour utiliser les six critères suivants pour évaluer les projets :

- La qualité scientifique du projet, évaluée sur une échelle à cinq niveaux (sq) ;
- La qualité rédactionnelle de la soumission, évaluée sur une échelle à cinq niveaux (wq) ;
- L’adéquation du projet avec les priorités du gouvernement, évaluée sur une échelle à trois niveaux (ad) ;
- L’expérience de l’équipe de recherche, évaluée sur une échelle à cinq niveaux (te) ;
- La participation de partenaires internationaux, évaluée de manière binaire (ic) ;
- Un score mesurant le niveau de publication des chercheurs, évalué sur une échelle [0,100] (ps).

Nous supposons que chacun des décideurs fournit 30 exemples d’affectation. Les tableaux 2.3 et 2.4 résument certaines de ces affectations.

Sur base de ces exemples d’affectation, le programme ICL est utilisé pour déterminer des profils consensuels. Cependant, aucune solution ne peut être trouvée, et nous utilisons alors le programme ICLV program qui permet d’inclure des vetos dans le modèle. Nous obtenons un résultat avec un veto sur le critère sq, $v_{sq}^{b_2} = 2$. Les résultats sont résumés dans le tableau 2.5. Le tableau 2.6 montre pour chaque décideur un vecteur de poids qui est compatible avec ses exemples d’affectation, le seuil de veto et les profils communs déterminés.

Nous pouvons aussi représenter pour chaque décideur les contraintes sur les poids qui découlent du choix des profils et du veto communs. Ces restrictions sont déterminées par

	sq	wq	ad	te	ic	ps
Prj 1	1	5	1	5	0	72
Prj 2	5	4	1	5	1	6
Prj 3	5	4	3	5	1	50
Prj 31	4	5	3	5	0	98
Prj 32	5	3	1	1	1	65
Prj 33	3	1	2	1	0	53
Prj 61	2	3	2	2	0	62
Prj 62	5	5	1	4	1	67
Prj 63	5	2	2	1	1	9
Prj 91	3	1	1	1	0	89
Prj 92	1	1	2	1	0	12
Prj 93	1	3	3	2	0	40

TABLE 2.4 – Performances de quelques exemples d'affectations.

Profil	sq	wq	ad	te	ic	ps
b_1	3	4	1	3	0	22
b_2	5	4	3	4	2	27
v^{b_1}	-	-	-	-	-	-
v^{b_2}	2	-	-	-	-	-

TABLE 2.5 – Profils inférés et veto.

décideur	sq	wq	ad	te	ic	ps	λ
1	0	0.2	0.2	0.4	0	0.2	0.5
2	0	0.2	0.2	0.4	0	0.2	0.5
3	0	0.2	0.2	0.3	0.3	0	0.6
4	0.2	0.2	0.2	0.4	0	0	0.5

TABLE 2.6 – Des jeux de poids et des seuils de majorité pour chaque décideur, compatibles avec les profils élicités et les exemples d'affectation.

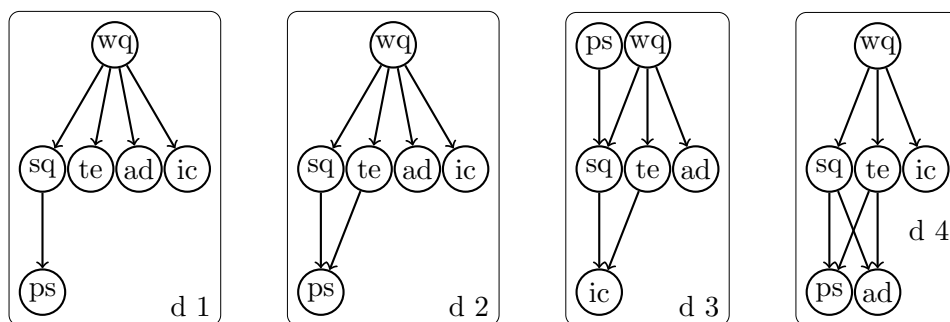


FIGURE 2.3 – Les restrictions sur les poids imposées par le choix des profils communs. Un arc d'un critère j à un critère j' signifie que le choix de ces profils implique que ce décideur devra choisir un poids pour j strictement plus élevé que pour j' .

	sq	wq	ad	te	ic	ps
poids	0.09	0.18	0.18	0.37	0.09	0.09
$\lambda = 0.5$						

TABLE 2.7 – Poids et seuil de majorité λ consensuels.

le programme CWR et sont représentées à la figure 2.3.

On peut facilement remarquer qu'il sera impossible aux décideurs de tomber d'accord sur les poids des critères (sq doit être plus important que ps pour les décideurs 1, 2 et 4, alors que ps doit être plus important que sq pour le décideur 3). Comme nous l'avons déjà évoqué précédemment, il va être nécessaire de remettre en cause certains exemples d'affectation lors d'une discussion entre les décideurs.

La méthode proposée par [Damart et al. \(2007\)](#) permet d'arriver à un consensus de poids de manière itérative. Le lecteur intéressé pourra se référer à [Cailloux et al. \(2012\)](#) pour une mise en oeuvre de cette procédure sur cet exemple. Un modèle des préférences commun peut être trouvé en modifiant 3 exemples d'affectation pour 2 décideurs. Les poids consensuels obtenus sont représentés dans le tableau 2.7.

Afin d'évaluer les performances des programmes mathématiques proposés dans le cadre de ce travail, nous les avons testés sur une série de problèmes créés artificiellement. Ainsi, nous avons généré 2000 problèmes de tailles différentes en faisant varier les paramètres suivants :

- Nombre de critères entre 3 et 10 ;
- Nombre de décideurs entre 1 et 4 ;
- Nombre de catégories entre 2 et 5 ;
- Nombre d'exemples d'affectation entre 1 et 700.

Les alternatives prennent leurs valeurs sur les critères dans l'intervalle $[0, 99]$.

Pour un problème donné, nous générons des limites de catégories que nous supposons partagées par tous les décideurs. Pour chaque décideur nous générons aussi un vecteur de poids aléatoire, ainsi que des exemples d'affectation. L'affectation de ces derniers se fait en fonction du poids généré et des profils des catégories définis à l'avance. Ensuite les programmes ICL et CWR sont appliqués au problème. Par construction il existe des profils communs aux décideurs qui satisfont tous les exemples d'affectation des décideurs.

Number of binary variables	Number of instances	Solved
[0, 399]	477	100%
[400, 799]	441	87%
[800, 1199]	362	80%
[1200, 1599]	290	78%
[1600, 1999]	268	75%
[2000, 2199]	121	69%

TABLE 2.8 – Pourcentage de problèmes résolus en moins de 90 minutes.

Nous observons ensuite le pourcentage de programmes ICL qui peuvent être résolus en moins de 90 minutes (sur un ordinateur portable de performance moyenne en 2012). Le tableau 2.8 présente ces valeurs en fonction de la taille des problèmes, i.e. le nombre de variables binaires impliquées. Ce dernier nombre est égal au produit entre le nombre de critères, le nombre de profils à déterminer et le nombre d'exemples d'affectation (Cailloux *et al.*, 2012).

Les problèmes impliquant moins de 400 variables binaires sont tous résolubles en moins de 90 minutes. Ces problèmes correspondent à des exemples avec environ 3 catégories, 5 critères et 40 exemples d'affectation. Une grande majorité des problèmes est résolue quand le nombre de variables binaires ne dépasse pas 1200. Ces problèmes correspondent à des problèmes de taille classique, comme par exemple 6 critères, 3 catégories et 100 exemples d'affectation. D'autres résultats peuvent être trouvés dans Cailloux *et al.* (2012).

Le programme CWR a aussi été testé lorsque des profils des catégories ont pu être déterminés. Les restrictions sur les poids sont calculables en quelques minutes pour les problèmes les plus difficiles, et en quelques secondes pour la majorité des problèmes.

Ces résultats montrent que le programme ICL est utilisable dans un contexte réel, où les calculs seraient effectués entre deux réunions des décideurs. Pour des problèmes de taille plus grande il sera probablement nécessaire de développer des méthodes approchées afin de déterminer les profils en un temps raisonnable.

2.3.5 Conclusions et perspectives

Ce travail ouvre un grand nombre de perspectives de recherche. En effet, des problèmes impliquant plus de 8 critères sont difficilement résolubles en un temps raisonnable. À côté de la construction de méthodes approchées, il serait aussi intéressant d'adapter l'algorithme en vue de rechercher des profils qui minimisent les restrictions sur les poids imposés aux décideurs. Par ailleurs, lorsque des profils communs ne peuvent être déterminés, il serait intéressant d'identifier des sous-ensembles de décideurs qui sont proches d'un consensus. Dans la même lignée, on pourrait imaginer que les décideurs sont capables de donner un degré de confiance à chacun de leur exemple d'affectation. La prise en compte de ces degrés, lorsqu'aucune solution ne peut être trouvée, pourrait aussi faciliter la résolution des conflits.

Pour terminer, nous aimerions à nouveau souligner l'importance de la validation empirique de la méthode que nous proposons. Elle permet justement de montrer ses limites dans des utilisations réelles, et génère un grand nombre de nouvelles questions de recherche.

Dans la section suivante, nous présentons une approche pour éliciter les poids des critères, qui tente de garantir une grande stabilité de la relation de surclassement. Une nouvelle fois, la procédure proposée sera validée de manière empirique, afin de pouvoir en déduire un processus d'élicitation des préférences adapté.

2.4 Elicitation stable des poids d'une relation de surclassement

Le contexte de ce travail est à nouveau celui des méthodes de comparaisons par paires. Classiquement, la détermination des informations préférentielles peut se faire suivant deux approches :

- soit via une *approche directe*, où les paramètres préférentiels sont déterminés par un questionnaire adéquat, avant que la relation de surclassement ne soit calculée ;
- ou, via une *approche indirecte* (désaggrégation), où de la connaissance partielle sur certaines situations de surclassement (ou la recommandation finale) est utilisée pour inférer les paramètres préférentiels du modèle.

Une des premières approches de désaggrégation a été proposée dans le contexte de théorie de la valeur multiattribut dans la méthode UTA de [Jacquet-Lagrèze et Siskos \(1982\)](#). Les auteurs y utilisent des techniques de programmation linéaire afin de déterminer des fonctions de valeur additives à partir d'un classement d'alternatives donné par le décideur.

Dans le cadre des méthodes de surclassement, comme nous l'avons déjà mentionné à la section précédente, différentes solutions ont été proposées pour inférer les paramètres de la méthode ELECTRE TRI. Par exemple, dans [Mousseau et Słowiński \(1998\)](#) et [Mousseau et al. \(2001\)](#), les auteurs proposent une approche interactive pour déterminer les paramètres préférentiels d'un modèle de tri à partir d'exemples d'affectation donnés par le décideur. D'autres travaux concernent notamment l'élicitation de seuils de veto ([Dias et Mousseau, 2006](#)).

Une des difficultés dans ces travaux est celle du choix de des valeurs des paramètres parmi toutes les valeurs compatibles avec l'information préférentielle fournie en entrée. Dans de nombreux cas, ce choix se fait de manière assez arbitraire, ou déportant ce choix sur le solveur. Une analyse de sensibilité peut alors être effectuée pour permettre de mesurer l'impact de ce choix algorithmique sur la recommandation de décision. Cependant dans la plupart des cas (comme par exemple dans [Triantaphyllou et Sanchez \(1997\)](#)), chaque critère est considéré indépendamment des autres et des variations de leurs poids sont testées séparément autour d'une solution "idéale". Une autre possibilité consiste à considérer tous les ensembles de valeurs des paramètres compatibles avec l'information fournie en entrée, afin de fournir une recommandation robuste, comme dans [Greco et al. \(2008, 2011\)](#). Cependant, dans certaines situations pratiques, ces approches très générales peuvent ne pas fournir des conclusions très riches, étant donné que l'ensemble des modèles compatibles peut être très vaste lorsque l'information préférentielle du décideur est pauvre.

Dans ce travail, que nous avons mené en collaboration avec Raymond Bisdorff⁵ et Thomas Veneziano⁶ (dans le cadre de l'encadrement de la thèse de doctorat de ce dernier), nous proposons d'éliciter un unique vecteur de poids des critères, avec comme objectif de maximiser la *stabilité* de la relation de surclassement sous-jacente. Ceci doit

5. Raymond Bisdorff, Université du Luxembourg, Luxembourg

6. Thomas Veneziano, Université du Luxembourg, Luxembourg

notamment réduire la dépendance de cette dernière à la fixation précise des valeurs des poids.

La suite de cette section s'organise comme suit : tout d'abord, nous présentons ce que nous entendons par la *stabilité* de la relation de surclassement en section 2.4.1. Ensuite, nous introduisons brièvement les procédures algorithmiques que nous avons mises en place pour l'élicitation des poids à la section 2.4.2, avant de passer à la validation empirique de notre proposition à la section 2.4.3 et une proposition d'un processus d'élicitation des poids.

2.4.1 Au sujet de la stabilité de la relation de surclassement

Pour rappel nous notons X l'ensemble des alternatives et J l'ensemble des critères. La relation de surclassement représente des comparaisons du type “est au moins aussi bon que” entre toutes paires d'alternatives de X . Elle se construit classiquement à partir d'un indice de concordance et d'un principe de discordance. L'indice de concordance mesure le degré selon lequel les comparaisons marginales sur les critères sont en accord avec la comparaison globale de deux alternatives. Les comparaisons marginales sont caractérisées par la fonction suivante, $\forall x, y \in X$ et $\forall i \in J$:

$$S_i(x, y) \begin{cases} = 1 & \text{si } x_i \text{ est clairement au moins aussi bon que } y_i, \\ = -1 & \text{si } x_i \text{ n'est clairement pas au moins aussi bon que } y_i, \\ \in]-1; 1[& \text{s'il n'est pas clair si } x \text{ est, ou n'est pas,} \\ & \text{au moins aussi bon que } y \text{ sur le critère } i, \end{cases}$$

où la performance de l'alternative x sur le critère i est notée x_i .

La transition de l'état totalement validé (+1) à l'état totalement non-validé (-1) peut se faire via une interpolation linéaire comme dans les méthodes ELECTRE (Figueira *et al.*, 2010), une fonction constante égale à la valeur médiane de l'intervalle comme dans Bisdorff (2002b) et Bisdorff *et al.* (2008), ou toute autre fonction décroissante monotone.

A chaque critère i est aussi associé un poids w_i qui représente son importance. Soit $w = (w_1, \dots, w_m)$ le vecteur des poids t.q. $0 < w_i$ ($\forall i \in J$) et soit $\mathcal{W}$ l'ensemble de ces vecteurs de poids. L'indice de concordance global est alors construit à partir d'une somme pondérée des concordances partielles.

Dans le paradigme du surclassement, dans des situations où les alternatives ont des profils très conflictuels, un principe de veto permet d'invalider selon un certain degré la concordance globale. Lorsqu'une telle situation de veto apparaît, la concordance est soit affaiblie (voir la discordance dans ELECTRE III (Roy et Bouyssou, 1993), complètement invalidée (voir le principe du veto dans ELECTRE I (Roy et Bouyssou, 1993)), ou, placée dans une situation d'indétermination (Bisdorff *et al.*, 2008).

Dans ce travail nous ne traitons que le second type de veto, qui lorsqu'il est activé, ne considère plus le vecteur de poids, et invalide complètement la concordance. Par conséquent, comme nous étudions la stabilité de la relation de surclassement par rapport à des variations des poids, sans perte de généralité, la relation de surclassement évaluée $\tilde{S}^w$ se résume ici à l'indice de concordance défini par :

$$\tilde{S}^w(x, y) = \sum_{w_i \in \mathcal{W}} w_i \cdot S_i(x, y), \quad \forall (x, y) \in X \times X.$$

Une *coupe majoritaire* de cette relation de surclassement évaluée permet de déterminer si une situation de surclassement est validée ou non. Nous disons qu'une alternative x surclasse (resp. ne surclasse pas) une alternative y si $\tilde{S}^w(x, y) > 0$, (resp. $\tilde{S}^w(x, y) < 0$), i.e. lorsqu'une majorité pondérée est en faveur (resp. n'est pas en faveur) de la situation " x est au moins aussi bon que y ". Nous noterons cette situation xS^wy (resp. $x\mathcal{S}^wy$). $\tilde{S}^w(x, y) = 0$ indique une situation de balance, où les critères en faveur du surclassement sont aussi importants que ceux en défaveur. Nous noterons cette situation $x?^wy$.

Soit $\geq_w$ le préordre⁷ sur J associé à la relation naturelle $\geq$ sur les valeurs des poids w_i du vecteur w .

Définition 4 (préordre-compatible). *Deux vecteurs de poids de critères $w, w' \in \mathcal{W}$ sont préordre-compatibles s'ils induisent le même préordre sur J .*

La notion de *stabilité* que nous avons étudié dans ce travail caractérise pour chaque $(x, y) \in X \times X$, la dépendance de la situation de surclassement au vecteur de poids $w \in \mathcal{W}$. Ce concept a été défini à l'origine par [Bisdorff \(2004\)](#). L'auteur y donne des conditions mathématiques pour évaluer cette dépendance aux poids. Soient x et y deux alternatives de X . La situation xS^wy (resp. $x\mathcal{S}^wy$) est considérée comme :

- *Indépendante* (par rapport aux poids) : si une majorité pondérée de critères est en faveur de (resp. n'est pas en faveur de) cette situation de surclassement, pour tous les vecteurs de poids de $\mathcal{W}$;
- *Stable* (par rapport aux poids) : si une majorité pondérée de critères est en faveur de (resp. n'est pas en faveur de) cette situation de surclassement entre x et y pour tout vecteur de poids qui est préordre-compatible avec w . Cette situation ne dépend que du préordre induit par w , et non de valeurs numériques particulières ;
- *Instable* (par rapport aux poids) : si une majorité pondérée de critères est en faveur de (resp. n'est pas en faveur de) cette situation de surclassement pour w , mais pas pour tous les vecteurs de poids préordre-compatibles avec w . Cette situation dépend de la précision des valeurs numériques des poids des critères.

Dans le contexte d'élicitation des préférences par une approche de désagrégation, tenter d'obtenir des poids qui maximisent le nombre de situations de surclassement stables permet une meilleure validation de la relation de surclassement par le décideur. En effet, il est possible de lui garantir que les situations stables ne changeront pas tant que le préordre induit par son vecteur de poids est respecté. Nous pensons que ce préordre peut être plus aisément validé par le décideur que des valeurs numériques précises des poids. Ainsi, si nous supposons que le décideur a validé ce préordre, il n'aura plus qu'à se focaliser sur la validation ou non des situations instables pour affiner les poids des critères.

Le lecteur intéressé pourra se référer à [Bisdorff et al. \(2013\)](#) et [Bisdorff \(2004\)](#) pour le calcul du niveau de stabilité associé à une situation de surclassement. Ci-après, nous illustrons ces concepts de stabilité par un petit exemple.

Considérons un cas avec 3 alternatives et 6 critères, dont le tableau de performance est donné dans la partie gauche du tableau 2.9, et pour lequel tous les critères sont à maximiser. Soit le vecteur de poids w qui induit le préordre $\{1, 2\} > \{3\} > \{4, 5, 6\}$ sur les critères. La relation de surclassement évaluée (dont les valeurs ont été normalisées entre -1 et 1) est représentée dans la partie droite du tableau et le niveau de stabilité y est indiqué entre parenthèses. Nous pouvons observer que les deux situations de surclassement bS^wa et cS^wb ne sont pas également stables, alors qu'elles ont la même valeur de surclassement (0.09). En

7. $>_w$ dénote la partie asymétrique de $\geq_w$, alors que $=_w$ dénote sa partie symétrique.

J	1	2	3	4	5	6	$\tilde{S}^w$ (stabilité)		
W	3	3	2	1	1	1	<i>a</i>	<i>b</i>	<i>c</i>
<i>a</i>	7	6	5	3	5	7	<i>a</i>	1.00 (ind.)	0.09 (ins.)
<i>b</i>	6	8	6	3	4	5	<i>b</i>	0.09 (sta.)	1.00 (ind.)
<i>c</i>	8	7	4	5	8	6	<i>c</i>	0.45 (sta.)	0.09 (ins.)

ind. : indépendant ; *sta.* : stable ; *ins.* : instable.

TABLE 2.9 – Tableau de performance, relation de surclassement $\tilde{S}^w$ et stabilité associée.

effet, le premier surclassement est stable et restera donc valide tant que le décideur accepte le préordre induit par les poids des critères. La faible majorité de 0.09 n'est donc pas aussi instable que sa seule valeur numérique ne le laissait penser. La deuxième situation par contre est une situation instable qui pourra potentiellement être invalidée si le décideur change légèrement les poids (même s'il respecte par exemple le préordre initial des poids).

2.4.2 Algorithme de détermination stable des poids

Dans cette section, nous ne donnerons que quelques idées générales sur la construction de la procédure qui permet de déterminer les poids garantissant une certaine stabilité de la relation de surclassement. Pour plus de détails, le lecteur pourra se référer à [Bisdorff et al. \(2013\)](#). La détermination d'un jeu de poids maximisant la stabilité de la relation de surclassement se fait via des techniques de programmation linéaire en nombres entiers. Le point de départ est l'information préférentielle qu'un décideur pourrait exprimer au sujet du problème sous-jacent. Elle peut prendre les formes suivantes :

- un ensemble $E \subseteq X \times X$ de paires d'alternatives (x, y) pour lesquelles un décideur est capable d'exprimer de la préférence stricte ou de l'indifférence ;
- un préordre partiel $\geq_N$ sur un sous-ensemble de critères $N \subseteq J$;
Exemple : le critère 1 est plus important que le critère 4 ;
- des contraintes sur des valeurs numériques de certains poids de critères ;
Exemple : le poids du critère 2 vaut 3, ou est entre 2 et 4 ;
- un préordre partiel entre certaines coalitions de critères, exprimant des préférences sur l'importance de ces ensembles ;
Exemple : la coalition de critères 1 et 3 est plus importante que 2 ;
- des ensembles de critères suffisants pour valider ou invalider une situation de surclassement ;
Exemple : si l'alternative x est au moins aussi bonne que y sur les critères 1, 2 et 3, le décideur considère que x est globalement au moins aussi bon que y .

Détaillons un peu la mise en œuvre du premier type d'information préférentielle. Soit P la relation de préférence stricte et I la relation d'indifférence. Si le décideur exprime que aPb , alors on en déduit que xS^wy et yS^wx , où w est le vecteur de poids recherché. De la même façon, si le décideur exprime que zIt , alors on en déduit que zS^wt et tS^wz .

Afin de fournir une solution aussi stable que possible au décideur, nous traitons les déclarations du type xPy comme suit :

- Nous garantissons par des contraintes linéaires que xS^wy et yS^wx ;
- Nous tentons d'obtenir la stabilité des deux surclassements (via des contraintes relaxées, voir [Bisdorff et al. \(2013\)](#) pour des détails).

De manière similaire nous traduisons des jugements d'indifférence par des contraintes qui garantissent le surclassement dans les deux sens, et nous tentons d'obtenir des situations stables par des contraintes relaxées.

Ces contraintes relaxées garantissent d'éviter un blocage en permettant au solveur de trouver une solution qui ne respecte pas totalement la stabilité de l'information préférentielle du décideur.

Les contraintes que nous déduisons pour les autres types d'information préférentielle sont plus classiques, et nous invitons à nouveau le lecteur intéressé à consulter notre article [Bisdorff et al. \(2013\)](#) pour plus de détails.

L'objectif du programme linéaire en nombres entiers qui découle de ces considérations est de déterminer un vecteur de poids w^* qui maximise le nombre de situations de surclassement stables tout en minimisant la somme des poids w_i^* , ce qui a tendance, en pratique, à minimiser le nombre de classes d'équivalence dans le préordre des poids.

2.4.3 Validation empirique et processus d'éllicitation

En vue de valider notre approche (que nous allons nommer STAB dans la suite), nous mettons en place deux expériences :

1. En partant d'une relation de surclassement complète, obtenue à l'aide d'un vecteur de poids w inconnu, nous recherchons un autre vecteur w^* qui permet de reconstruire la relation de surclassement coupée. L'objectif est de valider notre modèle et d'analyser son comportement lorsque l'information fournie en entrée est complète.
2. Dans la deuxième expérience, nous construisons itérativement un ensemble d'information préférentielle sur les alternatives qui permet de reconstruire la relation de surclassement coupée complète. Ce scénario doit mettre en évidence l'applicabilité en pratique de notre procédure, où l'information préférentielle en entrée n'est généralement que partielle.

Nous considérons 25 tailles de problèmes différentes, en faisant varier le nombre d'alternatives et de critères (7, 10, 13, 16 et 19). Pour chaque taille nous générons 500 tableaux de performances et vecteurs de poids. Ceux-ci nous permettent de construire le graphe de surclassement original, dont il s'agira de retrouver par notre procédure la version coupée.

Un algorithme de contrôle est également utilisé pour montrer l'intérêt de notre approche. Nous l'appelons A_{CON} et son objectif est de désagréger la relation de surclassement en vue de déterminer un vecteur de poids compatible, sans aucune contrainte de stabilité. Pour chacune des 3 relations de surclassement trouvées (génération aléatoire (Orig), A_{CON} , STAB), nous comptons le pourcentage de situations stables, en vue de montrer l'apport de notre approche.

Le tableau 2.10 résume, pour certaines tailles de problèmes, le pourcentage moyen de situations stables observées pour les 3 relations. Les temps de calcul indiqués sont pour un ordinateur portable moyen de 2012. Etant donné que l'algorithme de contrôle met au maximum 3 secondes pour trouver une solution, nous n'avons pas indiqué ses durées d'exécution.

Nous observons que chacune des méthodes analysées a tendance à améliorer la stabilité de la relation de surclassement originale (Orig). A_{CON} , en minimisant simplement la somme des poids, a tendance à minimiser le nombre de classes d'équivalence dans le préordre des poids, ce qui provoque une augmentation de 18% en moyenne de la stabilité. Ces résultats montrent aussi que pour des grands problèmes, les temps de calcul deviennent assez longs,

$ J $	$ X $	% moyen de surclassements stables			Temps de calcul moyen (s)
		Orig	A _{CON}	STAB	
7	7	62	86	90	0.0
7	13	68	78	85	0.3
7	19	65	72	78	1.8
13	7	48	76	86	0.0
13	13	46	62	83	4.8
13	19	46	56	76	25.1
16	7	40	71	95	0.1
16	13	41	67	91	44.5
16	19	42	60	80	412.3
19	7	38	71	86	0.0
19	13	36	55	85	11.9
19	19	39	53	77	459.7

TABLE 2.10 – Accroissements de stabilité et temps de calculs.

ce qui exclut une utilisation dans un processus d'élicitation en temps réel, en face à face avec le décideur. La deuxième expérience doit nous permettre de rendre ce processus plus opérationnel, en ne considérant en entrée qu'une partie de l'information préférentielle disponible.

Pour cette deuxième expérience, nous considérons un décideur fictif à qui on demande de fournir de l'information préférentielle sur des paires d'alternatives. Le but est d'obtenir itérativement assez d'informations pour retrouver les poids des critères qui permettent de rendre compte de toutes les comparaisons par paires (modélisées par la relation de surclassement originale). Ceci doit nous permettre d'estimer le nombre adéquat de paires d'alternatives auquel il faut confronter le décideur afin d'obtenir une solution acceptable en un temps raisonnable.

2 heuristiques sont testées pour le choix de ces paires d'alternatives :

- sélection aléatoire (SA),
- sélection de la paire ayant la valeur de surclassement la plus faible en valeur absolue (VSF), i.e. proche de la valeur médiane (et donc potentiellement très dépendante des poids).

Afin d'arrêter la construction itérative de la relation de surclassement, nous testons deux conditions : soit la reconstruction complète, soit une reconstruction à 95% conforme à la relation originale.

Le tableau 2.11 résume pour certaines tailles de problèmes, le nombre moyen de paires sélectionnées (i.e. une estimation du nombre de questions à poser au décideur avant de passer à l'exploitation de la relation de surclassement) et le nombre moyen de résolutions (si une paire sélectionnée est déjà stable, on passe à la suivante). L'heuristique SA ayant fourni de très mauvais résultats, nous ne présentons que les résultats de l'heuristique VSF.

Dans la partie gauche du tableau 2.11, nous détaillons les résultats pour une reconstruction complète de la relation de surclassement coupée. Ceci correspond à une version itérative de la première expérience, en ne considérant pas toutes les paires d'alternatives. Les temps de calcul sont significativement réduits, même pour des instances de grande taille (moins d'une seconde pour chaque itération sur un ordinateur portable standard). Ceci montre que ce processus est utilisable en temps réel pour l'élicitation des poids d'un

J	X	nb_paires [†]	reconstruction à 100%				reconstruction à 95%					
			nb_select*		nb_solve*	%stable [‡]	nb_select*		nb_solve*	%stable [‡]		
7	7	21	3.7	/	17.6%	1.2	83	2.0	/	9.5%	0.8	87
7	13	78	8.7	/	11.2%	2.2	76	4.5	/	5.7%	1.2	82
7	19	171	10.2	/	6.0%	2.4	76	5.0	/	2.9%	1.3	84
13	7	21	5.2	/	24.8%	2.1	79	2.9	/	13.8%	1.3	83
13	13	78	16.8	/	21.5%	4.5	70	7.3	/	9.4%	2.2	79
13	19	171	25.8	/	15.1%	6.2	63	8.0	/	4.7%	2.3	78
19	7	21	5.9	/	28.1%	2.5	77	3.8	/	18.1%	1.7	78
19	13	78	18.2	/	23.3%	5.6	70	7.8	/	10.0%	2.5	78
19	19	171	30.7	/	18.0%	8.4	59	10.3	/	6.0%	3.1	76

[†] nb_paires : nombre de paires d'alternatives pouvant être potentiellement considérées

* nb_select : nombre moyen / % de paires sélectionnées par l'algorithme

* nb_solve : nombre moyen de résolutions du programme linéaire

[‡] %stable : % moyen de situations stables dans la relation de surclassement obtenue

TABLE 2.11 – Résultats pour la reconstruction itérative, en utilisant l'heuristique de sélection VSF.

décideur.

Il est aussi intéressant de noter que l'heuristique VSF ne nécessite que 30% des paires possibles dans le pire des cas (beaucoup de critères et peu d'alternatives), et seulement 6% dans le meilleur des cas (peu de critères et beaucoup d'alternatives).

Ces expériences montrent clairement l'intérêt de l'algorithme proposé. Nous pouvons même en déduire une description d'un processus d'élicitation des poids comme suit :

Tout d'abord nous questionnons le décideur pour déterminer s'il peut exprimer un préordre partiel sur les poids de certains critères. Si c'est le cas, nous rajoutons ces contraintes au modèle mathématique. Ensuite nous le questionnons sur quelques paires à l'aide de l'heuristique VSF, au sujet desquelles il peut exprimer de la préférence, de l'indifférence ou de l'ignorance. Le préordre sur les poids est ensuite déterminé et présenté au décideur pour validation. S'il désire le modifier, le nouveau préordre est intégré au modèle, et nous représentons au décideur quelques paires instables (i.e. pour lesquelles au moins un arc entre les deux alternatives est instable) en vue de raffiner les poids. Pour terminer, la relation de surclassement peut être exploitée afin de répondre à la problématique de décision sélectionnée.

2.4.4 Conclusions et perspectives

La validation empirique de l'approche que nous avons proposée dans ce travail nous permet de supposer que le processus décrit devrait être applicable sur des problèmes réels sans trop de difficultés. Cependant, une vraie validation de ce travail ne peut passer que par une confrontation du protocole d'élicitation à des vrais décideurs. Dans le cadre de la thèse de Thomas Veneziano nous avons effectué des tests sur un exemple fictif avec des décideurs réels. Les résultats sont très encourageants, et le processus décrit dans cette section s'est montré très utile.

L'utilité de construire une relation de surclassement stable est assez évidente. Si la relation est bien validée par le décideur, l'exploitation qui en sera faite en vue de résoudre une problématique de décision n'en sera qu'améliorée, et la recommandation qui en découle plus robuste.

Ce travail a aussi fait l'hypothèse que les seuils de discrimination pour la construction de la concordance sont donnés à l'avance. Dans les expériences que nous avons faites

avec de vrais décideurs, nous avons remarqué la difficulté qu'ils ont à fixer ces seuils. Par conséquent nous travaillons actuellement sur une élicitation simultanée des poids et de ces seuils tout en maximisant la stabilité de la relation de surclassement.

Chapitre 3

Outils de soutien au processus d'AMCD

Sommaire

3.1 Le standard XMCD	38
3.1.1 Présentation	39
3.1.2 Illustration	40
3.1.3 Conclusion	43
3.2 Les services-web XMCD	44
3.2.1 Motivations	44
3.2.2 Détails techniques et accès	46
3.2.3 Conclusion	47
3.3 La plateforme diviz	48
3.3.1 Description et utilisation	48
3.3.2 Exemple d'utilisation	51
3.3.3 Conclusions et perspectives	55
3.4 Modélisation du processus d'AMCD	56
3.4.1 Aperçu du modèle	57
3.4.2 Une instance du processus d'AMCD	61
3.4.3 Perspectives	65

Ce chapitre présente le second axe de nos contributions, sous la forme de travaux sur des outils informatiques ou de modélisation qui permettent de soutenir le processus d'AMCD. Depuis le début de nos recherches en AMCD, nous avons toujours porté une attention particulière à ces outils qui contribuent au bon déroulement du processus d'aide à la décision.

Nous sommes partis du constat initial que la situation des logiciels en AMCD n'était pas satisfaisante, et que le domaine pourrait bénéficier d'outils plus performants et pérennes. Nous montrons donc dans la suite, comment nous avons initié plusieurs actions qui, ces dernières années, ont énormément fait évoluer le paysage des logiciels en AMCD.

A cette fin, nous introduisons dans la section 3.1 le format de données XMCD, dont l'objectif est de proposer une façon standard de représenter des données issues de l'AMCD en XML, en vue de permettre à différents programmes informatiques de devenir plus facilement interopérables (Meyer et Bigaret, 2013). Ensuite, dans la section 3.2 nous présentons l'effort de publication d'algorithmes d'AMCD sous la forme de services web, avant

de détailler, dans la section 3.3, le logiciel *diviz*, un outil de création de chaînes de traitement algorithmiques pour l'AMCD utilisant le format XMCDa ainsi que les services web (Meyer et Bigaret, 2012). Pour terminer, dans la section 3.4 nous décrivons nos travaux en cours sur la modélisation du processus d'AMCD en utilisant des techniques issues de l'ingénierie dirigée par les modèles. Il est intéressant de noter que les contributions de ce chapitre s'inscrivent dans le cadre du Decision Deck Consortium (Decision Deck Consortium, 2009a), un groupement de chercheurs européens, dont l'objectif est de proposer des outils informatiques pour venir en soutien au processus d'AMCD.

3.1 Le standard XMCDa

Les activités de recherche en AMCD se sont rapidement développées ces dernières années, et ont donné naissance à de nombreux outils informatiques qui implémentent diverses procédures d'agrégation multicritères. Malheureusement, lorsqu'il s'agit de se servir de ces logiciels en pratique, l'utilisateur est souvent confronté aux problèmes suivants :

1. des procédures d'agrégation différentes sont généralement implémentées dans des logiciels différents, avec des interfaces utilisateur hétérogènes ;
2. tester plusieurs procédures d'agrégation multicritère sur un même problème est souvent difficile, car les logiciels ne supportent pas les mêmes formats de données ;
3. un grand nombre d'algorithmes en AMCD qui sont publiés dans des articles scientifiques ne sont pas accessibles facilement, et ne sont par conséquent souvent utilisés que par leurs auteurs ;
4. beaucoup de logiciels d'AMCD ne sont pas libres (ni au sens financier, ni d'un point de vue *open-source*), ce qui peut être considéré comme un frein pour leur dissémination plus large.

Afin de tenter d'améliorer cette situation, un groupe de chercheurs européens a créé le Decision Deck Consortium (voir Decision Deck Consortium, 2009b), dont l'objectif est de développer de manière collaborative des logiciels *open-source* implémentant des techniques d'AMCD.

L'hétérogénéité des formats de données utilisés par les différents programmes empêche également de créer des combinaisons d'algorithmes d'AMCD, et rend difficile la mise en place de chaînes de traitement impliquant plusieurs outils informatiques pour soutenir le processus d'AMCD. Dans cette situation, la résolution d'un problème de décision complexe se résume en général à l'utilisation d'un seul outil d'AMCD, complété potentiellement par d'autres logiciels de pré- ou de post-analyse. Cette constatation représente notamment une grande frustration pour de nombreux analystes, chercheurs et enseignants en AMCD qui aimeraient tester plusieurs procédures d'agrégation sur une même instance d'un problème, sans devoir réencoder les données dans différents formats.

Dans cette section nous présentons de manière succincte le standard de données XMCDa, issu des travaux du Decision Deck Consortium. Ce format unique doit permettre aux logiciels qui l'adoptent de devenir interopérables tout en garantissant une facilité d'utilisation. XMCDa est principalement le fruit d'une collaboration avec 2 chercheurs du consortium, Raymond Bisdorff¹ et Thomas Veneziano². La définition de ce standard a pris plusieurs

1. Raymond Bisdorff, Université du Luxembourg, Luxembourg

2. Thomas Veneziano, Université du Luxembourg, Luxembourg

mois, et les versions intermédiaires ont été régulièrement confrontées aux autres membres du consortium en vue de leur validation. Ce travail de longue haleine commence à porter ses fruits, étant donné que le nombre de programmes compatibles avec XMCD est en augmentation permanente.

La suite de la section se structure de la façon suivante. Tout d’abord, en 3.1.1 nous présentons la philosophie générale qui nous a guidé dans la définition de XMCD, ainsi que les grandes lignes du standard. Pour une présentation plus exhaustive, le lecteur intéressé pourra se référer à Meyer et Bigaret (2013). Ensuite, la section 3.1.2 présente une mise en œuvre de XMCD via l’encodage d’un problème classique de la littérature d’AMCD.

3.1.1 Présentation

XMCD est un langage de balisage qui utilise la syntaxe de XML³. Il est défini par un Schema XML⁴ qui décrit sa structure et les contraintes qui régissent son contenu. Un document XML qui respecte le Schema XMCD est dit “XMCD valide”. La version officielle de XMCD au moment de la rédaction de ce texte est 2.2.0.

Avant toute chose, il est important de noter que le but de XMCD n’est pas de modéliser le processus d’AMCD qui est composé de plusieurs étapes et où plusieurs intervenants jouent des rôles bien définis (décideur, analyste, expert métier, ...). XMCD a été développé pour permettre à des algorithmes, procédures d’agrégation et outils d’analyse du domaine de l’AMCD de devenir interopérables et de partager un même standard de données. Par conséquent, XMCD se limite à modéliser des données et concepts nécessaires à ces procédures.

L’origine de XMCD remonte à l’automne 2007 où le Decision Deck Consortium s’est réuni à Paris pour réfléchir à un standard de données pouvant être utilisé par un grand nombre de logiciels d’AMCD. Les travaux qui en ont découlé ont donné lieu à plusieurs versions successives, chacune améliorant la précédente en corrigeant des erreurs ou en rajoutant de nouveaux concepts, jusqu’à la version actuelle.

Un fichier XMCD contient un certain nombre de balises qui décrivent les concepts et les données liées au problème d’aide à la décision sous-jacent. Lors de l’élaboration du standard, une des grandes difficultés était de réunir autour d’une même représentation, des concepts issus d’écoles de pensées différentes. Un exemple simple est le celui du poids d’un critère, qui a une sémantique différente dans le surclassement et les méthodes recherchant un critère de synthèse unique. La version actuelle de XMCD fait ainsi souvent abstraction de la sémantique précise liée à certains concepts, afin de tenter d’unifier un grand nombre d’approches méthodologiques. Ainsi, le poids d’un critère est représenté dans XMCD sous la forme d’une valeur numérique associée à un critère et le soin est donc laissé à l’utilisateur ou au logiciel d’attacher une certaine sémantique à cette donnée. Nous sommes évidemment conscients des abus que cette approche peut générer.

La balise racine de XMCD est appelée XMCD et contient de nombreuses sous-balises, chacune décrivant en détail des données ou des concepts liés au problème de décision. En résumé, ces balises peuvent être affectées à cinq catégories générales :

- la description du projet d’aide à la décision courant ou du fichier XMCD ;
- les définitions des concepts élémentaires d’AMCD comme les attributs, les critères, les catégories et les alternatives ;

3. <http://www.w3.org/XML/>

4. <http://www.w3.org/XML/Schema>

ID voiture	marque	coût ($g1$, €)	accel. ($g2$, s)	reprise ($g3$, s)	freins ($g4$)	tenue de route ($g5$)
a01	Tipo	18342	30.7	37.2	2.33	3
a02	Alfa	15335	30.2	41.6	2	2.5
a03	Sunny	16973	29	34.9	2.66	2.5
a04	Mazda	15460	30.4	35.8	1.66	1.5
a05	Colt	15131	29.7	35.6	1.66	1.75
a06	Corolla	13841	30.8	36.5	1.33	2
a07	Civic	18971	28	35.6	2.33	2
a08	Astra	18319	28.9	35.3	1.66	2
a09	Escort	19800	29.4	34.7	2	1.75
a10	R19	16966	30	37.7	2.33	3.25
a11	P309-16	17537	28.3	34.8	2.33	2.75
a12	P309	15980	29.6	35.3	2.33	2.75
a13	Galant	17219	30.2	36.9	1.66	1.25
a14	R21t	21334	28.9	36.7	2	2.25

TABLE 3.1 – Tableau de performance pour le problème de choix de Thierry.

- la description des préférences liées aux critères, alternatives, attributs ou catégories (soit fournies en entrée par le décideur, ou produites en sortie par une procédure) ;
- les données du tableau de performance ;
- des messages produits par les algorithmes (messages d’erreur ou d’information) et des paramètres d’entrée pour ces procédures.

Un fichier XMCDa valide ne contient pas forcément des éléments de toutes ces catégories et peut n’être constitué que d’une seule balise.

Le lecteur intéressé trouvera une documentation détaillée du standard sur le site web du Decision Deck Consortium⁵. Un guide d’utilisation avec une description des principales balises peut être trouvé dans Meyer et Bigaret (2013). Nous préférons ici tenter de donner une intuition sur XMCDa à travers un exemple illustratif à la section suivante.

3.1.2 Illustration

Le problème de décision, dont nous présentons le codage XMCDa ici, concerne le choix d’une voiture par un étudiant belge. Il a été décrit dans Bouyssou *et al.* (2000, chapitre 6), où le lecteur intéressé trouvera des informations complémentaires. Il s’agit du même exemple que celui traité dans la section 2.2.3 sur la classification non supervisée.

En 1993, Thierry, un étudiant âgé de 21 ans, est passionné par les voitures de sport et aimerait acheter une voiture d’occasion puissante. Il sélectionne cinq critères liés au coût (critère $g1$), à la performance du moteur (critères $g2$ and $g3$) et à la sécurité (critères $g4$ et $g5$). La liste des alternatives et de leurs évaluations sur ces cinq critères est présentée dans le tableau 3.1. Le critère de “coût” (€) et les critères de performance “accélération” (secondes) et “reprise” (secondes) doivent être minimisés, tandis que les critères de sécurité “freinage” et “tenue de route” doivent être maximisés. Les valeurs prises par ces deux derniers critères sont des évaluations moyennes obtenues à partir d’évaluations qualitatives multiples qui ont été recodées en entiers entre 0 et 4.

5. <http://www.decision-deck.org/xmcda>

Comme dans Bouyssou *et al.* (2006, chapitre 7), nous supposons ici que Thierry a une certaine connaissance de plusieurs voitures, et qu'il est capable d'en classer quelques unes d'après ses préférences :

P309-16 $\succ$ Sunny $\succ$ Galant $\succ$ Escort $\succ$ R21t.

En XMCD, un extrait de la partie définissant les alternatives est donné ci-après :

```
<alternatives name="Les voitures potentielles">
  <alternative id="a12" name="P309">
    <description>
      <comment>Peugeot 309</comment>
    </description>
  </alternative>
  [...]
  <alternative id="a14" name="R21t">
    <description>
      <comment>Renault 21</comment>
    </description>
  </alternative>
</alternatives>
```

Chaque alternative est décrite au minimum par son identifiant (id). La balise *description* contient potentiellement d'autres balises que *comment*, mais nous nous limiterons ici à celle-là.

Les critères sont définis par l'extrait de code suivant :

```
<criteria>
  <criterion name="cout" id="g1">
    <description>
      <comment>Cout en euros</comment>
    </description>
    <scale>
      <quantitative>
        <preferenceDirection>min</preferenceDirection>
      </quantitative>
    </scale>
  </criterion>
  [...]
  <criterion name="tenue de route" id="g5">
    <description>
      <comment>0: mauvaise, 4: bonne</comment>
    </description>
    <scale>
      <quantitative>
        <preferenceDirection>max</preferenceDirection>
        <minimum><real>0</real></minimum>
        <maximum><real>4</real></maximum>
      </quantitative>
    </scale>
  </criterion>
</criteria>
```

Comme pour les alternatives, chaque critère est au minimum décrit par son identifiant (id). Les directions de préférences peuvent être spécifiées avec la balise *preferenceDirection*, ainsi que les bornes de l'échelle d'évaluation, si nécessaire.

Les performances des voitures sur les différents critères sont stockés dans l'extrait de code suivant :

```

<performanceTable>
  <alternativePerformances>
    <alternativeID>a11</alternativeID>
    <performance>
      <criterionID>g1</criterionID>
      <value><real>17537</real></value>
    </performance>
    [...]
    <performance>
      <criterionID>g5</criterionID>
      <value><real>2.75</real></value>
    </performance>
  </alternativePerformances>
  [...]
  <alternativePerformances>
    <alternativeID>a14</alternativeID>
    <performance>
      <criterionID>g1</criterionID>
      <value><real>21334</real></value>
    </performance>
    [...]
    <performance>
      <criterionID>g5</criterionID>
      <value><real>2.25</real></value>
    </performance>
  </alternativePerformances>
</performanceTable>

```

Chaque balise *alternativePerformances* correspond à une ligne du tableau de performances.

Le classement fourni par Thierry sur quelques voitures est résumé dans l'extrait suivant :

```

<alternativesValues name="Rangs des alternatives">
  <description>
    <comment>Classement a priori de Thierry</comment>
  </description>
  <alternativeValue>
    <alternativeID>a11</alternativeID>
    <value>
      <integer>1</integer>
    </value>
  </alternativeValue>
  [...]
  <alternativeValue>
    <alternativeID>a14</alternativeID>
    <value>
      <integer>5</integer>
    </value>
  </alternativeValue>
</alternativesValues>

```

Ce morceau de XMCD A représente les rangs des alternatives dans le classement de Thierry.

Pour terminer, nous présentons aussi un extrait d'une sortie d'une procédure d'élicitation de préférences. Nous avons choisi pour cela la technique de désagrégation UTA de Jacquet-Lagrèze et Siskos (1982) qui, à partir d'un classement d'alternatives, fournit des fonctions de valeur qui modélisent les préférences du décideur.

```

<criteria name="fonctions de valeur">
  <criterion id="g1">
    <criterionFunction>

```

```

<points>
  <point>
    <abscissa>
      <real>21334.0</real>
    </abscissa>
    <ordinate>
      <real>0.0</real>
    </ordinate>
  </point>
  <point>
    <abscissa>
      <real>17587.5</real>
    </abscissa>
    <ordinate>
      <real>0.5</real>
    </ordinate>
  </point>
  <point>
    <abscissa>
      <real>13841.0</real>
    </abscissa>
    <ordinate>
      <real>0.5</real>
    </ordinate>
  </point>
</points>
</criterionFunction>
</criterion>
[...]
</criteria>

```

Cette portion de données représente, pour le critère de coût, une fonction linéaire par morceaux à deux segments, définie par 3 points.

3.1.3 Conclusion

Les avantages de proposer un standard de données pour l'AMCD sont donc très nombreux. En particulier, les données de l'exemple précédent peuvent être utilisées par n'importe quel logiciel d'AMCD qui est capable d'interpréter cette norme. Cela représente donc une avancée significative par rapport à la situation initiale des outils informatiques en AMCD.

Par ailleurs, grâce à des outils de transformation comme XSLT⁶, les fichiers XMCD A peuvent être aisément transformés en d'autres formats, comme par exemple du HTML, qui peut être visualisé dans n'importe quel navigateur web. La figure 3.1 montre une possible représentation du tableau de performance de l'exemple précédent dans un navigateur (obtenue à partir d'une transformation du code XMCD A).

Il est également possible de générer directement des interfaces utilisateurs à partir du Schema XML de XMCD A. Ceci doit faciliter la création d'interfaces graphiques de nouveaux logiciels d'AMCD.

Pour terminer, XMCD A permet la construction de chaînes de traitement de procédures d'AMCD. En effet, si l'analyste dispose de plusieurs logiciels d'AMCD qu'il souhaite mettre les uns à la suite des autres, ce standard facilitera ce travail, étant donné qu'aucune

6. <http://www.w3.org/Style/XSL>

	g1	g2	g3	g4	g5
a11	17537	28.3	34.8	2.33	2.75
a03	16973	29	34.9	2.66	2.5
a07	18971	28	35.6	2.33	2
a12	15980	29.6	35.3	2.33	2.75
a10	16966	30	37.7	2.33	3.25
a08	18319	28.9	35.3	1.66	2

FIGURE 3.1 – Le tableau de performance de Thierry au format HTML.

transformation de données intermédiaire n'est nécessaire. Nous reviendrons sur ce point dans les deux sections suivantes.

*

La version actuelle de XMCDa est utilisée par plusieurs logiciels ou bibliothèques d'AMCD. Par exemple on peut citer des bibliothèques de calculs comme ws-RXMCDa de [Bigaret et Meyer \(2009-2010\)](#), Digraphs de [Bisdorff \(2007\)](#) et J-MCDa de [Cailloux \(2010\)](#). Le Decision Deck Consortium propose également des outils utilisant ce standard. Il s'agit des services-web XMCDa qui facilitent l'accès à des ressources algorithmiques d'AMCD. La section suivante présente cette initiative qui démontre clairement l'intérêt pratique de du standard XMCDa.

3.2 Les services-web XMCDa

D'un point de vue général, un service-web est une application à laquelle il est possible d'accéder via Internet, et qui est exécutée sur un système distant. Les grands avantages de ces programmes en ligne sont leur disponibilité et leur facilité d'utilisation.

L'idée initiale d'utiliser des services-web pour proposer des procédures d'AMCD revient à Raymond Bisdorff. Par la suite, nous avons contribué à les mettre en place à plus grande échelle avec Sébastien Bigaret⁷, avec qui nous avons également participé à l'amélioration de leur architecture générale. Dans la suite, en section 3.2.1 nous présentons d'abord les motivations qui justifient le choix des services-web, avant de détailler leur fonctionnement et leur utilisation à la section 3.2.2.

3.2.1 Motivations

L'option des services-web s'est imposée pour un certain nombre de raisons qui sont fondées sur les exigences exprimées par le Decision Deck Consortium.

Tout d'abord, il s'agit de faciliter l'*interopérabilité* entre les différents programmes d'AMCD produits. Celle-ci entraîne la possibilité de créer des enchaînements de procédures en vue de faciliter le processus d'AMCD. La résultante de cette exigence est la naissance de XMCDa, et de manière logique, les services-web proposés sont basés sur ce standard.

7. Sébastien Bigaret, Télécom Bretagne, France

Ensuite, le consortium voudrait *capitaliser* sur le long terme les efforts d'implémentation. De nos jours, beaucoup d'implémentations d'algorithmes ont un cycle de vie très court, car, par exemple, ils sont intégrés dans des logiciels qui ne sont plus maintenus par leurs développeurs. Ainsi leurs implémentations ont tendances à disparaître et le même algorithme se voit reprogrammé de manière récurrente. La technologie des services-web répond bien à ces soucis de capitalisation et de réutilisation.

Il s'agit également de *maximiser les contributeurs* potentiels et de *minimiser leurs efforts de programmation*. En AMCD, comme dans beaucoup d'autres domaines de recherches, les personnes qui implémentent des algorithmes, ne sont pas forcément des informaticiens. Pour ces chercheurs, l'informatique n'est qu'un outil de support pour répondre à une question de recherche. Il en découle que si on leur demande de participer à un effort de programmation, leur imposer un langage de programmation est souvent perçu comme une contrainte difficile. Afin de répondre à ce problème, l'architecture de services-web mise en place permet d'y inclure des algorithmes écrits dans une grande variété de langages de programmation.

Le consortium soutient également l'idée qu'il est intéressant de produire de *petits composants algorithmiques* à la place de grandes boîtes noires. En effet, les logiciels ou procédures d'AMCD sont souvent des séquences d'algorithmes plus élémentaires. Dans de nombreux programmes ceci est peu visible, alors que les résultats intermédiaires de ces différentes étapes peuvent avoir un intérêt pour mieux comprendre le problème de décision. Les outils proposés via les services-web doivent donc correspondre à ces briques plus élémentaires, qui le cas échéant peuvent être recombinaées pour créer la procédure d'AMCD originale. Cela génère deux effets positifs de manière immédiate : l'effet boîte noire de certaines procédures est supprimé et la création de variantes de procédures est simplifiée.

Pour terminer, il est important pour le consortium de *simplifier l'accès* aux outils. Contrairement aux statistiques ou à la fouille de données, le domaine de l'AMCD souffre clairement d'un manque de disponibilité des procédures de résolution. La publication de ces ressources de calcul sous la forme de services-web accessibles au plus grand nombre doit répondre à ce problème d'accès.

Au moment de l'écriture de ce manuscrit, le Decision Deck Consortium met à disposition des utilisateurs une centaine de services-web XMCD. Bon nombre d'entre-eux ont été développés par des chercheurs ou des étudiants et se basent sur quelques librairies d'AMCD existantes, comme le package Kappalab développé par Grabisch *et al.* (2008), la librairie PyXMCD de Veneziano (2010), la librairie UTAR de Leistedt (2010) ou la librairie J-MCD de Cailloux (2010). Une liste exhaustive des services-web XMCD existants peut être trouvée sur <http://www.decision-deck.org/ws>, accompagnée d'une documentation des algorithmes implémentés, de leurs entrées et sorties au format XMCD, et des coordonnées des contributeurs.

D'un point de vue plus opérationnel, ces services permettent de reconstruire des méthodes classiques d'AMCD de la littérature. Par exemple, en combinant un certain nombre de ces services entre eux (par des appels successifs), la plupart des procédures ELECTRE (Roy, 1968) peuvent être obtenues. De même, un bon nombre des outils de la série PROMETHEE (Brans et Vincke, 1985) sont disponibles, ainsi que des techniques de désagrégation du type UTA (Jacquet-Lagrèze et Siskos, 1982). Des procédures plus récentes sont évidemment aussi disponibles sous la forme de services-web, et de nombreux contributeurs nous fournissent régulièrement de nouvelles briques algorithmiques.

3.2.2 Détails techniques et accès

Dans cette section nous présentons le principe de fonctionnement des services-web XMCD. Tout d'abord, la figure 3.2 résume l'architecture supportant ces services.

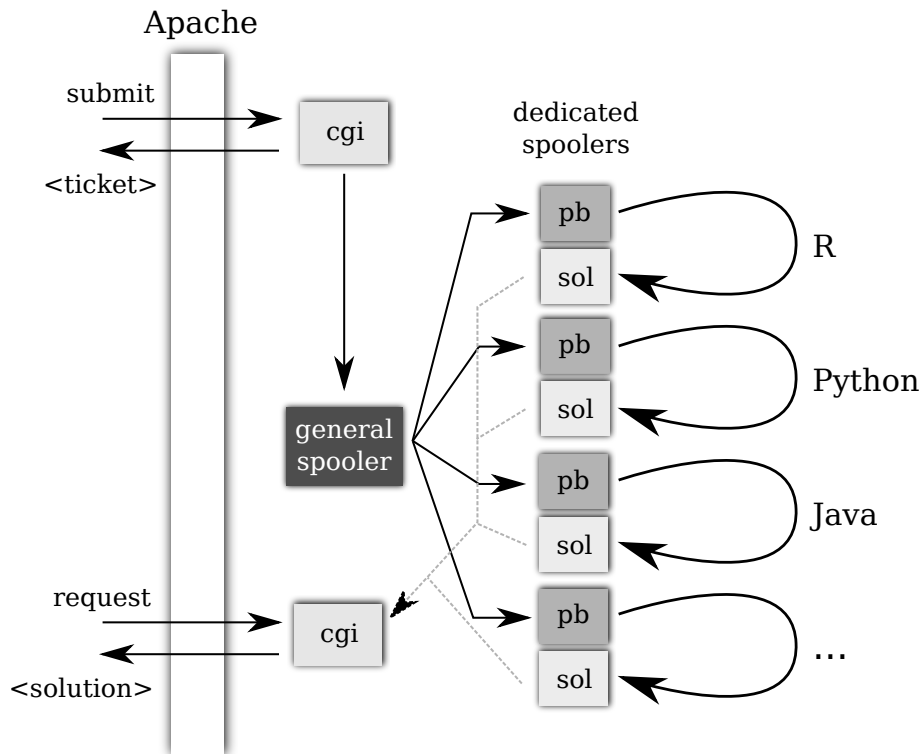


FIGURE 3.2 – Architecture des services-web XMCD.

Les calculs effectués par certains services ne nécessitent que quelques instants, alors que pour d'autres, cette exécution peut prendre plusieurs heures. Afin de tenir compte de cette spécificité, nous avons opté pour une implémentation asynchrone. Le fonctionnement de l'architecture de support des services-web est présenté ci-après :

- Un utilisateur soumet un problème à un service-web XMCD ;
- La requête est stockée dans la file d'attente générale qui contient les tâches qui doivent être traitées ;
- Le serveur retourne un *numéro de ticket* à l'utilisateur. Il identifie de manière unique la tâche qui a été soumise, et il est nécessaire pour consulter le résultat par la suite ;
- Un gestionnaire de tâches est dédié à cette file d'attente. Lorsqu'une nouvelle tâche est soumise, il la transmet à la file d'attente du service-web que l'utilisateur veut appeler ;
- Ce gestionnaire de tâches (un par service-web) s'occupe d'exécuter la tâche, en mettant d'abord en place l'environnement nécessaire à l'exécution du programme sous-jacent au service, et en le démarrant ensuite. Ce mécanisme permet l'exécution de programmes écrits en toutes sortes de langages, en leur affectant un environnement spécifique (chargement de bibliothèques, etc) ;
- Lorsque le programme se termine, son résultat est stocké ;
- Quand l'utilisateur sollicite le résultat, il présente le numéro de ticket qui lui a été fourni lors de la soumission, et récupère ensuite le résultat.

Tous les services-web sont accessibles via une interface unique, composée de trois méthodes :

- **hello** : une méthode sans paramètres, qui retourne simplement une chaîne de caractères contenant le nom du service, son auteur et sa version. Cette méthode peut être utilisée pour vérifier qu'un service est disponible.
- **submitProblem** : soumet de nouvelles tâches. Ses paramètres dépendent du service appelé, certains pouvant être optionnels. Chacun des paramètres est documenté sur la page web des services-web XMCD de Decision Deck. La méthode retourne une chaîne de caractères contenant le numéro du ticket affecté à cette tâche.
- **requestSolution** : étant donné un numéro de ticket, retourne les résultats qui ont été produit. Leur nombre dépend à nouveau du service appelé.

Les données échangées (problème soumis et résultat récupéré) sont évidemment au format XMCD.

Les services-web sont actuellement accessibles par des requêtes SOAP⁸. Il existe des implémentations SOAP pour quasiment tous les langages de programmation. Les url (adresses web) pour accéder aux services sont du type :

`http://webservices.decision-deck.org/soap/service.py,`

où **service** doit être remplacé par le nom du service. La liste des noms des services-web se trouve à l'adresse <http://www.decision-deck.org/ws>.

3.2.3 Conclusion

Dans cette section nous avons présenté les services-web XMCD. Nous avons contribué à la mise en oeuvre de leur architecture de support et nous développé un grand nombre d'entre-eux. Leur gestion quotidienne est effectuée par Sébastien Bigaret, qui s'occupe également d'implanter toute nouvelle proposition de service-web.

L'avantage de rassembler des algorithmes et procédures d'AMCD sous cette forme est indéniable et il est désormais très simple d'accéder à ces ressources de calcul et de les intégrer, par exemple, dans des logiciels.

Un tel exemple de réutilisation peut être trouvé dans le système d'information géographique open-source Quantum GIS ([Quantum GIS Development Team, 2009](#)). Depuis peu, il permet d'utiliser la méthode ELECTRE TRI et une procédure d'élicitation de profils et de poids ([Sobrie, 2011](#)). L'utilisation de ces méthodes se fait via des appels aux deux services-web XMCD correspondants, *ElectreTriExploitation* et *ElectreTriBMInference*. Pour le développeur, l'avantage a été de ne pas devoir s'occuper de l'implémentation de ces deux procédures, ce qui lui a permis de se focaliser sur la préparation des données dans le système d'information géographique.

*

La difficulté qui reste à résoudre est celle de la composition facile de ces services, en vue de créer des procédures d'AMCD capables de répondre aux besoins d'un analyste confronté à un problème d'aide à la décision. La section suivante décrit la plateforme diviz qui sert à créer de manière très intuitive des chaînes de traitement de services-web XMCD.

8. <http://www.w3.org/TR/soap12-part1/>

3.3 La plateforme diviz

La plateforme diviz est une initiative du Decision Deck Consortium, qui facilite l'utilisation des services-web XMCD, en simplifiant leur appel, et en permettant de les combiner facilement pour créer ce qu'on appelle des *workflows* d'algorithmes. Ces workflows peuvent par exemple représenter des méthodes d'AMCD, leurs variantes, ou des expériences menées sur des procédures d'agrégation.

L'outil diviz est né il y a 4 ans lors de notre rencontre avec Sébastien Bigaret⁹. Depuis, il a beaucoup évolué, et commence à être utilisé pour l'enseignement à travers toute l'Europe dans des institutions d'enseignement supérieur, et dans le cadre de la recherche pour la dissémination de résultats d'AMCD. Dans cette section nous présentons d'abord la plateforme à la section 3.3.1, avant de passer à un exemple illustratif à la section 3.3.2.

3.3.1 Description et utilisation

La figure 3.3 présente l'espace de travail de diviz.

- Dans la partie supérieure gauche, un arbre affiche la liste des workflows ouverts, avec leurs résultats d'exécution ;
- Le panneau central supérieur présente le workflow sélectionné : il montre soit le *panneau de création* qui permet de modifier le workflow, ou le *panneau de résultats* qui affiche les différents résultats intermédiaires de l'exécution ;
- Le panneau central inférieur n'apparaît que lorsque l'on analyse un résultat, et présente les sorties des différents éléments du workflow ;
- Sur la droite, tous les programmes disponibles dans diviz (les services-web XMCD) sont affichés et regroupés par thèmes.

Le logiciel peut être téléchargé à l'adresse <http://www.diviz.org>, où le lecteur intéressé trouvera aussi des exemples de workflows qu'il pourra importer facilement dans diviz.

Le design d'un workflow dans diviz se fait à l'aide d'une interface graphique intuitive, où chaque procédure de calcul est représentée par une boîte qui peut être connectée à des données ou d'autres éléments de calcul à l'aide de connecteurs (voir figure 3.4 pour une vue plus détaillée du panneau de création). La création de workflows complexes d'algorithmes ne nécessite donc pas de compétences de programmation, mais uniquement la connaissance du fonctionnement de chacun des modules utilisés (dont la documentation peut être trouvée sur le site web de diviz).

Les entrées et les sorties de ces composants de calcul peuvent être multiples et correspondre à différents concepts ou éléments de données d'AMCD. En vue d'illustrer ceci, considérons l'exemple suivant.

Exemple *diviz permet d'utiliser un composant appelé `weightedSum`. Cet élément calcule la somme pondérée d'évaluations d'alternatives par rapport à un vecteur de poids associé aux critères. Par conséquent, `weightedSum` requiert quatre entrées : la description des critères, celle des alternatives, le tableau de performances contenant les évaluations numériques des alternatives sur les critères, et les poids numériques associés aux critères. La principale sortie de ce module est le vecteur des valeurs des alternatives à travers cet*

9. Sébastien Bigaret, Télécom Bretagne, France

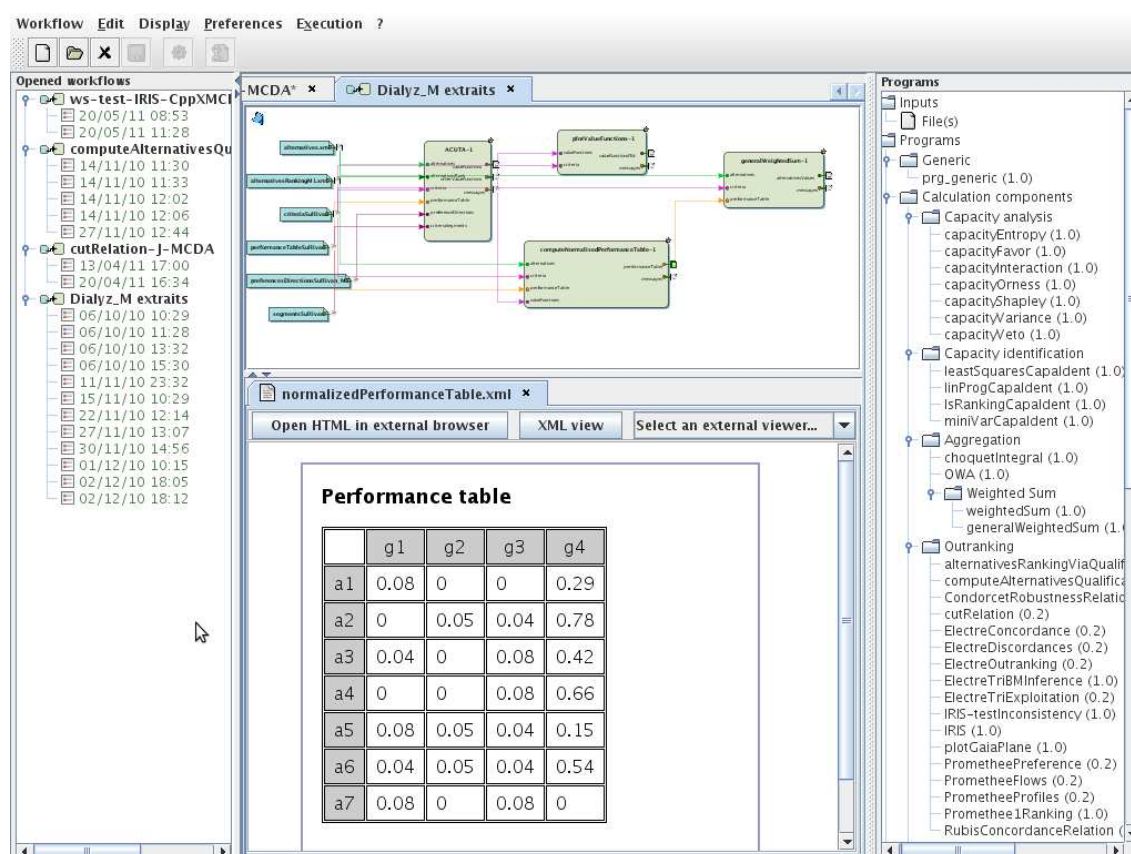


FIGURE 3.3 – Une session de travail typique dans diviz.

opérateur d'agrégation (voir la figure 3.4 pour un exemple d'utilisation du module *weightedSum*).

Pour construire un nouveau workflow, l'utilisateur choisit un ou plusieurs modules de la liste de droite et les glisse et les dépose dans l'espace de travail central. Ensuite il y ajoute les fichiers de données, et les connecte aux entrées des éléments de calcul. Pour terminer, il relie les entrées et les sorties des composants afin de définir la structure du workflow.

Lorsque le design du workflow est terminé, l'utilisateur peut l'exécuter afin d'obtenir les sorties des différents algorithmes. Ces calculs sont effectués par les services-web XMCD, et par conséquent, diviz ne contient aucune ressource de calcul locale, mais nécessite une connexion à internet pour y accéder.

Lorsque l'exécution du workflow est terminée, les sorties de chacun des modules peuvent être visualisées et analysées par l'utilisateur. L'utilisateur a donc accès à tous les résultats intermédiaires, ce qui lui permet de facilement adapter les paramètres des algorithmes.

Exemple Considérons le workflow suivant (représentant une méthode de désagrégation typique à la UTA) : un premier module détermine des fonctions de valeur linéaires par morceaux sur base d'un classement d'alternatives fourni par le décideur ; un second module transforme le tableau de performance en lui appliquant ces fonctions ; un troisième élément calcule la somme de ces performances pour chaque alternative ; un quatrième programme

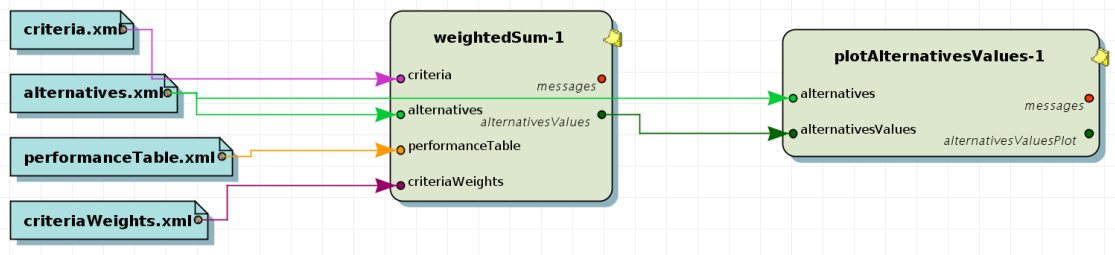


FIGURE 3.4 – Un workflow représentant les données d’entrée (à gauche), et un module de somme pondérée (au centre) qui est combiné à un module de représentation graphique du résultat (à droite).

représente graphiquement le résultat final, i.e. le classement des alternatives obtenu. Les résultats intermédiaires sont les fonctions de valeur, le tableau de performance transformé et les scores des alternatives. La disponibilité de ces informations permet à l'utilisateur de diviz d'acquérir une connaissance approfondie de la technique mise en œuvre, et lui facilite les réglages des différents paramètres (comme par exemple le nombre de segments des fonctions de valeur).

Dans diviz, l'historique des exécutions est stocké dans le logiciel et peut être consulté à tout moment. Plus précisément, si un workflow est modifié, ses versions antérieures sont toujours accessibles, ainsi que les résultats qui y sont associés. Cette fonctionnalité contribue également à la meilleure compréhension des procédures construites et facilite le calibrage des paramètres des différents éléments.

En plus de permettre de créer et d'exécuter des workflows, diviz peut aussi être utilisé pour comparer les sorties de différentes procédures d'AMCD sur un même jeu de données. Ainsi, si l'utilisateur construit à l'intérieur de son espace de travail différents workflows représentant plusieurs méthodes d'AMCD, il peut les connecter aux mêmes données d'entrée et facilement comparer leurs sorties. Cette fonctionnalité innovante est illustrée dans la section 3.3.2 via une application. Il faut cependant noter que cette possibilité doit en pratique être utilisée de manière prudente, étant donné que les paramètres préférentiels de différentes procédures d'AMCD peuvent avoir des significations très différentes. Nous pensons cependant qu'un utilisateur averti peut exploiter cette fonctionnalité pour générer de nouvelles pistes de recherche.

Pour terminer cette présentation de diviz, nous voudrions souligner une dernière propriété très intéressante de diviz, qui permet d'exporter les workflows, avec ou sans les données, sous la forme d'une archive. Cette dernière peut ensuite être partagée avec d'autres utilisateurs de diviz, qui peuvent la réimporter dans leur logiciel, et par exemple continuer le développement du workflow, ou l'exécuter sur les données originales ou leurs données personnelles. Par conséquent, diviz peut être utilisé comme un outil de diffusion très performant : tout d'abord, en combinaison avec un article de recherche, l'auteur d'une nouvelle procédure d'AMCD ou d'une expérience peut proposer le workflow diviz correspondant au téléchargement. Ensuite, dans un contexte applicatif, un analyste pourrait être tenté de partager les traitements mis en place avec les différents acteurs du processus d'AMCD. De manière plus générale, cette fonctionnalité doit participer à la dissémination plus large des résultats d'AMCD et des nouveaux algorithmes, et faciliter ainsi leur adoption par des chercheurs et des utilisateurs.

3.3.2 Exemple d'utilisation

Dans cette section nous illustrons l'utilisation de diviz sur l'exemple du choix d'une voiture que nous avons déjà introduit à la section 3.1 (voir aussi Bouyssou *et al.* (2000, chapitre 6)).

Le but de Thierry est de déterminer quelle voiture du tableau 3.1 il doit acheter. L'analyste propose par conséquent de déterminer un classement de toutes ces voitures en fonction des préférences de Thierry. Afin de mettre en avant le potentiel de diviz, nous étendons l'exemple classique comme suit :

- Dans une première étape, Thierry exprime des préférences sur quelques voitures qu'il connaît bien, sous la forme d'un classement :

$$P309-16 \succ \text{Sunny} \succ \text{Galant} \succ \text{Escort} \succ \text{R21t};$$

- Ensuite, dans une deuxième étape, après une discussion avec l'analyste, il change d'avis et n'est plus très confiant dans ce classement. Il préfère donner des informations préférentielles sur l'importance relative des différents critères et des seuils de discrimination. Thierry est conscient du fait que ces nouvelles préférences ne sont pas forcément compatibles avec le classement fourni précédemment, mais il désire comparer les deux résultats.

Afin de déterminer les deux recommandations et d'arriver à leur comparaison, l'analyste construit le workflow suivant dans diviz :

- Une première partie du workflow représente une variante de la technique UTA (Jacquet-Lagréze et Siskos, 1982) appelée ACUTA (Bous *et al.*, 2010), et qui permet de déterminer un classement des alternatives sur base du classement initial de Thierry (cadre 1 de la figure 3.5);
- Une deuxième partie du workflow représente la méthode PROMETHEE (Brans et Vincke (1985)) en vue de déterminer un classement des alternatives sur base d'information intra- et inter-critère (cadre 2 de la figure 3.5);
- Pour terminer, une troisième partie du workflow est utilisée pour comparer les sorties des deux méthodes (cadre 3 de la figure 3.5).

De plus amples détails sur la partie "ACUTA" du workflow sont visibles à la figure 3.6. Ce workflow contient 7 éléments algorithmiques. En résumé, la procédure ACUTA détermine des fonctions de valeur linéaires par morceaux, sur base d'un classement d'alternatives fourni en entrée, dans le cadre d'un modèle additif. Par conséquent, le module *ACUTA* requière en entrée la description des alternatives et des critères, le tableau de performance, un classement d'un sous-ensemble des alternatives, les directions de préférence des critères, et le nombre de segments des fonctions à déterminer. La sortie principale de ce composant est un ensemble de fonctions de valeur compatible avec le classement fourni en entrée.

Ces fonctions sont alors utilisées par *computeNormalisedPerformanceTable* pour transformer le tableau de performance original en un tableau qui représente les *valeurs de satisfaction* (entre 0 et 1) associées par Thierry à chacune des évaluations. La figure 3.7 représente un graphique de ces fonctions de valeur à l'aide du module *plotValueFunctions*.

De ces premiers calculs, l'analyste déduit que le critère de coût est le plus important pour Thierry, et qu'il a une forte préférence pour des voitures qui coûtent moins de 18000 €.

Le module *generalWeightedSum* est ensuite utilisé pour calculer la valeur agrégée de chaque alternative. Cette sortie est représentée sous la forme d'un diagramme à barres via

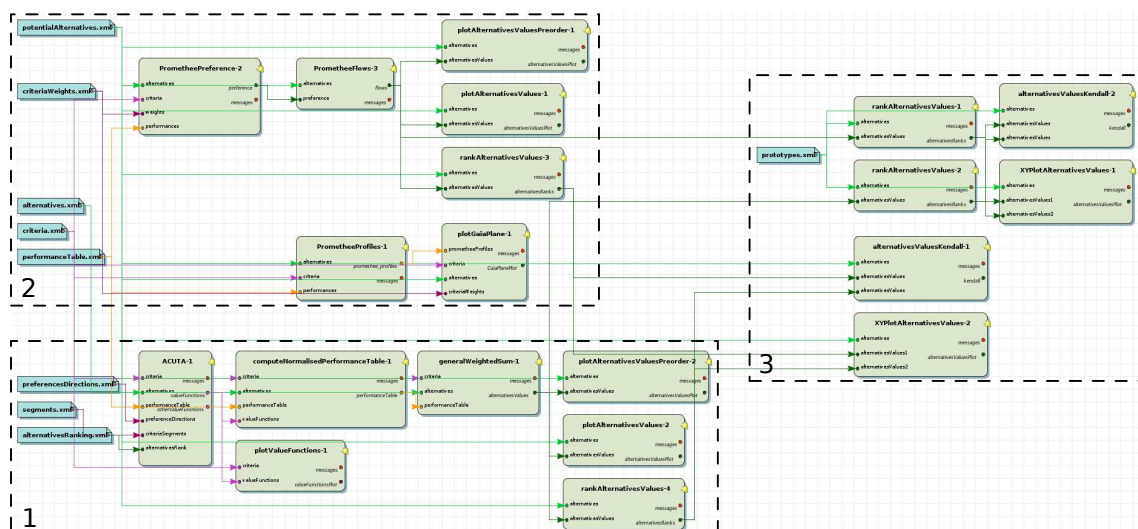


FIGURE 3.5 – Le workflow pour le problème de sélection d’une voiture.

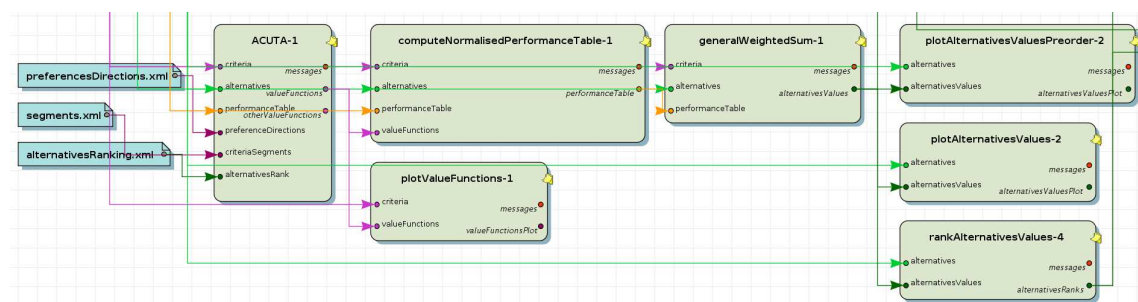


FIGURE 3.6 – La partie “ACUTA” du workflow.

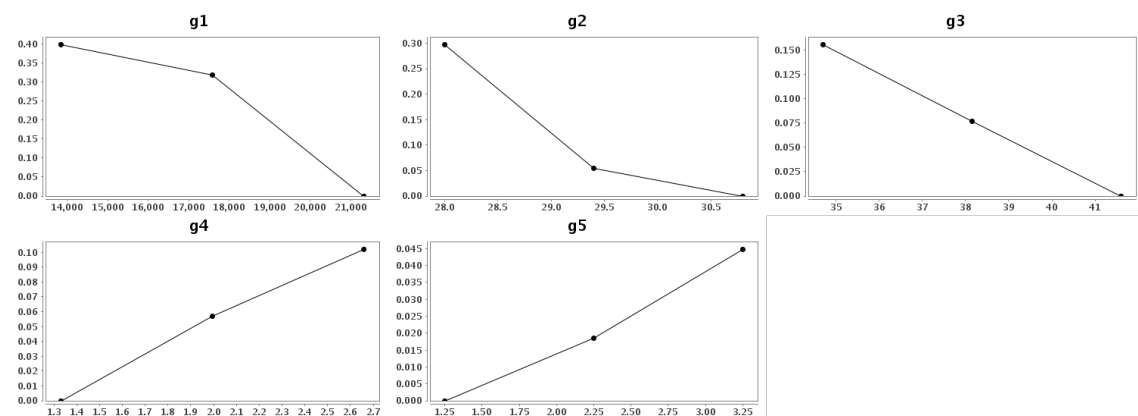


FIGURE 3.7 – Fonctions de valeur marginales.

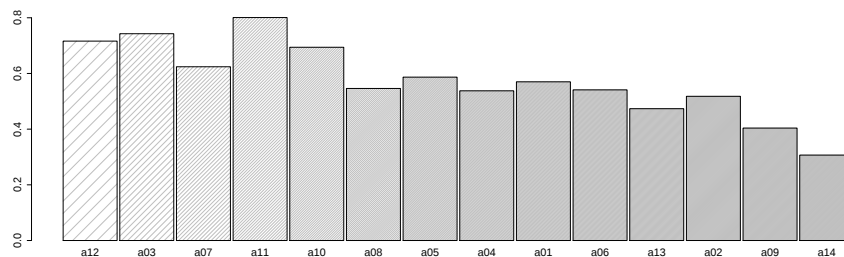


FIGURE 3.8 – Diagramme à barres des scores des alternatives.

le composant *plotAlternativesValues* (voir figure 3.8), ainsi que sous la forme d'un classement via *plotAlternativesValuesPreorder*. Pour terminer, le module *rankAlternativesValues* est utilisé pour obtenir les rangs des alternatives en fonction de leur score global.

On peut observer que l'alternative a11 (P309-16) obtient la valeur la plus élevée et peut ainsi être considérée comme le meilleur compromis pour Thierry dans ce modèle.

Plus de détails sur la partie PROMETHEE du workflow sont visibles à la figure 3.9. Ce sous-workflow est composé de 7 éléments. Le module *PrometheePreference* calcule, sur

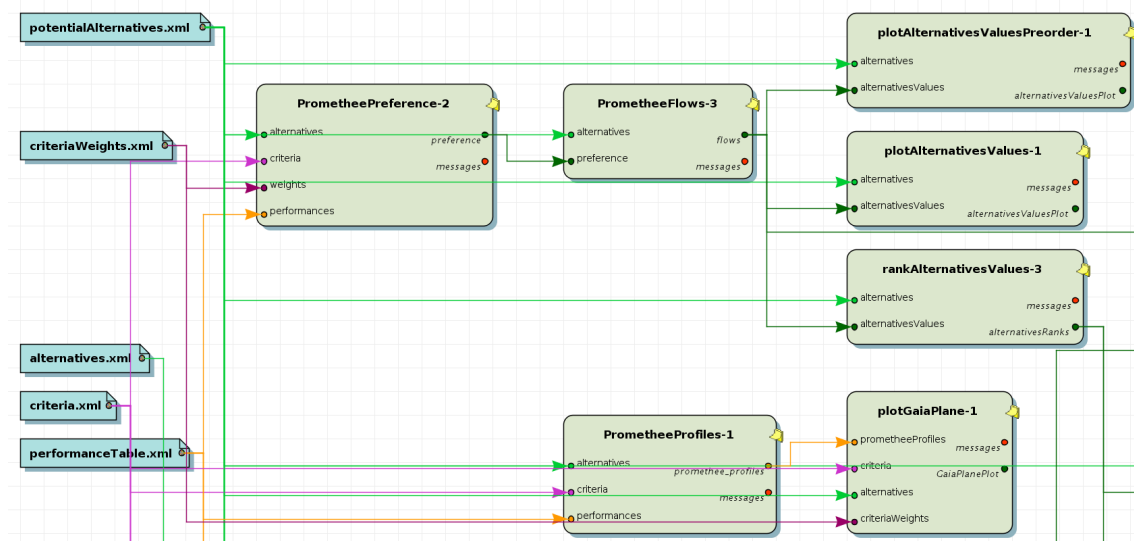


FIGURE 3.9 – La partie PROMETHEE du workflow.

base d'un tableau de performances, de critères, d'alternatives, de seuils de discrimination et de poids de critères, un indice de préférence entre toutes paires d'alternatives. L'analyste élicite ces préférences directement auprès de Thierry en respectant la sémantique spécifique à la procédure PROMETHEE de ces paramètres. Le résultat de cette détermination est résumé dans le tableau 3.2. L'indice de préférence est ensuite utilisé pour calculer le flot net via le composant *PrometheeFlows*, qui est utilisé pour classer les alternatives. Comme précédemment, le préordre des alternatives est représenté graphiquement à l'aide de *plotAlternativesPreorder*, un diagramme à barres des flots nets est construit via le module *plotAlternativesValues* (voir figure 3.10), et les rangs des alternatives sont obtenus par le composant *rankAlternativesValues*.

	coût ($g1$, €)	accél. ($g2$, s)	reprise ($g3$, s)	freins ($g4$)	tenue de route ($g5$)
poids (%)	40	20	20	10	10
indifférence	500	1	0.5	1	1
préférence	2000	1.5	1	2	2

TABLE 3.2 – Les informations intra- et inter-critères.

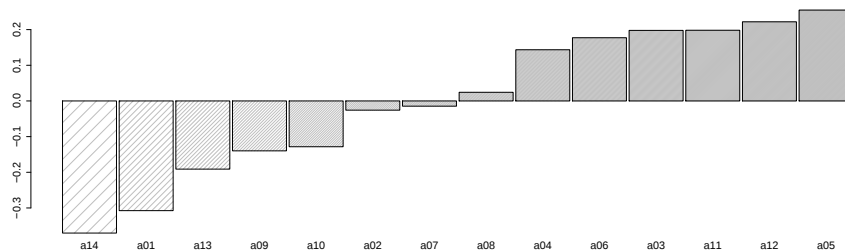


FIGURE 3.10 – Diagramme à barres des flots nets.

D'après les flots nets obtenus par la méthode PROMETHEE, la voiture a5 (Colt) est considérée comme étant le meilleur compromis pour Thierry, étant donné ce modèle de ses préférences.

La troisième partie du workflow consiste en la comparaison de ces deux résultats. Elle est représentée de manière plus détaillée à la figure 3.11. Dans une première étape (haut

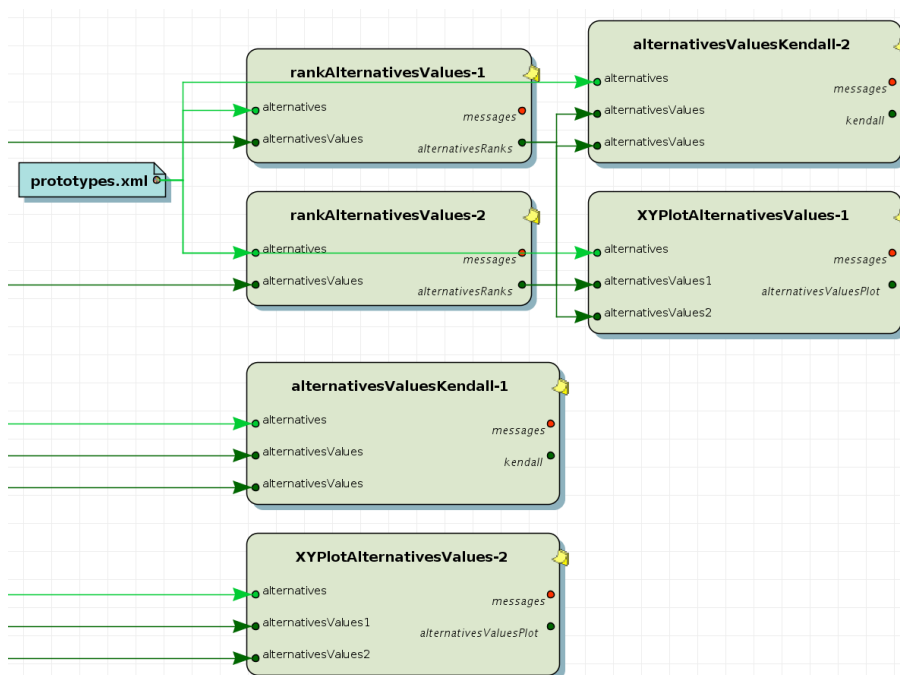


FIGURE 3.11 – Comparaison des deux méthodes.

de la figure 3.11) Thierry est intéressé de savoir comment les 5 voitures qu’il a classé initialement sont ordonnées par la méthode PROMETHEE. A cette fin, l’analyste restreint le classement ACUTA à ces 5 voitures. Le taux de Kendall est ensuite calculé entre ces deux préordres à l’aide du module *alternativesValuesKendall*. On obtient une valeur de 0.8, qui signifie qu’il existe une inversion dans les deux classements (a09 est placé avant a18 par PROMETHEE, alors que Thierry avait indiqué préférer a18 à a09).

Ensuite, les deux classements de PROMETHEE et d’ACUTA sont comparés dans le bas de la figure 3.11. Une nouvelle fois le taux de Kendall est calculé, et les inversions entre les deux classements sont visualisés à l’aide du module *XYPlotAlternativesValues* (voir figure 3.12). L’axe des abscisses représente les rangs par PROMETHEE, alors que l’axe des ordonnées représente les rangs par ACUTA. Un grand nombre d’inversions sont visibles entre ces deux classements, ce qui est aussi confirmé par le taux de Kendall de 0.47.

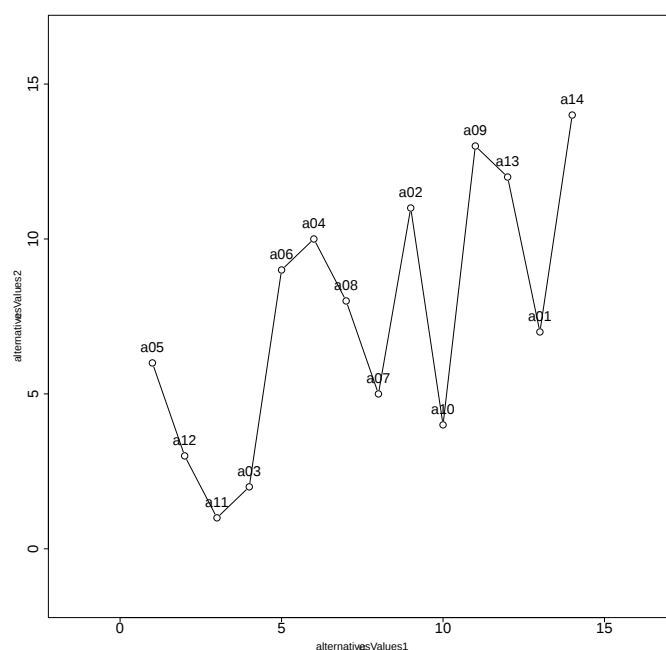


FIGURE 3.12 – Graphe XY représentant les deux classements (abscisse : rangs par la méthode PROMETHEE, ordonnée : rangs par ACUTA).

Plus de détails sur diviz et cet exemple peuvent être trouvés dans Meyer et Bigaret (2012). Par ailleurs, le workflow présenté dans cette section peut être téléchargé sur le site web de diviz¹⁰ et le lecteur intéressé pourra facilement l’importer dans diviz.

3.3.3 Conclusions et perspectives

Les efforts que nous menons sur le développement et la promotion de diviz sont très gratifiants. En effet, de plus en plus d’enseignants adoptent cet outil pour leurs cours d’AMCD. Personnellement nous utilisons diviz dans plusieurs formations, avec des publics différents, et nous remarquons clairement que la compréhension des procédures d’AMCD à la fin des cours en est nettement améliorée. Concernant les chercheurs, l’article Meyer et

10. <http://www.decision-deck.org/diviz/workflow.URPDM2010SpecialIssue.html>

Bigaret (2012) a incité plusieurs d'entre-eux à proposer d'inclure leurs algorithmes dans l'environnement de services-web XMCDa ainsi que dans diviz.

La plateforme diviz n'est pas un outil figé. La prochaine étape de développement consistera à intégrer des modèles de workflows plus complexes, permettant de créer des boucles itératives ou des nœuds de tests conditionnels.

*

La prochaine section traite d'un travail en cours. Nous estimons cependant qu'il serait intéressant de montrer ici les résultats préliminaires, étant donné qu'ils font appel aux outils que nous avons présenté dans ce chapitre. Il s'agit d'une tentative de modélisation du processus d'AMCD à l'aide de techniques issues de l'ingénierie dirigée par les modèles.

3.4 Modélisation du processus d'AMCD

Le processus d'AMCD est en général guidé par l'analyste, qui doit faire preuve d'une grande expertise lors de la gestion de ses différents moments clés. Le déroulement du processus dépend donc très fortement de son expérience personnelle, et de sa connaissance des techniques d'AMCD qu'il met en œuvre.

Un grand nombre de sujets de recherche en AMCD se focalisent sur les procédures d'agrégation ou d'élicitation des préférences. Celles-ci ne représentent cependant qu'une petite partie du processus d'aide. Celui-ci est en général composé d'un certain nombre d'activités, comme par exemple, la structuration du problème, la définition des alternatives de décision, la structuration d'une hiérarchie de critères, l'évaluation des alternatives sur les critères par des experts, l'élicitation des préférences du ou des décideurs, l'utilisation d'une procédure d'agrégation, l'étude de la robustesse de la conclusion et la proposition d'une recommandation de décision. On peut donc facilement imaginer que ce processus est assez complexe, et dépend fortement de la situation pratique et de la méthode de résolution choisie.

Cependant, à notre connaissance, peu de travaux ont été menés à ce jour pour décrire ce processus et le formaliser. Il est intéressant de citer Tsoukiàs (2007), qui a abordé ce problème en proposant un modèle assez général du processus et des concepts clés qui y interviennent. Sans désir d'exhaustivité, d'autres travaux comme ceux de Belton et Stewart (2001), Bana e Costa *et al.* (1999), Brown (1989) et Pidd (1988) proposent également des recommandations pour un processus d'aide à la décision. Cependant, aucune de ces tentatives ne propose un modèle détaillé, supporté par des outils opérationnels.

Les avantages d'avoir un processus clairement documenté et modélisé sont nombreux. Tout d'abord, l'analyste aurait à sa disposition une aide pour le guider dans le processus, et qui lui permettrait de clarifier son discours face aux différents intervenants. Ensuite, en cas de doutes, le décideur serait plus enclin à accepter la recommandation de décision, si elle se basait sur un processus clairement décrit. En effet, une traçabilité des différentes étapes lui donnerait accès à des justifications qui expliqueraient les conclusions retenues.

Ces constats sont le point de départ des travaux que nous présentons dans cette section, et que nous avons mené en collaboration avec Vanea Chiprianov¹¹ et Jacques Simonin¹².

11. Vanea Chiprianov, Télécom Bretagne, France

12. Jacques Simonin, Télécom Bretagne, France

Nous proposons de modéliser d'abord le processus d'AMCD en utilisant des techniques issues de l'ingénierie dirigée par les modèles, i.e. via un langage standard de modélisation. Ce processus générique peut ensuite être instancié pour des problèmes de décision particuliers. Par ailleurs, le modèle doit permettre l'appel d'outils de support à différents moments du processus.

La section se structure de la façon suivante. Tout d'abord, en section 3.4.1 nous présentons quelques éléments clés du modèle ainsi que la façon dont il est en train d'être créé. Le détail de la version actuelle du modèle peut être trouvé dans le rapport technique [Chiprianov et al. \(2012\)](#). Ensuite nous présentons une instanciation du processus à l'exemple classique du choix de voiture à la section 3.4.2 et illustre l'utilisation d'outils informatiques pour soutenir le processus.

3.4.1 Aperçu du modèle

Un *modèle de processus* est défini comme une représentation de l'enchaînement d'activités ou d'actions exécutées par des acteurs sur des données d'entrée en vue de développer des artefacts de sortie (d'après les définitions de processus de [Ramsin et Paige \(2008\)](#) et de modèle de [Rothenberg \(1989\)](#)).

Nous choisissons ici le langage de modélisation appelé SPEM (Object Management Group's Software Process Engineering Metamodel) de l'OMG (2008) pour décrire le processus d'AMCD. Il s'agit d'un langage qui permet de représenter des activités et des tâches d'un processus, tout en faisant intervenir de multiples rôles.

Les objectifs de cette modélisation sont multiples :

- fournir une représentation aussi complète que possible du processus, en fédérant plusieurs écoles de pensées, ainsi que différentes façons de travailler des analystes ;
- augmenter la lisibilité pour le décideur et l'analyste à travers une représentation graphique du processus, qui peut être consultée à tout moment de son déroulement ;
- guider l'analyste, qui se voit libéré de certaines responsabilités ;
- donner la possibilité de connecter facilement des services ou workflows de services au processus, en vue d'en automatiser une partie ;
- augmenter la traçabilité et l'audit de recommandations de décision.

La méthode de construction du modèle est une *spirale itérative* (voir figure 3.13) avec deux activités principales qui se suivent : la modélisation et la confrontation à des experts de l'AMCD. La première version du modèle (V01) se base sur les travaux de [Tsoukiàs \(2007\)](#). Celle-ci a ensuite été confrontée à des chercheurs en AMCD lors de trois itérations majeures. Lors de la dernière itération, le modèle a été présenté dans le cadre d'un d'un workshop d'AMCD ([Bigaret et al., 2012](#)). Après chaque itération, des améliorations et révisions proposées par les "experts" sont intégrées au modèle. Il est important de noter que le travail est loin d'être abouti, et nous pensons que plusieurs versions supplémentaires seront encore nécessaires avant d'arriver à un consensus stable et accepté par un grand nombre de chercheurs du domaine.

SPEM offre la possibilité d'organiser un processus de manière hiérarchique, à travers trois types de niveaux. Le premier type est composé de *sous-processus*, le second d'*activités* et le troisième de *tâches*. Le premier type de niveaux est utilisé pour décrire les grands sous-processus qui composent le processus général d'AMCD. Le second type de niveaux est utilisé pour décrire des activités composées de sous-activités et de tâches, alors que le celui des tâches représente le niveau final qui ne doit plus être redécomposé. Ce dernier contient aussi des détails sur les données manipulées ou les intervenants dans le processus,

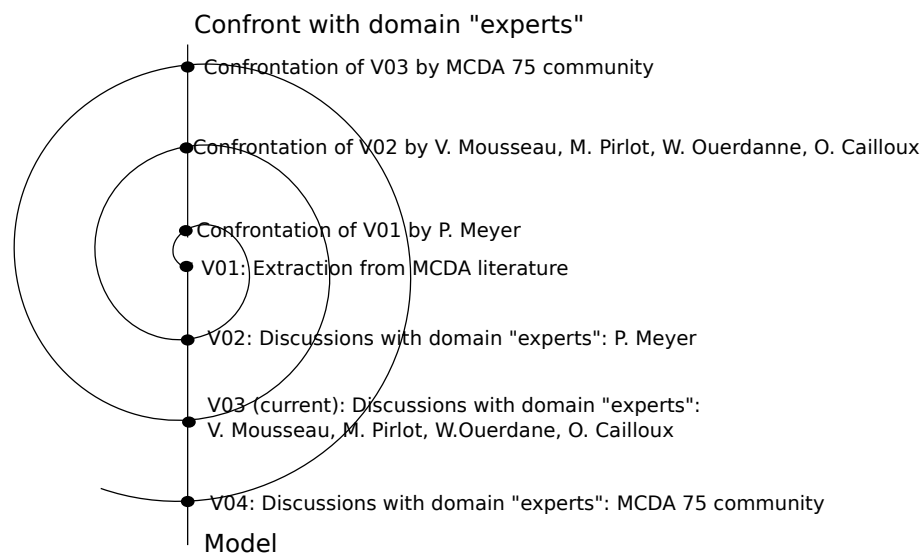
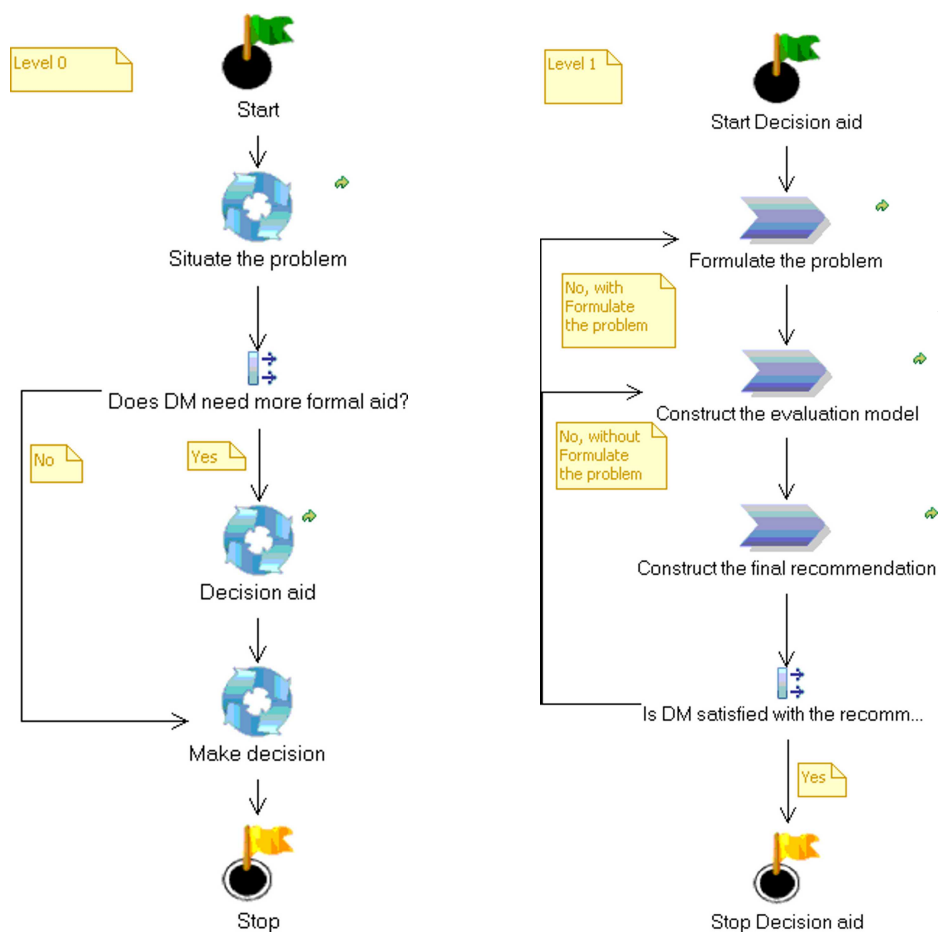


FIGURE 3.13 – La méthode de construction du modèle du processus d'AMCD.

ainsi que la connexion à des procédures d'AMCD (comme par exemple les services-web XMCD).

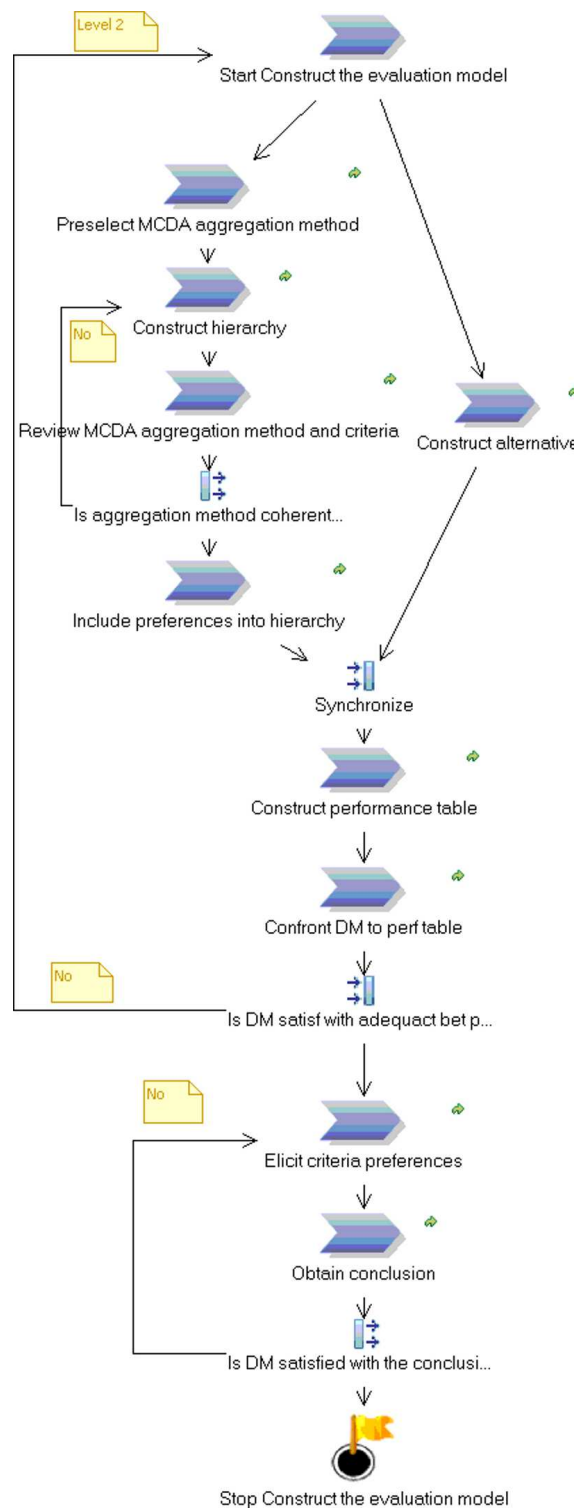
FIGURE 3.14 – Le processus général (gauche) et le sous-processus *Decision aid* (droite).

Le niveau le plus abstrait du modèle est présenté dans la partie gauche de la figure 3.14. Après que le processus d'aide ait démarré, un premier sous-processus débute (*Situating the problem*) qui consiste entre autres à déterminer les rôles de décideur et d'analyste ainsi que les préoccupations générales du décideur. Ensuite, si le décideur pense avoir besoin de plus d'aide avant de prendre une décision, le sous-processus formel *Decision aid* est entamé. Si non, il peut prendre sa décision, et le processus général se termine.

Pour rappel, le détail des sous-processus et des activités est disponible dans [Chiprianov et al. \(2012\)](#). Nous ne présentons ici qu'une description assez haut niveau du sous-processus *Decision aid* (à droite de la figure 3.14). En premier lieu, dans ce sous-processus, le décideur est confronté à une activité appelée *Formulate the problem*. Il s'agit d'y déterminer entre autres à quel type de problématique le décideur est confronté, ainsi que les premiers éléments du problème de décision, comme certaines options ou points de vue importants pour le décideur. Le processus nous amène ensuite dans l'activité *Construct the evaluation model*, qui est détaillée à la figure 3.15.

Cette activité représente l'activité principale du processus d'AMCD, dans laquelle les alternatives et les critères sont construits, les préférences du décideur sont élicitées et une conclusion est fournie. Une présentation des différentes activités qui la composent est donnée ci-après :

1. *Start Construct the evaluation model* : le point de départ de l'activité de construction du modèle d'évaluation ;
2. *Preselect MCDA aggregation method* : l'analyste, en fonction de l'information recueillie précédemment, choisit un certain nombre de procédures d'agrégation d'AMCD (ou une seule) adaptées ;
3. L'activité *Construct hierarchy* : la construction complexe de la hiérarchie des critères, sur base de différents indicateurs et attributs ;
4. *Review MCDA aggregation method and criteria* : après la construction de la hiérarchie de critères, l'analyste a probablement une meilleure connaissance du problème, et serait peut-être tenté de réviser l'ensemble des procédures d'agrégation choisies (en vue de le réduire ou de l'augmenter) ;
5. *Is aggregation method coherent ?* : si certaines méthodes d'agrégation ne sont pas cohérentes avec la hiérarchie construite, réviser la hiérarchie ou mettre à jour les procédures retenues ;
6. *Include preferences into hierarchy* : les préférences du décideur sont incluses dans la hiérarchie de critères (par exemple sur la façon d'agréger certains indicateurs) ;
7. *Construct alternative* : indépendamment des activités concernant les critères, cette activité construit l'ensemble des alternatives potentielles ;
8. *Construct performance table* : le tableau de performance est construit en évaluant les alternatives sur les critères ;
9. *Confront DM to perf table* : le décideur est confronté au tableau ou à une représentation synthétique de celui-ci ;
10. *Elicit criteria preferences* : l'analyste détermine ensemble avec le décideur ses préférences inter- et intra-critères, en fonction des procédures d'agrégation retenues ;
11. *Obtain conclusion* : une ou plusieurs conclusions sont obtenues en utilisant la ou les procédures d'agrégation retenues ;
12. *Is DM satisfied with the conclusion ?* : si le décideur n'est pas satisfait avec les conclusions proposées, des modifications peuvent être apportées aux préférences élicitées ;

FIGURE 3.15 – L'activité *Construct the evaluation model*.

13. *Stop Construct the evaluation model* : l'activité *Construct the evaluation model* se termine.

La sortie principale de cette activité est une conclusion, qui représente une version préliminaire de la recommandation de décision. Dans une étape suivante, l'activité *Construct the final decision recommendation* sera utilisée pour construire la recommandation finale,

en faisant des post-analyses (robustesse, stabilité, ...).

Si le décideur est satisfait avec la recommandation finale le sous-processus se termine. Si non, une nouvelle itération peut être envisagée pour modifier le modèle d'évaluation ou même la formulation initiale du problème (cf. figure 3.14 à droite).

Il est évident que toutes les activités que nous avons décrites ici se décomposent en un grand nombre de sous-activités et de tâches. Elles dépendent d'un certain nombre de choix que l'analyste fait tout au long du processus. Leur description exacte, ainsi que la façon dont elles s'enchaînent font partie d'un travail de recherche que nous continuons d'alimenter.

Un certain nombre de tâches de ces activités générales peuvent être directement connectées à des outils informatiques de soutien au processus d'AMCD. Nous illustrons cette possibilité dans la section suivante.

3.4.2 Une instance du processus d'AMCD

L'instanciation du processus général présenté à la section précédente à un problème de décision concret permet de guider un analyste dans la résolution de ce dernier. L'exemple choisi ici est à nouveau celui de l'étudiant Thierry qui aimerait choisir une voiture, qui a déjà été présenté et traité aux sections 2.2.3, 3.1.2 et 3.3.2.

Sans rentrer dans plus de détails, le sous-processus *Situate the problem* de la partie gauche de la figure 3.14 aide à déterminer dans une première étape les rôles de décideur (Thierry) et d'analyste. Ensuite nous supposons que Thierry désire avoir une aide pour son problème, ce qui nous mène au sous-processus *Decision aid* sur la droite de la figure 3.14.

L'activité *Formulate the problem* permet à l'analyste d'identifier qu'il est face à un problème de classement (classer les voitures de la meilleure à la plus mauvaise) ou à un problème de choix (sélectionner la meilleure voiture pour Thierry). Dans une discussion avec Thierry, il identifie également que celui-ci a trois préoccupations majeures pour ce problème : minimiser les coûts, maximiser la performance de la voiture, et maximiser sa sécurité.

L'analyste passe alors à l'activité *Construct the evaluation model* qui est détaillée à la figure 3.15. Dans l'activité *Preselect MCDA aggregation method*, sur base de la discussion initiale que l'analyste a eu avec Thierry (celui-ci parle de "compenser" des points forts avec des points faibles), il décide d'utiliser un modèle d'agrégation se basant sur la théorie de la valeur multi-attribut. Pour commencer, il choisit d'utiliser un modèle d'agrégation additif.

L'analyste passe ensuite à l'activité *Construct hierarchy*, dans laquelle il met en place les différents critères qui permettent à Thierry de comparer les voitures entre-elles. La discussion mène Thierry à définir un critère de coût (g_1) comme une agrégation de divers indicateurs comme le prix d'achat de la voiture, des coûts annuels, des taxes, et des estimations de coûts liés à l'assurance du véhicule et à sa consommation d'essence. Comme Thierry veut acheter une voiture puissante, un critère mesurant son accélération (g_2) (en secondes) et un critère de reprise (g_3) (en secondes) sont rajoutés. Pour terminer, comme indiqué précédemment dans l'activité *Formulate the problem*, Thierry porte également de l'importance à la sécurité, ce qui amène l'analyste à rajouter deux critères mesurant la qualité du freinage (g_4) et la qualité de la tenue de route (g_5). Ces derniers critères sont des agrégats d'un certain nombres d'indicateurs choisis par Thierry (voir Bouyssou *et al.* (2000) pour de plus amples détails sur la construction de cette hiérarchie).

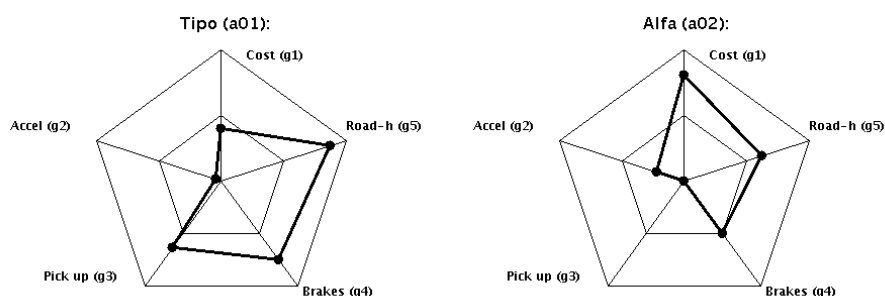


FIGURE 3.16 – Graphes en étoile de deux alternatives obtenus par le service-web `plotStarGraphPerformanceTable`.

L’analyste vérifie alors si la hiérarchie ainsi construite est compatible avec un modèle de fonction de valeur additif et décide ensuite de passer à l’activité *Include preferences into the hierarchy*. Dans celle-ci, l’analyste demande à Thierry si les 5 critères sont à maximiser ou à minimiser sur leurs échelles d’évaluation. Thierry indique que le critère de “coût” et les critères de performance “accélération” et “reprise” doivent être minimisés, alors que les critères de sécurité “freinage” et “tenue de route” sont à maximiser.

En parallèle, l’analyste et Thierry doivent décider de l’ensemble des alternatives potentielles pour l’achat dans l’activité *Construct alternative*. Comme Thierry est un étudiant, il n’a pas le budget pour s’acheter une voiture neuve, et décide d’explorer les voitures d’occasion de moins de 4 ans du marché. Comme par ailleurs il vit en ville et qu’il n’a pas de garage, il ne veut pas une voiture trop attractive pour les voleurs. Ces contraintes mènent à restreindre le choix de Thierry aux 14 voitures du tableau 3.1.

L’activité suivante est de construire le tableau de performance. En pratique, ceci signifie que les 14 voitures doivent être évalués sur tous les indicateurs et les critères définis précédemment. Le résultat de cette évaluation se trouve dans le tableau 3.1.

Ensuite Thierry est confronté à cette représentation de son problème dans l’activité *Confront DM to performance table*. L’analyste lui présente les données, et par ailleurs, fait appel au service-web `plotStarGraphPerformanceTable` qui permet de représenter chaque alternative sous la forme d’un *graphe en étoile* (voir figure 3.16). Cette activité permet à Thierry de mieux appréhender le problème sous-jacent et de comparer les alternatives entre-elles, en vue d’augmenter la compréhension de ses préférences.

Le service-web est appelé à l’aide de la plateforme diviz, où un workflow des services est construit progressivement tout au long du processus d’aide à la décision. À cette étape, ce workflow contient en plus des données du problème uniquement le module `plotStarGraphPerformanceTable` (voir figure 3.17).

Nous supposons ici que Thierry est satisfait avec les données et que l’analyste considère que le modèle additif est toujours approprié pour la modélisation des préférences. L’activité suivante du processus d’AMCD est *Elicit criteria preferences*, où les préférences de Thierry sont ajoutées aux définitions des critères. Dans le cas du modèle retenu par l’analyste, il s’agit d’élucider des fonctions de valeur qui représentent la satisfaction de Thierry par rapport aux différentes valeurs prises par les alternatives sur les critères.

Dans sa discussion avec Thierry, l’analyste détermine que Thierry connaît bien certaines voitures et qu’il est capable d’exprimer des préférences sous la forme d’un classement de certaines d’entre-elles (Bouyssou *et al.*, 2006) :

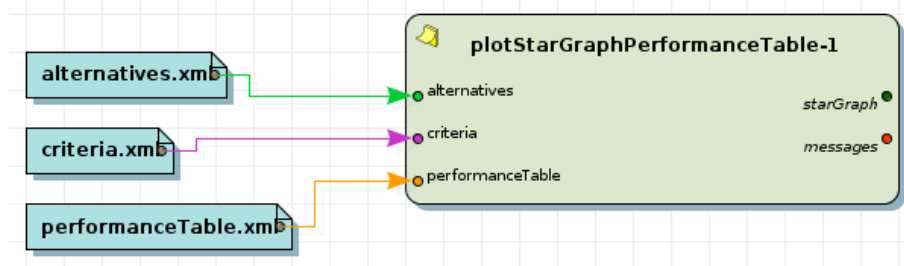


FIGURE 3.17 – La première étape de la construction progressive du workflow diviz supportant le processus d'AMCD.

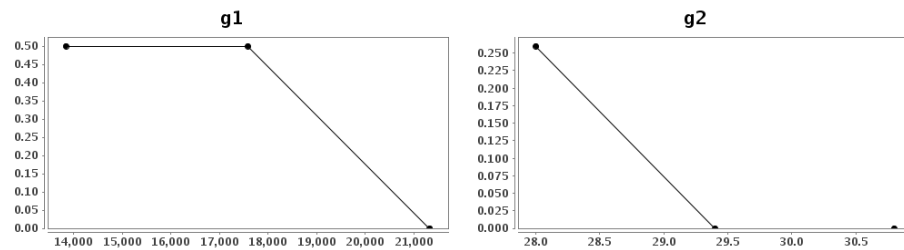


FIGURE 3.18 – Représentations graphiques des fonctions de valeur pour les critères g_1 (coût) et g_2 (accélération), obtenues à l'aide du service `plotValueFunctions`.

P309-16 $\succ$ Sunny $\succ$ Galant $\succ$ Escort $\succ$ R21t;

Ceci amène l'analyste à choisir l'approche indirecte UTA de [Jacquet-Lagrèze et Siskos \(1982\)](#) d'élicitation des fonctions de valeur. Il continue la construction du workflow diviz en y ajoutant le module UTA. La sortie de cette procédure donne un ensemble de fonctions de valeur qui représentent les préférences de Thierry par rapport au classement des 5 voitures fourni en entrée. L'analyste confronte alors Thierry à une représentation graphique de ses préférences à l'aide du module `plotValueFunctions` dans diviz (voir figure 3.18) afin de les valider. Le workflow diviz obtenu à cette étape est représenté à la figure 3.19.

Afin d'obtenir un classement des 14 voitures potentielles, le modèle de fonction de valeurs additives doit être appliqué au tableau de performances. Ceci est fait dans le

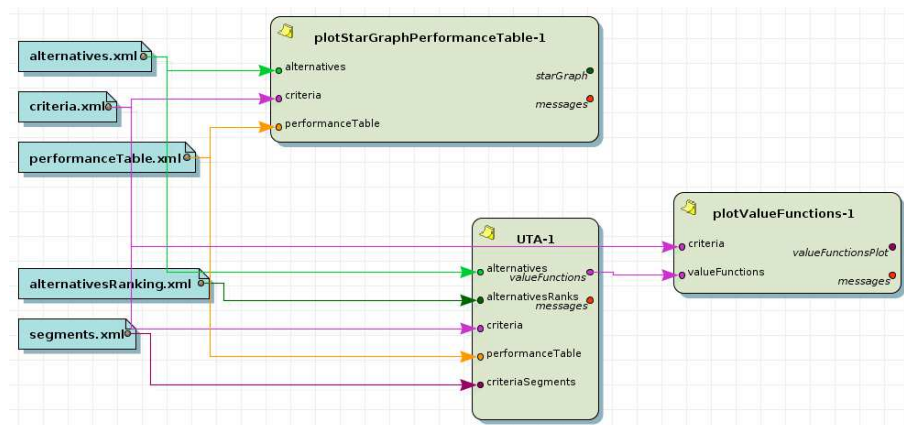


FIGURE 3.19 – La deuxième étape de la construction progressive du workflow diviz.

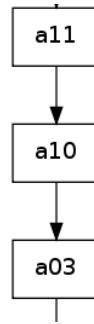


FIGURE 3.20 – Début du classement des voitures.

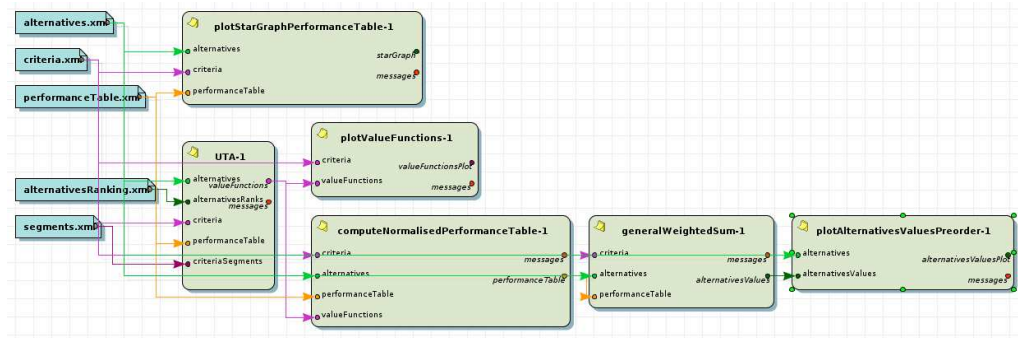


FIGURE 3.21 – La troisième étape de la construction progressive du workflow diviz.

processus d'AMCD à l'activité *Obtain conclusion*. Le workflow diviz est complété afin de transformer les données du tableau de performance en des degrés de satisfaction de Thierry (à l'aide de `computeNormalizedPerformanceTable`), avant qu'une agrégation additive des performances des voitures ne soit effectuée (grâce à `generalWeightedSum`), dont la sortie est le score de chaque voiture à travers le modèle choisi.

Thierry est ensuite confronté à une représentation graphique du résultat (grâce à `plotAlternativesValuesPreorder`), et nous supposons qu'il est satisfait avec la première voiture du classement (a11, la Peugeot 309-16) de la figure 3.20. Le workflow diviz construit par l'analyste à cette étape est représenté à la figure 3.21.

Il est intéressant de noter la correspondance entre le processus que nous proposons et le workflow diviz construit de manière progressive. Celle-ci est présentée dans le tableau 3.3. Dans ce cas, il n'y a que 3 activités du processus qui ont des services correspondants. D'autre part, une des activités nécessite une série de services.

process activity	diviz service
<i>Confront the DM to performance table</i>	<code>plotStarGraphPerformanceTable</code>
<i>Elicit criteria preferences</i>	<code>UTA</code> ; <code>plotValueFunctions</code>
<i>Obtain conclusion</i>	<code>computeNormalizedPerformanceTable</code> ; <code>generalWeightedSum</code> ; <code>plotAlternativesValuesPreorder</code>

TABLE 3.3 – Correspondance entre les activités du processus d'AMCD et les services-web XMCDa pour le problème de Thierry.

3.4.3 Perspectives

La construction de ce modèle du processus d'AMCD est un travail de longue haleine. Les activités présentées ici sont décomposées pour la plupart en des sous-activités, qui ensuite sont décomposées en des tâches, en fonction de certains choix que l'analyste a fait dans des étapes précédentes. Une des grandes difficultés consiste à proposer un modèle qui réunit les deux grands courants méthodologiques de l'AMCD.

D'autre part, l'opérationnalisation de ce processus requiert qu'on mette en correspondance un certain nombre de tâches avec des logiciels d'AMCD. Nous avons montré cette correspondance au niveau d'activités assez génériques, mais en pratique, elle se fait au niveau des tâches (qui ne sont pas présentées ici). L'étude de cette correspondance s'appelle leur "alignement" en ingénierie des modèles. Il permet de mesurer la correspondance entre une partie du modèle et des outils informatiques existant, en vue de le valider le processus.

Chapitre 4

Applications de l'AMCD

Sommaire

4.1	Au sujet de la construction d'indicateurs de bien-être en économie	68
4.1.1	Cadre expérimental	70
4.1.2	Résultats	73
4.1.3	Conclusions et perspectives	76
4.2	Construction d'une échelle de risque territorial prenant en compte plusieurs critères et experts	76
4.2.1	Méthodologie multicritère et multi-décideurs pour la construction d'une échelle de risque territorial	77
4.2.2	Exemple illustratif	79
4.2.3	Conclusion	83

Ce chapitre représente le troisième volet de notre recherche, et détaille deux de nos contributions à l'utilisation de techniques d'AMCD dans des applications. Cette démarche est importante pour valoriser nos travaux ainsi que pour faire connaître l'AMCD au delà des frontières de la recherche opérationnelle. Nous commençons par présenter une étude exploratoire sur la construction d'indicateurs de bien-être en économie dans la section 4.1 (Meyer et Ponthière, 2011). Ensuite, dans la section 4.2, nous présentons comment on pourrait utiliser la méthodologie proposée en section 2.3 pour construire une échelle de risque territoriale (Cailloux *et al.*, 2013) tenant compte de l'avis de plusieurs experts.

À côté de cela, nous avons également participé ces dernières années à la mise en œuvre de l'AMCD dans deux autres applications. Tout d'abord nous avons proposé d'inclure les préférences des opérateurs dans les modèles de contrôle automatiques de drones. Ce travail que nous avons mené en collaboration avec Pritesh Narayan¹ et Duncan Campbell² a produit des résultats très encourageants qui sont publiés dans Narayan *et al.* (2013). Nous y montrons que l'inclusion des préférences des opérateurs dans ce contrôle automatique améliore sensiblement les trajectoires de vol, et augmente par la même occasion la confiance que les opérateurs ont dans l'algorithme de pilotage autonome. Ensuite, une deuxième application concerne l'évaluation économique de techniques de télémédecine. L'introduction des technologies de l'information et de la communication au sein des organisations de

1. Pritesh Narayan, University of the West of England, Bristol, Grande Bretagne

2. Duncan Campbell, Queensland University of Technology, Brisbane, Australie

santé représente un enjeu stratégique important dans le contexte actuel des transformations profondes auquel est confronté le système de santé français. Dans ce travail, que nous avons mené avec [Gérald-Réparate Retali](#)³ et [Myriam Le Goff Pronost](#)⁴, nous avons tenté de démontrer que la prise en compte des préférences des patients et des médecins au sujet de différentes options de télémédecine a un intérêt pour aider un décideur institutionnel à choisir quelle technologie de santé sera opérationnalisée. La méthodologie proposée a été testée sur des patients et des médecins dans le cadre de l'implantation d'unités de dialyse médicalisées télésurveillées en Côtes d'Armor. Ces travaux s'inscrivent dans le cadre des travaux de thèse de doctorat de [Gérald-Réparate Retali](#) (voir aussi [Retali et al., 2011a,b](#)). Nous nous permettons ici de ne pas détailler ces deux dernières applications, car elles nécessitent de longues introductions aux domaines métiers.

4.1 Au sujet de la construction d'indicateurs de bien-être en économie

Les économistes s'intéressent depuis bien longtemps à la mesure du bien-être des humains. Régulièrement ils tentent de répondre aux questions suivantes :

- Est-ce que toutes les sociétés humaines se valent ?
- Si non, est-ce que le bien-être est perçu différemment d'une société à une autre ?

Les principales difficultés liées à leur étude sont l'interprétation subjective des questions et la nature multidimensionnelle des “objets” qu'il s'agit de mesurer.

Concernant la nature subjective de la qualité de la vie, il n'est pas surprenant de considérer que deux individus perçoivent de manière très différente le concept de “bonne société”. Même si on arrivait à définir ce terme de manière très objective, la perception qu'en auraient des individus serait biaisée par leurs systèmes de valeurs leur expériences passées différents. Cette hétérogénéité dans la perception des individus des conditions de vie est évidemment très problématique pour la construction d'indicateurs de bien-être. Cela tend à remettre en question des indicateurs basés sur des agrégations qui ne tiennent pas compte des préférences des individus. Des indicateurs de ce type incluent le *United Nations' Human Development Index* – le HDI – ([United Nations Development Program \(UNDP\), 1990](#)) et l'*Index of Economic Well-Being* d'[Osberg et Sharpe \(2002, 2005\)](#).

Il est communément accepté qu'une bonne qualité de vie est liée à plusieurs facteurs, comme par exemple un pouvoir d'achat élevé, beaucoup de temps libre, et une longue vie en bonne santé. Cette décomposition traduit bien la nature multidimensionnelle des standards de vie et on peut facilement imaginer des interactions entre ces facteurs (complémentarité ou redondance).

Pour illustrer ces propos, on peut prendre l'exemple du HDI, qui est construit en agrégeant un indice d'espérance de vie, un indice d'éducation, et un indice de PIB via une moyenne. Le [tableau 4.1](#), qui présente le classement des 10 premiers pays en fonction de leur HDI en 2007, souligne les difficultés liées à la mesure des standards de vie.

Tout d'abord, comme cela a été souligné par [Dasgupta \(1993\)](#), le classement fourni par le HDI souffre de l'arbitraire lié aux poids des 3 dimensions. Il n'est par exemple pas évident de comprendre pourquoi des poids égaux sont affectés à ces 3 indicateurs. Cette critique est renforcée par le fait qu'une petite variation dans ces poids peut avoir un impact significatif sur le classement obtenu.

3. [Gérald-Réparate Retali](#), Télécom Bretagne, France

4. [Myriam Le Goff Pronost](#), Télécom Bretagne, France

pays	espérance de vie	éducation	PIB	HDI
Islande	0.941	0.978	0.985	0.968
Norvège	0.913	0.991	1.000	0.968
Australie	0.931	0.993	0.962	0.962
Canada	0.921	0.991	0.970	0.961
Irlande	0.890	0.993	0.994	0.959
Suède	0.925	0.978	0.965	0.956
Suisse	0.938	0.946	0.981	0.955
Japon	0.954	0.946	0.959	0.953
Pays Bas	0.904	0.988	0.966	0.953
France	0.919	0.982	0.954	0.952

TABLE 4.1 – Classement international, Human Development Index, 2007

Ensuite, il n'est pas clair pourquoi l'HDI repose sur une moyenne pondérée des trois indices. En effet, on peut facilement imaginer que dans la perception des individus, il existe certaines interactions entre les différentes dimensions qui composent le bien-être, ce qui ne peut être traduit par une moyenne pondérée. Par exemple, on pourrait se demander pourquoi un pays qui est très bon en revenu, moyen en enseignement et mauvais en longévité est considéré comme équivalent à un pays qui est moyen sur tous les facteurs. On pourrait en effet facilement imaginer qu'il existe une interaction positive entre le pouvoir d'achat et la longévité, de façon à ce que le deuxième pays soit classé avant le premier. Cependant, pour prendre en compte ces interactions, il faut s'écarter de l'agrégation par la moyenne pondérée.

L'objectif de cette étude exploratoire, que nous avons menée avec Grégory Ponthière⁵, est de réexaminer la construction d'indices de qualité de vie, en portant une attention particulière à la subjectivité et la nature multidimensionnelle des objets à mesurer. A cette fin, nous proposons d'utiliser la théorie de la valeur multiattribut (MAVT) (Keeney et Raiffa, 1976), en vue d'extraire, de manière empirique, des poids plausibles pour les différentes dimensions qui composent la qualité de la vie. Ces poids modélisent les préférences des individus interrogés. Afin de tenir compte de l'interaction possible entre les différentes dimensions, nous optons pour un modèle utilisant l'intégrale de Choquet comme opérateur d'agrégation (Choquet, 1953). Elle représente une extension naturelle de la moyenne pondérée classique, tout en permettant de modéliser un certain type d'interaction entre les dimensions. Le lecteur intéressé pourra se référer à Grabisch *et al.* (2008) pour une présentation détaillée de l'intégrale de Choquet et de plusieurs méthodes d'élicitation des préférences permettant d'identifier une capacité sur base d'un classement d'alternatives fourni en entrée.

Afin d'illustrer l'utilisation de MAVT pour la construction d'indices de qualité de vie, nous mettons en place une petite expérience, où des étudiants doivent classer des sociétés fictives décrites sur 5 dimensions (longévité, santé, consommation, temps pour les loisirs, qualité de l'environnement). À partir de ces classements, nous extrayons les poids de ces dimensions. Cette expérience nous permet de déduire quelques observations intéressantes qui sont évidemment à valider à plus grande échelle. Ce travail n'a aucune prétention de généralité, mais nous sert uniquement de point de départ pour une étude plus vaste.

Nous terminons cette introduction par une illustration de la diversité des préférences qui peuvent être représentées par une intégrale de Choquet. A cette fin, imaginons un petit exemple composé de 3 individus ayant des préférences différentes sur des sociétés décrites par leur consommation (mesuré en euros par mois) et la longévité (mesure en terme d'espérance de vie à 65 ans).

5. Grégory Ponthière, Ecole normale supérieure, France

Les préférences des 3 individus A , B et C sont résumées dans le tableau 4.2 en terme de capacités dans l'espace consommation/longévité (con/lon). Pour simplifier le discours, nous supposons ici que les fonctions de valeur marginales ont une forme particulière, que nous appelons “en s”. Ces fonctions sont convexes pour des évaluations en dessous d'un niveau de référence, et concaves au dessus. Pour notre cas, ces niveaux sont fixés à 13000 euros par mois pour la consommation et 19 ans pour l'espérance de vie à 65 ans.

individu	fonction de valeur	$\mu(\emptyset)$	$\mu(\text{con})$	$\mu(\text{lon})$	$\mu(\text{con}, \text{lon})$
A	en s	0	0.1	0.2	1
B	en s	0	0.2	0.8	1
C	en s	0	0.8	0.9	1

TABLE 4.2 – Description des préférences des trois individus A , B et C

Les préférences de ces 3 individus diffèrent en terme d'interactions entre les 2 dimensions considérées. A considère que la consommation et la longévité sont complémentaires (i.e. $\mu(\text{con}, \text{lon}) > \mu(\text{con}) + \mu(\text{lon})$), B ne perçoit pas d'interactions entre les 2 dimensions (i.e. $\mu(\text{con}, \text{lon}) = \mu(\text{con}) + \mu(\text{lon})$), alors que C considère que la consommation et la longévité sont redondants (i.e. $\mu(\text{con}, \text{lon}) < \mu(\text{con}) + \mu(\text{lon})$).

Les graphes de la figure 4.1 présentent les courbes d'indifférence des trois individus (l'axe vertical indique la valeur totale d'une société). Même si une consommation et une longévité élevée mènent évidemment à une valeur élevée dans toutes les situations, on observe cependant que les courbes d'indifférence des trois individus sont très différentes.

Le premier graphe de la figure 4.1 souligne l'existence de *complémentarités* entre la consommation et la longévité pour l'individu A . Un haut niveau de valeur totale peut être obtenu uniquement si la consommation et la longévité sont toutes les deux suffisamment élevées. Par exemple pour A , une consommation basse est difficilement compensable par des années d'espérance de vie.

Le second graphe montre les préférences de l'individu B , pour qui il n'existe *aucune interaction* entre les 2 dimensions. D'après lui, il est possible de compenser un niveau de consommation bas par un complément de longévité, afin de garder une valeur élevée pour la société.

Pour terminer, le troisième graphe de la figure 4.1 présente les courbes d'indifférence de l'individu C , pour qui il y a de la *redondance* entre les deux dimensions. C préfère une société qui est très bonne sur une des deux dimensions, par rapport à une société moyenne sur les deux dimensions.

La suite de cette section se structure de la façon suivante. Nous présentons tout d'abord en 4.1.1 le cadre général de l'expérience ainsi que son implémentation pratique. Ensuite, en section 4.1.2 nous présentons quelques résultats et montrons sur un petit exemple leur implication sur un classement de sociétés réelles.

4.1.1 Cadre expérimental

Afin d'illustrer l'utilisation de MAVT pour la construction d'indicateurs de bien-être, nous avons mis en place une expérience que nous avons appelée “*Making the World Better : a simple classroom experiment*”, dans laquelle nous avons soumis à un petit groupe d'individus, un questionnaire standardisé, leur demandant de classer des sociétés fictives multidimensionnelles. De ces réponses, nous testons d'abord l'existence d'un modèle additif pour représenter leurs préférences individuelles. Après avoir montré l'inadéquation de

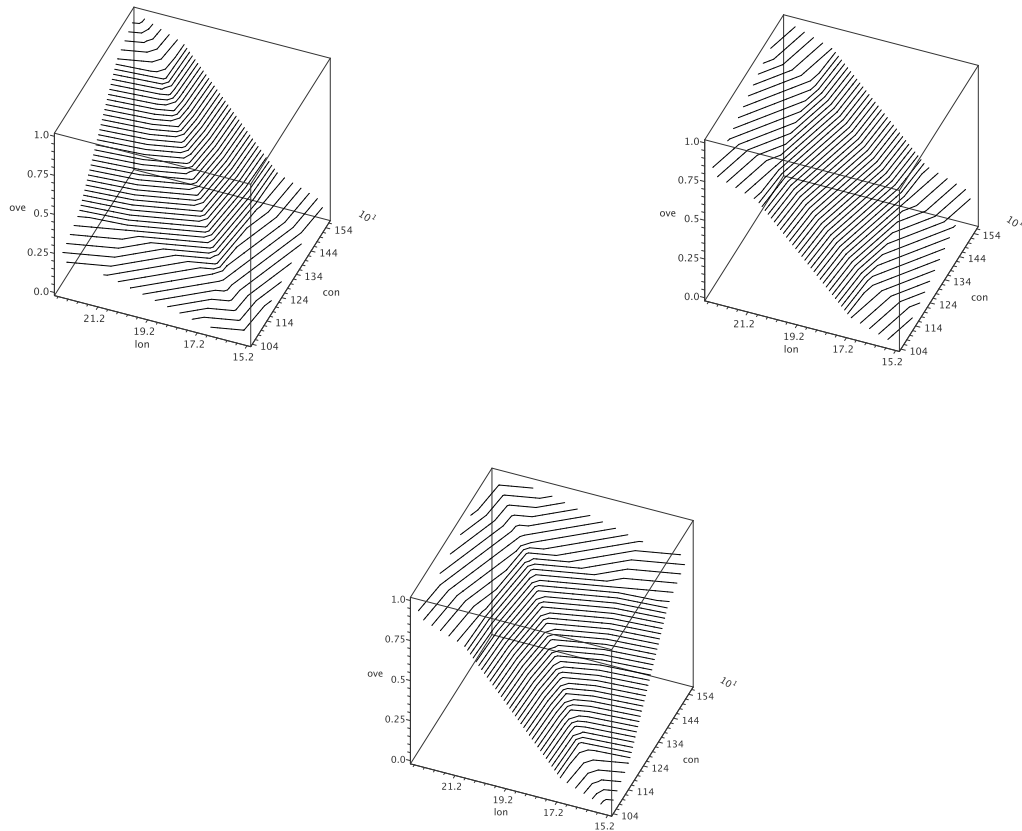


FIGURE 4.1 – Courbes d'indifférence pour les 3 individus fictifs.

ce type de modèle, nous identifions leurs préférences individuelles en utilisant un modèle à base de l'intégrale de Choquet.

Pour préparer cette expérience, nous choisissons d'abord des dimensions pertinentes pour mesurer la qualité de vie. Afin de permettre aux répondants de s'exprimer par rapport aux sociétés qui leur seront présentées, nous sélectionnons les 5 dimensions suivantes (voir aussi l'argumentaire de [Miller \(1956\)](#) sur le nombre de dimensions à retenir) : la *consommation*, la *qualité de l'environnement*, la *santé*, le *temps de loisir* et la *longévité*. Nous sommes tout à fait conscients que la restriction à ces dimensions est une simplification de la situation réelle, mais nous pensons que les sociétés fictives qui en découlent permettront aux répondants de bien se projeter par rapport à leurs préférences.

Pour chacune de ces 5 dimensions, un indicateur est choisi (qu'on appellera par la suite critère). La consommation est mesurée en euros par mois, la qualité environnementale en tonnes de CO₂ émis par habitant, la santé est mesurée par l'espérance de vie en bonne santé à 65 ans (en années), le temps de loisir est mesuré par le nombre d'heures de travail par semaine, alors que la longévité est mesurée par l'espérance de vie à 65 ans (en années). D'autres choix auraient pu être faits, mais ces indicateurs ont l'avantage d'être clairs et faciles à interpréter.

Pour cette expérience nous choisissons de présenter des sociétés fictives aux répondants. Si on leur demandait de classer des pays réels, leurs réponses seraient probablement biaisées par d'autres facteurs que nous ne pourrions pas maîtriser. La construction de ces sociétés

se fait en partant d'une société de référence, dont les évaluations sur les différents critères correspondent à des valeurs moyennes dans la partie du monde que nous étudions ici (l'Europe de l'ouest). Pour chaque critère, quatre niveaux supplémentaires sont introduits. Le niveau "mauvais" (resp. "très mauvais") correspond à une évaluation à 90% (resp. 80%) de l'évaluation de référence, alors que le niveau "bon" (resp. "très bon") correspond à 110% (resp. 120%) de l'évaluation de référence. Ceci nous garantit que les sociétés générées à partir de ces niveaux sont plausibles. Cette construction est résumée dans le tableau 4.3.

Critères	"très mauvaise" société (-20 %)	"mauvaise" société (-10 %)	société de référence	"bonne" société (+ 10 %)	"très bonne" société (+ 20 %)
Consommation (€ par mois)	1040 €	1170 €	1300 €	1430 €	1560 €
Environnement (tonnes de CO ₂ par habitant)	12 tonnes	11 tonnes	10 tonnes	9 tonnes	8 tonnes
Santé (espérance de vie en bonne santé à 65 ans)	12 années	13.5 années	15 années	16.5 années	18 années
Temps de loisir (heures de travail par semaine)	45.6 hours	41.8 hours	38 hours	34.2 hours	30.4 hours
Longévité (espérance de vie à 65 ans)	15.2 années	17.1 années	19 années	20.9 années	22.8 années

TABLE 4.3 – Construction des sociétés fictives

La modélisation des préférences des individus se fait en partie à travers des fonctions de valeur. La construction de ces fonctions peut se faire par exemple en utilisant la méthodologie MACBETH de [Bana e Costa et Vansnick \(1999\)](#) dans le cas d'une agrégation additive, ou son extension proposée dans [Labreuche et Grabisch \(2003\)](#) dans le cas d'une agrégation par l'intégrale de Choquet. Cependant cette tâche de questionnement direct n'est pas compatible avec une expérience telle que nous la mettons en place, et qui est fortement contrainte par le temps. La méthodologie adoptée est décrite par la suite.

Afin de vérifier si un modèle additif existe pour représenter les préférences des individus, nous utilisons une technique inspirée de [Greco et al. \(2008\)](#) (voir aussi la section 2.1). Ainsi, aucune hypothèse ne sera faite a priori pour ce modèle sur la forme de ces fonctions de valeur marginales. Si aucune solution ne peut être trouvée, l'hypothèse d'additivité des préférences peut être rejetée.

Pour l'élicitation des modèles se basant sur l'intégrale de Choquet, l'identification simultanée des poids des coalitions de critères et des fonctions de valeur marginales est problématique, car elle génère des contraintes non linéaires dans les programmes mathématiques. Afin d'éviter cette difficulté, nous décidons de fixer la forme de ces fonctions une fois pour toute. La forme choisie est celles des fonctions "en s" présentées dans l'exemple introductif. Il s'agit de fonctions linéaires par morceaux qui sont convexes en dessous du niveau de référence, et concaves au dessus. Ces fonctions particulières vérifient la première loi de Gossen (i.e. la loi de l'intensité décroissante des besoins) pour des gains par rapport à l'évaluation de référence, et modélisent le fait que de petites pertes par rapport à la société de référence provoquent des pertes de satisfaction décroissantes. Nous sommes conscients que tout en étant plausible, cette hypothèse forte sur les fonctions de valeurs marginales est un point faible de cette étude expérimentale. Cependant, dans une étape de post-analyse, nous perturbons légèrement ces fonctions, afin d'étudier la robustesse de

nos conclusions. Le tableau 4.4 résume ces fonctions de valeurs partielles.

fonction de valeur	“très mauvais”	“mauvais”	‘référence’	“bon”	“très bon”
en s	0	0.125	0.5	0.875	1

TABLE 4.4 – Fonction de valeur “en s”

La détermination des paramètres de l’intégrale de Choquet peut se faire de manière indirecte via une méthode d’identification des capacités. Dans [Grabisch *et al.* \(2008\)](#) les auteurs présentent une grande variété de telles procédures. Nous utilisons dans cette expérience celle qui détermine une capacité minimisant sa variance (voir aussi [Kojadinovic, 2007](#)). Celle-ci exploitera en effet au mieux, en moyenne, tous ses arguments, et choisir cette capacité revient à utiliser l’intégrale de Choquet qui est la “plus proche” de la moyenne arithmétique. Par ailleurs, la fonction objectif dans cette programmation mathématique est strictement convexe, ce qui garantit une solution unique pour la capacité. Nous recherchons également des représentations numériques avec le plus petit niveau de k -additivité ([Grabisch, 1997](#)) afin de d’obtenir la représentation la plus simple, et donc la plus générale possible.

Le questionnaire soumis aux étudiants est structuré comme suit. Tout d’abord les deux premières sections présentent le but de l’étude (i.e. une exploration des préférences des sujets concernant des sociétés), tout en détaillant les tâches à réaliser par les répondants (i.e. classer des sociétés en fonction de leurs préférences). La section 3 du questionnaire présente les sociétés à comparer et les localise par rapport à la société de référence. La section 4 est divisée en 2 parties : tout d’abord, on présente 10 groupes de 6 sociétés à classer au répondant. A l’intérieur d’un groupe, les sociétés sont égales sur deux critères, et varient sur les 3 critères restants. Les améliorations et détériorations marginales par rapport à la société de références sont identiques dans les 10 groupes. Dans la seconde partie, les répondants doivent d’abord recopier les 10 sociétés qu’ils ont préférées dans la première partie, et classer les 5 premières, avant de refaire la même chose avec les 10 sociétés les moins attractives. La dernière section du questionnaire contient quelques questions de contrôle, ainsi qu’une estimation du degré de confiance dans les réponses fournies.

Le groupe de répondants est constitué de 14 étudiants du Master en Sciences Environnementales et de Gestion de l’Université de Liège en Belgique, durant l’année académique 2007-2008.

4.1.2 Résultats

Nous ne présentons ici que les principaux résultats de cette étude exploratoire. Pour plus de détails, nous conseillons au lecteur de se référer à [Meyer et Ponthière \(2011\)](#).

Tout d’abord il faut noter que les modèles de préférences que nous discutons ici sont basés sur les réponses à la seconde partie du questionnaire. La première raison de ce choix est qu’après avoir traité les 10 groupes de sociétés, nous pouvons supposer que les répondants sont familiers avec les données, et que cette connaissance leur permettra de donner un classement assez fiable dans la seconde partie. La deuxième raison concerne les limites des modèles étudiés. En effet, étant donné le grand nombre de sociétés classées par les étudiants dans la première partie, il n’est pas étonnant qu’aucun modèle à base de fonction de valeur additive ou même d’intégrale de Choquet (contrainte par nos fonctions

marginales “en s”) n'existe pour représenter ces préférences. Nous utilisons cependant cette première partie pour valider nos observations.

En premier lieu le modèle additif général est testé pour chacun des répondants. Pour 6 répondants un tel modèle existe, alors que pour les 7 autres, la flexibilité de ce modèle n'est pas suffisante pour représenter leurs préférences. Il est aussi intéressant de noter qu'un répondant a systématiquement violé la dominance de Pareto, ce qui nous a obligé de l'exclure des résultats (répondant 5).

Ce premier résultat nous permet de rejeter le modèle de fonctions de valeur additives pour la représentation générale des préférences des répondants. Dans le cadre particulier défini par cette expérience, les indicateurs classiques (comme par exemple le HDI) ne sont donc pas forcément de bons représentants des préférences des individus.

Nous testons ensuite le modèle à base d'intégrale de Choquet, et nous trouvons un modèle pour chacun des répondants analysés (sous les contraintes que nous avons décrites précédemment). Le niveau de k -additivité minimal est présenté dans le tableau 4.5.

01	02	03	04	06	07	08	09	10	11	12	13	14
2	3	2	3	2	2	4	3	2	2	2	4	2

TABLE 4.5 – Niveaux de k -additivity minimaux pour les 13 répondants analysés.

L'indice d'importance de [Shapley \(1953\)](#) permet de donner une indication de l'importance de chacun des critères. Le tableau 4.6 présente les valeurs de ces indices pour chaque répondant.

rép.	con.	env.	san.	loi.	lon.	rép.	con.	env.	san.	loi.	lon.
01	0.26	0.24	0.15	0.20	0.15	09	0.19	0.24	0.17	0.22	0.18
02	0.21	0.12	0.16	0.31	0.20	10	0.20	0.22	0.21	0.23	0.15
03	0.19	0.19	0.20	0.23	0.19	11	0.23	0.19	0.21	0.20	0.16
04	0.29	0.16	0.26	0.15	0.15	12	0.20	0.16	0.19	0.18	0.27
06	0.20	0.21	0.18	0.24	0.18	13	0.22	0.17	0.15	0.24	0.21
07	0.20	0.22	0.18	0.20	0.20	14	0.22	0.21	0.13	0.20	0.24
08	0.23	0.23	0.18	0.18	0.18						

TABLE 4.6 – Les valeurs des indices de Shapley pour les 13 répondants analysés.

Nous observons tout d'abord que les valeurs prises par cet indice d'importance sont très différentes suivant les critères. Ensuite, le tableau 4.6 nous permet aussi de voir qu'il existe une grande hétérogénéité parmi les répondants au sujet de l'importance qu'ils accordent aux différentes dimensions de bien-être (différence sur l'ordre d'importance et la variance des valeurs). Une nouvelle fois, dans le contexte particulier de cette expérience, cela tend à montrer que les indices utilisant des poids égaux pour les différentes dimensions du bien-être, ne représentent pas fidèlement la variété des préférences des humains. Cette observation est renforcée par le fait que nous estimons ici une capacité minimisant la variance de ses arguments, donc aussi proche que possible de la moyenne arithmétique.

De la même façon, un indice d'interaction ([Murofushi et Soneda, 1993](#)) permet de calculer, pour une capacité donnée, l'interaction entre les différents critères considérés. Le tableau 4.7 présente le nombre d'interactions positives et négatives observées entre les différents critères parmi les 13 répondants analysés.

On peut observer qu'il existe à nouveau une grande hétérogénéité parmi les sujets et il est difficile d'en extraire des tendances générales. Mais assez clairement, de la complémentarité et de la redondance apparaissent entre quasiment toutes les paires de critères.

	complémentarité (> 0)				redondance (< 0)			
	env.	san.	loi.	lon.	env.	san.	loi.	lon.
con.	6	4	6	7	4	8	6	5
env.		6	7	2		6	3	7
san.			6	0			5	10
loi.				6				7

TABLE 4.7 – Nombre d’interactions strictement positives et négatives pour les 13 répondants.

Des résultats plus détaillés sont présentés dans [Meyer et Ponthière \(2011\)](#), où nous discutons plus amplement les similitudes et les différences entre les différents répondants. Il faut noter que nous avons également entrepris une petite étude de robustesse, afin de valider le choix des fonctions de valeur marginales “en s”. Cette étude, détaillée dans [Meyer et Ponthière \(2011\)](#), nous permet de garantir une certaine validité de nos conclusions, en faisant varier sensiblement la forme des fonctions de valeur marginales. Dans ce travail nous avons également appliqué les modèles trouvés à la première partie du questionnaire, pour chacun des individus. L’objectif est de mesurer à quel point les représentations numériques déterminées via la seconde partie du questionnaire permettent de rendre compte des classements exprimés lors de la première phase. Les résultats sont malheureusement un peu décevants, et montrent que pour certains individus, il existe une très grande différence entre le modèle et les préférences exprimées dans la première tâche. Le modèle d’intégrale de Choquet avec les restrictions que nous nous sommes imposées n’est donc pas assez flexible pour modéliser parfaitement nos répondants.

Pour terminer nous avons appliqué les modèles des préférences des individus interrogés au classement de sociétés réelles. Les pays retenus sont l’Autriche, le Danemark, l’Allemagne, la Grèce, l’Italie, les Pays Bas, et l’Espagne, dont les performances sur les 5 critères retenus sont résumées dans le tableau 4.8.

Country	Con	Env	San	Loi	Lon
Autriche	1557	8	15.95	39.9	18.60
Danemark	1428	10	15.70	35.6	17.40
Allemagne	1459	10	15.50	36	18.45
Grèce	1281	9	15.30	43	17.95
Italie	1305	8	15.65	38.80	19.15
Pays Bas	1492	9	16.35	30.80	18.00
Espagne	1255	8	14.15	39.60	19.00

TABLE 4.8 – Performances des 7 pays sur les 5 critères.

Les classements obtenus en appliquant les modèles élicités pour les 13 étudiants sont résumés dans le tableau 4.9.

Pays	01	02	03	04	06	07	08	09	10	11	12	13	14
Autriche	2	2	2	2	2	2	2	2	2	2	2	2	2
Danemark	5	4	4	4	5	5	5	5	3	5	5	5	5
Allemagne	3	3	3	3	3	4	3	4	4	3	4	3	4
Grèce	7	6	7	6	7	7	7	7	7	7	7	7	7
Italie	4	5	5	5	4	3	4	3	5	4	3	4	3
Pays Bas	1	1	1	1	1	1	1	1	1	1	1	1	1
Espagne	6	7	6	7	6	6	6	6	6	6	6	6	6

TABLE 4.9 – Rangs des 7 pays pour chacun des répondants

Comme le montre ce tableau, tous les répondants classent les Pays Bas en premier et l'Autriche en deuxième. Ceci s'explique tout d'abord par le fait que les répondants mettent tous beaucoup d'importance sur les critères de consommation, de loisir et de qualité environnementale, et qu'ensuite il existe beaucoup d'interactions positives entre ces critères. Pour le reste du classement, il existe des désaccords entre les répondants, ce qui s'explique par les différences entre les préférences des individus.

4.1.3 Conclusions et perspectives

Cette petite étude exploratoire nous montre surtout que la mesure de la qualité de la vie est un problème complexe et difficile. Ce travail nous a permis de souligner certains problèmes liés à la subjectivité et la multidimensionalité de la mesure de cette qualité perçue.

La première observation que nous faisons est que, sur base du petit échantillon de répondants, le modèle additif n'a pas l'air d'être un bon modèle pour mesurer la qualité de la vie. Cependant nous ne le rejetons pas catégoriquement, car nous pourrions nous contenter d'un modèle approché, reclassant une partie des sociétés correctement (ce qui nécessiterait moins d'hypothèses sur le modèle sous-jacent que pour l'intégrale de Choquet). Par ailleurs, il est évidemment possible que ce modèle additif puisse être validé à plus grande échelle. Nous observons aussi qu'une des raisons de ce rejet pourrait être l'existence de certaines complémentarités et de redondances entre les indicateurs considérés. Ensuite, nous soulignons qu'il existe une grande hétérogénéité dans les préférences exprimées par les répondants sur ces sociétés hypothétiques. Nous montrons également comment on peut appliquer les modèles des préférences obtenus à des sociétés réelles afin de créer un classement de pays représentant la qualité de vie individuelle.

Cette application n'a évidemment aucune prétention de généralité, mais illustre simplement l'intérêt d'appliquer des outils de MAVT à la problématique des indices de bien-être. En particulier un certain nombre de difficultés restent à résoudre : tout d'abord au sujet de l'aspect arbitraire des fonctions de valeurs "en s" (voir aussi la remarque précédente au sujet du modèle additif), ensuite concernant la possibilité de mener ce type d'expérience à plus grande échelle, et pour terminer concernant l'agrégation de ces préférences en un indice plus général. Nous comptons continuer cette étude exploratoire dans un futur proche afin d'être capable de faire des recommandations plus concluantes.

*

Dans la section suivante, nous présentons une possible application de la méthodologie d'aide au tri pour un groupe de décideurs que nous avons présentée à la section 2.3. Il s'agit d'une réflexion sur la mise en place d'outils pour construire une échelle de risque territoriale en collaboration avec un groupe d'experts. Il est important de noter qu'il ne s'agit que d'une proposition, étant donné que la solution envisagée n'a pas encore été mise en œuvre à ce jour.

4.2 Construction d'une échelle de risque territorial prenant en compte plusieurs critères et experts

Évaluer un niveau de risque associé à une zone géographique peut représenter une tâche complexe. Ce risque concerne souvent des points de vue multiples et conflictuels : une zone

peut être peu risquée sur un critère, tout en étant exposée à un risque critique d'après un autre point de vue. D'autre part, l'affectation d'un niveau de risque à une zone peut exiger l'avis de plusieurs intervenants qui pourraient avoir des objectifs et des systèmes de valeurs différents. En conséquence, la construction de telles échelles de risque nécessite l'agrégation de différents points de vue, tout en tenant compte d'opinions multiples.

Cette activité d'évaluation du risque peut donc être vue comme une tâche classique d'AMCD, ce qui permet de mettre en œuvre une approche formelle pour le problème d'agrégation sous-jacent, tout en modélisant et en incluant les préférences des différents intervenants du processus.

Plus classiquement, l'analyse du risque est traitée dans la littérature depuis longtemps à travers différentes approches méthodologiques. Pour la plupart d'entre elles, le concept de probabilités joue un rôle central pour représenter, analyser et évaluer le niveau de risque. [Aven \(2003, 2008\)](#) fournit une revue détaillée des théories qui permettent d'évaluer la notion de risque. La littérature propose des méthodes comme les réseaux bayésiens, ([Jensen, 1997](#); [Singpurwalla, 2006](#)), la décision dans l'incertain de [von Neumann et Morgenstern \(1944\)](#); [Wakker \(2012\)](#), l'analyse par les arbres de défaillance ([Molak, 1996](#)), etc.

Dans ce travail, que nous avons mené avec Vincent Mousseau⁶, Olivier Cailloux⁷ et Brice Mayag⁸, la perspective de l'analyse du risque que nous adoptons ne nécessite pas forcément de modéliser l'incertitude sous la forme de probabilités, mais est aussi basée sur une évaluation qualitative. Il faut noter que de l'analyse du risque basée sur une méthodologie mêlant des aspects qualitatifs et quantitatifs n'implique pas forcément que le résultat de l'analyse est moins précis que par une approche purement quantitative. En effet, des quantités numériques et des probabilités représentant le risque sont aussi sujettes à de l'imprécision. D'autre part, la modélisation du risque en utilisant des probabilités, est souvent associée à une évaluation objective. Cependant, comme le mentionne [Aven \(2010, page 110\)](#), une mesure objective du risque ne peut en général pas être obtenue. L'approche que nous proposons permet donc d'intégrer la subjectivité des décideurs et experts dans les modèles de risque.

Dans la suite, en section 4.2.1 nous motiverons d'abord de manière plus détaillée l'intérêt d'utiliser des techniques d'AMCD pour la création d'échelles de risque territorial, et plus particulièrement de la technique d'éllicitation que nous avons présentée en section 2.3. Ensuite à la section 4.2.2, nous présenterons un exemple illustratif implémentant notre approche et qui se base sur les outils informatiques présentés au chapitre 3.

4.2.1 Méthodologie multicritère et multi-décideurs pour la construction d'une échelle de risque territorial

La plupart des concepts utilisés dans les problèmes d'évaluation du risque peuvent assez facilement être modélisés dans la problématique du tri multicritère.

Ce qu'on appelle un ensemble de catégories prédéfinies en AMCD correspond à l'échelle de risque par rapport à laquelle les différentes zones doivent être évaluées. Ces catégories peuvent par exemple être {risque élevé, risque moyen, risque bas}.

Les zones géographiques correspondent aux alternatives de décision, alors que les différents points de vue qui interviennent dans l'évaluation du risque correspondent aux critères

6. Vincent Mousseau, École centrale Paris, France

7. Olivier Cailloux, École centrale Paris, France

8. Brice Mayag, Université Paris Dauphine, France

d'évaluation dans un problème de tri multicritère. On peut par exemple citer la présence ou non d'une école dans la zone ou le pourcentage d'habitants vulnérables.

Par conséquent, déterminer une échelle qualitative de risque revient à construire un modèle de tri multicritère.

Chaque zone étant décrite par un vecteur de facteurs de risque associés aux points de vue impliqués dans le problème, l'objectif est donc d'affecter ces zones à un ensemble de catégories représentant les différents niveaux de risque. Le modèle de tri mis en place contient aussi des données plus subjectives représentant les préférences des experts considérés, comme par exemple l'importance relative des différents critères.

Les paramètres préférentiels peuvent être élicités de manière directe, mais ceci est en général assez difficile, car cela nécessite que le décideur ait une connaissance de la procédure d'aggrégation choisie. Nous proposons plutôt de déduire les paramètres préférentiels de manière indirecte, en demandant aux décideurs de proposer des exemples de zones pour chacun des niveaux de risque retenus, i.e. pour chacune des catégories prédéfinies.

Le fait que plusieurs décideurs interviennent dans la création de l'échelle pose évidemment une difficulté supplémentaire. En effet, chacun d'eux peut potentiellement considérer un même facteur comme plus ou moins dangereux, et il n'y a donc pas forcément consensus entre tous ces experts. La méthodologie présentée en section 2.3 permet cependant de bien répondre à ce problème, et aide à éliciter une partie des paramètres préférentiels de manière consensuelle.

Par conséquent, l'approche proposée ici se situe dans le paradigme des relations de surclassement. Cette modélisation se prête bien à notre problème, étant donné que certains critères déduits des points de vue seront évalués sur des échelles qualitatives, et que l'échelle de risque globale est purement ordinale. Pour la construction de l'échelle nous proposons donc d'utiliser la méthode ELECTRE TRI présentée à la section 2.3.1 et qui se base sur les travaux de Bouyssou et Marchant (2007b,c). Pour rappel, elle utilise uniquement la règle d'affectation pessimiste, sans seuils d'indifférence ou de préférence, et avec un veto binaire, qui invalide un surclassement, peu importe la situation de concordance.

Le processus de construction que nous proposons se détaille comme suit :

- Obtenir de chaque expert des zones typiques qui correspondent aux niveaux de risques prédéfinis. Ces zones sont définies par leurs évaluations sur les critères et seront utilisées comme information indirecte pour l'élicitation du modèle d'évaluation du risque ;
- Rechercher un modèle ELECTRE TRI sans vetos représentant les exemples de zones à l'aide de l'algorithme ICL (voir aussi la section 2.3.2) ;
- Si aucun modèle d'ELECTRE TRI sans vetos ne peut être trouvé, rechercher un modèle avec vetos à l'aide de l'algorithme ICLV (voir aussi la section 2.3.2) ;
- Si aucun modèle ne peut être trouvé, rechercher un ensemble maximal d'exemples de zones compatible avec le modèle en utilisant une technique de Mousseau *et al.* (2003, 2006). Les experts doivent dans ce cas individuellement dire s'ils acceptent d'enlever certaines zones ou de modifier leur évaluation de risque, afin de récupérer la compatibilité avec le modèle ELECTRE TRI ;
- Lorsque des limites de catégories consensuelles ont été trouvées, elles doivent être présentées aux experts pour validation. S'ils ne sont pas d'accord, il est possible qu'ils rajoutent des exemples de zones avec leurs niveaux de risque, afin de converger itérativement à des limites de catégories satisfaisantes.
- A tout moment du processus, les experts peuvent aussi spécifier directement des contraintes sur les valeurs des paramètres du modèle (par exemple des valeurs de

veto). Les algorithmes ICL et ICLV sont capables de tenir compte de ce type de contraintes additionnelles.

- A cette étape, les experts sont d'accord sur un ensemble de limites de catégories, et des poids consensuels peuvent être trouvés à l'aide de la méthode proposée par [Damart et al. \(2007\)](#).

La sortie du processus sont des paramètres préférentiels consensuels, partagés par tous les experts : un vecteur de poids des critères, des limites des catégories, et potentiellement des seuils de veto. Ceux-ci peuvent maintenant être utilisés avec le modèle d'ELECTRE TRI particulier afin de déterminer les catégories des zones restantes, i.e. afin de les évaluer sur l'échelle de risque définie.

Une des principales caractéristiques de l'approche proposée est qu'elle s'applique à un groupe d'experts et qu'elle ne suppose pas qu'une partie du modèle de préférence est connu à l'avance.

Dans la section suivante, nous présentons un exemple illustratif, et montrons comment il peut être résolu à l'aide des outils que nous avons introduit au chapitre 3.

4.2.2 Exemple illustratif

Un groupe de 4 experts désire mettre en place une échelle qui permet d'évaluer le risque de zones territoriales autour d'une installation industrielle. Chaque zone à évaluer doit donc être affectée à une des 3 catégories suivantes {risque élevé $\prec$ risque moyen $\prec$ risque bas}. Chacun de ces niveaux de l'échelle est associé avec des mesures de précaution (par exemple, évacuation de la population). La construction de ce type d'échelles de risque liée à des installations industrielles est classique en France dans le cadre du PPRT (Plan de Prévention des Risques Technologiques) ([PPRT, 2011](#)).

Les 4 membres du groupe d'experts considèrent que les six critères suivants doivent être utilisés pour évaluer le risque associé à chaque zone. Chaque échelle d'évaluation des critères est définie de façon à ce qu'une valeur élevée corresponde à une "bonne" évaluation, et donc un risque moindre.

- (pr) Pourcentage de la population qui n'est pas vulnérable (les enfants de moins de 15 ans et les personnes âgées sont considérées comme vulnérables) ;
- (sc) Présence ou absence d'une école dans la zone ; évaluation binaire (pas d'école : 1, école : 0) ;
- (i) Impact sur d'autres installations industrielles pouvant mener à des effets de cascade (échelle ordinale de 5 niveaux) ;
- (bv) Vulnérabilité des bâtiments de la zone (échelle ordinale de 5 niveaux) ;
- (pe) Présence d'atouts techniques ou environnementaux (échelle ordinale de 5 niveaux) ;
- (d) La distance de la zone à l'installation industrielle (échelle ordinale de 3 niveaux ; moins de 500 mètres : 1, entre 500 et 2000 mètres : 2, plus de 2000 mètres : 3).

Nous supposons que chaque expert est capable de fournir 30 exemples de zones avec leurs risques associés. Ces exemples peuvent correspondre à des zones réelles que les experts ont évaluées dans le passé, ou à des zones fictives définies par leurs vecteurs d'évaluation. Une partie de ces exemples est représentée dans le tableau 4.10. Les données utilisées dans cet exemple illustratif sont disponibles dans le workflow *diviz* à l'adresse <http://www.diviz.org/workflow.ressArticle.html>.

À partir de ces exemples de zones, l'algorithme ICL est utilisé pour déterminer les limites des catégories consensuelles, partagées par tous les experts. Le résultat de cette

expert	zone	pr	sc	i	bv	pe	d	catégorie
dm1	Zone08	91	1	4	5	4	1	Low
dm1	Zone11	93	0	3	5	3	2	Medium
dm1	Zone12	43	1	2	1	3	3	High
dm1	Zone13	91	0	5	2	4	3	High
dm2	Zone00	64	1	4	3	3	2	Medium
dm2	Zone01	84	0	3	2	5	3	High
dm2	Zone03	4	1	3	2	1	2	High
dm2	Zone05	14	1	4	5	5	3	Low
dm3	Zone00	64	1	4	3	3	2	Medium
dm3	Zone01	84	0	3	2	5	3	Medium
dm3	Zone06	9	1	2	2	2	1	High
dm3	Zone08	91	1	4	5	4	1	Low
dm4	Zone02	69	1	5	3	1	2	Medium
dm4	Zone05	14	1	4	5	5	3	High
dm4	Zone09	38	1	3	5	4	3	High
dm4	Zone10	36	0	3	2	4	1	High

TABLE 4.10 – Partie des exemples de zones données par les experts.

procédure est présenté dans le tableau 4.11.

frontière	pr	sc	i	bv	pe	d
b_1	41	1	3	3	2	2
b_2	84	1	4	4	4	3

TABLE 4.11 – Limites des catégories inférées. b_1 sépare les catégories “risque élevé” et “risque moyen”, alors que b_2 sépare les catégories “risque moyen” et “risque bas”.

Les valeurs numériques de ces frontières sont facilement interprétables pour les experts. Par exemple, l'évaluation de la frontière b_2 sur le point de vue de vulnérabilité des bâtiments (bv) vaut 4. Ceci signifie qu'une évaluation de 4 pour une zone donnée sur ce critère compte comme un argument en faveur de l'affectation de cette zone à une catégorie meilleure que b_2 , i.e. la catégorie “risque bas”. Si la zone a une évaluation entre 1 et 3, ce point de vue plaide en faveur de l'affectation de la zone à la catégorie “risque moyen”, alors qu'une évaluation de 0 à 2 tend à placer cette zone dans la catégorie “risque élevé”.

Le tableau 4.11 présente donc des frontières de catégories compatibles avec tous les exemples d'affectation des experts. A ce stade, les poids des experts ne sont cependant pas encore consensuels. Le tableau 4.12 présente pour chaque expert un tel jeu de poids qui correspond à ses affectations.

expert	pr	sc	i	bv	pe	d	λ
dm1	0.05	0.204	0.05	0.348	0.254	0.094	0.649
dm2	0.05	0.273	0.05	0.223	0.13	0.273	0.774
dm3	0.237	0.237	0.237	0.144	0.05	0.094	0.572
dm4	0.556	0.05	0.05	0.244	0.05	0.05	0.804

TABLE 4.12 – Des vecteurs de poids pour chaque expert, déterminés par le programme ICL.

Cependant, lors de la validation de ces limites de catégories, le groupe d’expert estime que la valeur de la frontière b_2 pour le critère d’impact sur d’autres installations (i) doit valoir 5, et qu’une évaluation de 1 sur ce critère doit interdire à la zone d’appartenir à la catégorie de “risque bas” (peu importe les autres évaluations). Ce second point se modélise par un veto, qui est ajouté en contrainte au programme ICLV. Ce dernier est utilisé pour déterminer des profils qui sont compatibles avec ces contraintes supplémentaires, i.e. $v_i^2 = 1$ et $g_i(b_2) = 5$. Les nouvelles limites des catégories sont présentées au tableau 4.13, et des jeux de poids compatibles avec les affectations et ces limites sont résumées dans le tableau 4.14.

frontière	pr	sc	i	bv	pe	d
b_1	41	1	3	3	2	2
b_2	82	1	5	4	4	3
v^1	-	-	-	-	-	-
v^2	-	-	1	-	-	-

TABLE 4.13 – Limites des catégories inférées lors du second essai. b_1 sépare les catégories “risque élevé” et “risque moyen”, alors que b_2 sépare les catégories “risque moyen” et “risque bas”.

expert	pr	sc	i	bv	pe	d	λ
dm1	0.05	0.221	0.05	0.365	0.221	0.094	0.633
dm2	0.05	0.512	0.05	0.144	0.05	0.194	0.854
dm3	0.237	0.237	0.237	0.144	0.05	0.094	0.572
dm4	0.244	0.05	0.05	0.556	0.05	0.05	0.804

TABLE 4.14 – Des vecteurs de poids pour chaque expert, déterminés par le programme ICLV.

Une fois que ces limites ont été déterminées et que tous les experts les acceptent, l’approche suggérée par Damart *et al.* (2007) peut être utilisée pour construire de manière itérative des poids consensuels. Nous supposons ici que la sortie de cette procédure génère les poids du tableau 4.15, qui permettent de valider 96 des 120 exemples de zones fournies par les experts (tout en respectant les frontières, ainsi que la valeur de veto du tableau 4.13).

expert	pr	sc	i	bv	pe	d	λ
tous	0.36	0.18	0.18	0.18	0.05	0.05	0.68

TABLE 4.15 – Un vecteur de poids consensuel respectant 96 des exemples d’affectation.

Supposons maintenant que les 4 experts doivent évaluer le risque associé à 6 zones réelles (ZoneA à ZoneF du tableau 4.16). Ils peuvent à cette fin utiliser la règle d’affectation ELECTRE TRI en se basant sur les paramètres préférentiels élicités ci-avant. Les évaluations de ces zones sont données dans la dernière colonne du tableau 4.16.

Avant de terminer, nous montrons comment le processus d’aide à la décision de cet exemple illustratif peut être implémenté à l’aide des outils présentés au chapitre 3, et plus particulièrement diviz. La figure 4.2 représente le workflow construit pour venir en soutien au processus d’élicitation, ainsi qu’à l’affectation des zones réelles. Le workflow contient deux instances du module *ElectreTri1GroupDisaggregationSharedProfiles* (1 et 2) pour l’élicitation des profils consensuels durant les 2 essais de l’exemple. La sortie du

zone	pr	sc	i	bv	pe	d	évaluation du risque par ELECTRE TRI
ZoneA	68	1	3	3	2	2	moyen
ZoneB	50	0	2	1	4	3	haut
ZoneC	85	1	4	4	1	2	bas
ZoneD	10	0	2	3	1	2	haut
ZoneE	20	0	3	2	1	3	haut
ZoneF	15	1	3	5	4	1	haut

TABLE 4.16 – Zones réelles à évaluer en fonction de leur risque, et l'affectation par ELECTRE TRI.

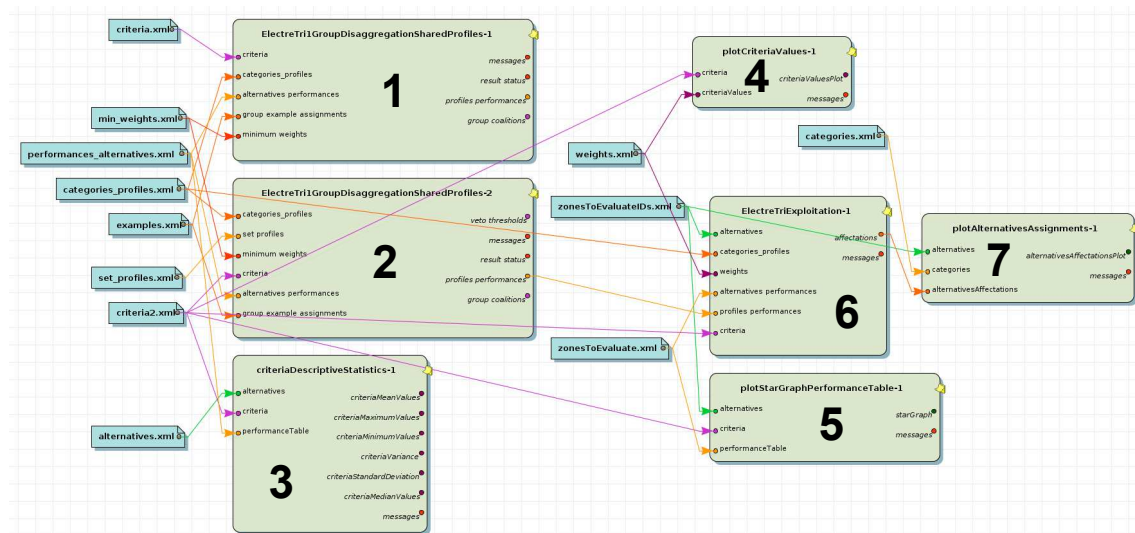


FIGURE 4.2 – Le workflow diviz implémentant le processus de l'exemple illustratif.

premier module n'est pas réutilisée, étant donné que les experts n'étaient pas satisfaits du résultat. Le module *criteriaDescriptiveStatistics* (3) est utilisé par les experts durant la phase d'éllicitation pour comprendre les décisions auxquels ils sont confrontés. Ce module affiche des indicateurs statistiques élémentaires sur les exemples d'affectation (pour chaque critère, moyenne, écart type valeurs maximales et minimales).

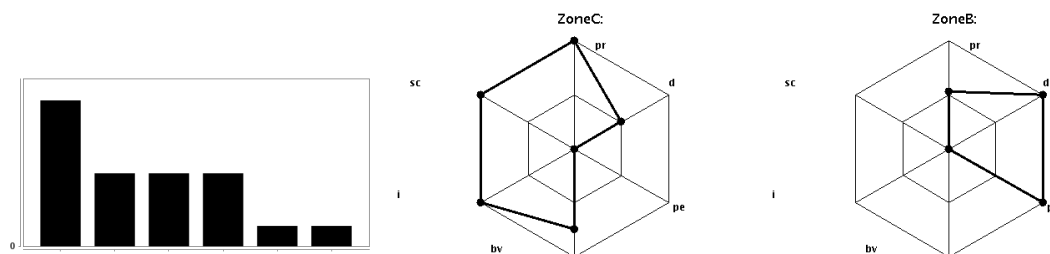


FIGURE 4.3 – Diagramme à barres des poids consensuels et diagramme en étoile de la zone ZoneC (resp. ZoneB), affectée à la catégorie “risque bas” (resp. “risque élevé”) d’après la règle d’affectation ELECTRE TRI.

Une fois que la sortie du second module d’éllicitation (2) est validée, la procédure sug-

gérée par Damart *et al.* (2007) est utilisée (sans diviz) pour obtenir des poids consensuels, dont les valeurs sont représentés dans la partie gauche de la figure 4.3 à l'aide du module *plotCriteriaValues* (4).

Les six zones réelles sont représentées graphiquement à l'aide du module *plotStargraphPerformanceTable* (5). Sa sortie pour ZoneC (resp. ZoneB) est donnée au centre (resp. à droite) de la figure 4.3. La sortie du module d'élicitation, combinée aux poids consensuels, est alors utilisé par le module *ElectreTriExploitation* (6), qui permet d'affecter ces 6 zones réelles à leur niveaux de risque respectifs. Pour terminer, le module *plotAlternativesAssignment* (7) représente graphiquement ces affectations (voir figure 4.4).

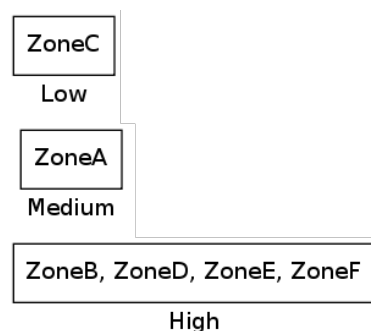


FIGURE 4.4 – Affectation automatique des 6 zones réelles à l'aide du modèle ELECTRE TRI.

Ce workflow peut être téléchargé ensemble avec les données à l'adresse <http://www.diviz.org/workflow.researchArticle.html>. Le lecteur intéressé peut l'importer dans diviz, et ainsi reproduire les calculs de cet exemple illustratif.

4.2.3 Conclusion

Il est évident que la détermination d'une échelle de risque territorial basée sur des points de vues multiples, et prenant en compte les préférences de plusieurs experts, est une tâche difficile. Nous pensons que la méthodologie présentée dans ce travail peut être intéressante dans une application réelle, et peut permettre à un groupe d'experts d'arriver à un consensus.

Dans l'exemple illustratif nous n'avons pas utilisé de probabilités. Ceci ne signifie pas que les critères probabilistes ne peuvent être considérés par notre approche. Nous pensons plutôt qu'un savant mélange de critères quantitatifs, probabilistes ou non, et de critères qualitatifs constitue la bonne démarche pour résoudre ce problème.

Pour terminer nous pensons avoir démontré une nouvelle fois l'intérêt d'un outil intégré comme diviz pour venir en soutien au processus d'aide à la décision proposé ici.

Chapitre 5

Conclusion et recherches futures

A l'issue de ce mémoire, le lecteur devrait avoir perçu les contours du champ d'investigation auquel nous portons notre attention dans nos travaux de recherche.

Nous avons souligné l'importance de fournir des procédures d'agrégation et d'élicitation multicritère dont la validité est minimalement garantie par des confrontations à un grand nombre de problèmes de décision artificiels. Dans certaines situations, ces expérimentations permettent également de produire des recommandations sur utilisation des procédures dans des conditions réelles. Nous avons également présenté un écosystème cohérent d'outils informatiques qui peuvent être utilisés par l'analyste dans le processus d'AMCD. La conséquence principale de cette partie de nos travaux est qu'elle fédère un grand nombre de chercheurs, qui, comme nous, pensent qu'une plateforme libre qui propose une grande variété algorithmes d'AMCD est un atout important pour ce domaine. Nous avons aussi montré l'intérêt de promouvoir l'AMCD au-delà de ses frontières, afin d'intégrer les préférences d'humains dans des domaines comme l'économie ou le contrôle de drones.

Au fur et à mesure de notre présentation, nous avons évoqué des problèmes ouverts et des pistes de recherche qu'il nous paraît intéressant de poursuivre. Ci-après nous détaillons tout d'abord les principales questions de recherche découlant de nos travaux passés.

Recommandations d'utilisation de procédures d'AMCD

Comme nous l'avons montré tout au long du chapitre 2, un point important de notre travail consiste à valider les procédures proposées via une confrontation à des problèmes de décision générés artificiellement. Au-delà de cette validation, ces expérimentations peuvent également servir pour des recommandations d'utilisation des procédures dans des situations de décision réelles.

Ainsi, nous aimerions développer cette idée afin d'en déduire des pistes pour la modélisation du processus d'AMCD, et en particulier, tenter de déterminer, dans le cadre de procédures d'élicitation, une estimation de la quantité d'information à recueillir pour obtenir un modèle satisfaisant des préférences du décideur. Nous avons déjà illustré la faisabilité de cette idée dans la section 2.4, où nous avons déterminé une estimation du nombre de comparaisons par paires à renseigner pour garantir une bonne stabilité de la relation de surclassement.

Étude de l'expressivité de modèles de surclassement

La section 2.1 traite de l'expressivité de modèles d'agrégation dans le contexte de la théorie de la valeur multiattribut. Il serait certainement intéressant de traiter le cas des procédures issues de l'approche du surclassement. La difficulté ici viendra du fait que pour aboutir à une recommandation finale, la relation de surclassement doit en général passer par une phase d'exploitation, pour laquelle la relation entre la sortie et les paramètres préférentiels d'entrée n'est pas toujours facilement interprétable. Du coup, il serait peut-être plus judicieux de partir de relations de surclassement évaluées, et d'analyser différentes façon de les construire, en vue d'en déduire une mesure de l'expressivité de ces modèles.

Prise en compte de décideurs multiples et d'évaluations non consensuelles

Dans la section 2.3 nous avons présenté des travaux en AMCD lorsqu'un groupe de décideurs est impliqué. L'inclusion de ces décideurs multiples dans les modèles d'agrégation n'est pas une chose aisée, mais nous pensons que les besoins liés aux applications réelles doivent nous faire réfléchir sur ces sujets complexes. De manière générale, deux voies peuvent être imaginées pour ces modèles de décision de groupe :

- Obtenir un modèle des préférences pour chaque décideur séparément (par exemple un classement), et ensuite créer le modèle social qui les agrège (par exemple par une règle majoritaire) ;
- Obtenir des paramètres préférentiels consensuels des décideurs, avant de créer un modèle du groupe.

La première voie permet d'utiliser des procédures d'AMCD initialement prévues pour des décideurs uniques. La seconde voie par contre, nécessite, comme nous l'avons montré à la section 2.3, de développer des techniques permettant d'obtenir un consensus sur les paramètres préférentiels. Une difficulté supplémentaire apparaît lorsque les alternatives de décision ne sont pas évaluées de la même façon par tous les décideurs, et qu'il n'y a pas de consensus sur les critères à utiliser. Nous avons récemment rencontré ce type de difficulté dans deux problèmes de décision réels (en architecture de logiciels et dans des problèmes spatiaux), et nous pensons que cette voie nécessite un travail de recherche assez conséquent.

Développement de modèles de workflows diviz plus complexes

Actuellement, les workflows d'algorithmes d'AMCD pouvant être créés avec diviz sont de nature linéaire, et l'exécution est purement séquentielle. Ceci rend difficile l'expression d'algorithmes dynamiques, où l'action d'une période est fonction de l'état de la période précédente. En vue de la valorisation de la plateforme, il est essentiel de pouvoir modéliser des processus de traitement plus complexes, et plus généralement de formaliser la notion de "modèle de calcul", indépendante du graphe de dépendance.

L'inclusion de conditions d'arrêt, de boucles et de tests conditionnels dans ce modèle de workflows permettra ainsi de mieux répondre aux besoins des utilisateurs de diviz, et d'apporter un soutien plus avancé à l'analyste dans le processus d'AMCD.

Validation empirique de procédures d'AMCD dans diviz

Afin de promouvoir la démarche de confrontation de procédures d'AMCD à un gros volume de problèmes de décision artificiels, nous pensons qu'il serait important d'intégrer des outils de validation dans l'environnement diviz. Plus concrètement, il s'agirait tout d'abord de développer des modules permettant de générer des données sous un certain nombre de contraintes, de profiter des modèles de workflows plus complexes afin de créer des exécutions itératives, et de produire des outils d'analyse de résultats permettant d'extraire des indicateurs de performance des modèles.

Développement et mise en place d'une plateforme communautaire de partage de workflows diviz

Le partage de workflows se fait actuellement de manière “manuelle”, en exportant un workflow dans diviz et en l'envoyant à ses collaborateurs ou en le publiant sur un site internet personnel. L'objectif de cette tâche est de créer une plateforme de partage (de type réseau social), où les contributeurs pourraient facilement publier des workflows, les classer par type de contenu, et où ils seraient automatiquement indexés pour des recherches ultérieures. Les utilisations de ce genre d'outil seraient multiples :

- Une plateforme de ce type sera un lieu de stockage de référence, et les chercheurs publiant des expériences utilisant diviz pourront référencer ces workflows dans leurs articles scientifiques.
- Le développement collaboratif de workflows dans un projet de recherche ou pour une application spécifique sera facilité. Une fonctionnalité de gestion de versions permettra d'accéder à tout l'historique de la construction du workflow.
- En enseignement, la plateforme permettra de créer des salles de classes virtuelles, afin que les étudiants d'un même cours puissent travailler collaborativement sur des workflows et que l'enseignant ait accès à toutes les ressources nécessaires pour l'évaluation de leur travail.

Poursuite de l'effort sur la modélisation du processus d'AMCD

Les travaux préliminaires que nous avons présenté à la section 3.4 sur la modélisation du processus d'aide à la décision démontrent la faisabilité du projet, mais soulignent également un certain nombre de difficultés auxquelles nous serons confrontés dans cette tâche.

En particulier, l'inclusion de différentes écoles de pensée et la prise en compte de plusieurs façons différentes d'aborder le problème constituent les deux obstacles majeurs qu'il s'agira de surmonter. Nous restons néanmoins convaincus que ces travaux aboutiront sur un outil méthodologique important qui permettra notamment d'enseigner l'utilisation des techniques d'AMCD de manière plus rigoureuse, et pourra servir de guide à des analystes confrontés à des problèmes de décision réels. Le “bon” alignement des tâches identifiées dans le modèle avec des services d'AMCD constituera un défi supplémentaire qu'il s'agira de relever dans cette recherche.

Étude des indicateurs de bien-être à plus grande échelle

L'étude exploratoire que nous avons présenté à la section 4.1 souligne l'intérêt de tenir compte des préférences individuelles dans la construction d'indicateurs de bien-être.

Cependant, elle fait également émerger un grand nombre de difficultés. Il n'est notamment pas clair si le modèle additif est à rejeter de manière générale ou non, et la construction des fonctions de valeur individuelles dans le cas d'une agrégation par l'intégrale de Choquet nécessite également une réflexion plus poussée.

Nous aimerions donc reproduire une expérience similaire à plus grande échelle. Le questionnaire actuel ayant démontré ses limites, nous pensons qu'il serait intéressant d'inclure une plus grande variété de sociétés fictives, qui ressembleraient plus à des sociétés réelles, sans être identifiables pour autant. De premiers essais montrent qu'il pourrait être intéressant de prendre comme sociétés à classer des représentants de classes de sociétés réelles. Par ailleurs, le nombre de sociétés à classer est un paramètre qu'il serait intéressant d'étudier. Ce dernier point rejoint d'ailleurs la première perspective de recherche présentée ci-avant.

*

Au delà de ces travaux qui représentent la suite logique de nos efforts, notre récente prise de responsabilité à la tête de l'équipe DECIDE de l'UMR 6285 Lab-STICC nous donne l'occasion de présenter deux objectifs de recherche qui ont pour but de fédérer les travaux de l'équipe et d'animer les échanges entre ses chercheurs.

Passage à l'échelle des méthodes d'AMCD

Les principales thématiques de DECIDE sont la fouille de données et l'AMCD. Des tentatives de rapprochement entre ces deux domaines ont déjà porté leurs fruits (voir par exemple [Lenca et al. \(2008\)](#) ou [Bisdorff et al. \(2011\)](#)), mais nous pensons qu'une étape supplémentaire devrait être franchie pour permettre aux algorithmes d'AMCD de traiter des volumes de données plus conséquents.

En particulier, suivant la problématique qui est à résoudre, certaines méthodes de comparaisons par paires ne sont pas adaptées à de grands volumes de données. L'utilisation de méthodes de résolution heuristiques doit probablement être recommandée dans ces situations (comme nous le montrons d'ailleurs dans la section 2.2 sur la classification non supervisée). Un domaine d'application où ces gros volumes de données peuvent poser problème est celui des systèmes d'information géographiques, dans lesquels les alternatives représentent des zones géographiques, qui peuvent être très nombreuses.

Les travaux qui traitent de ce passage à l'échelle en AMCD n'en sont qu'à leurs débuts. Nous pensons qu'au sein de l'équipe DECIDE nous pourrions apporter des réponses à ces questions difficiles, en fédérant les compétences de chercheurs des deux disciplines.

Observation de décideurs dans des situations réelles

L'UMR Lab-STICC dispose d'un laboratoire d'observation des usages des technologies de l'information et de la communication (LOUSTIC)¹. Cette plateforme permet notamment d'étudier des testeurs confrontés à des dispositifs comme des logiciels ou d'autres services innovants.

Nous pensons que ce dispositif peut être utilisé pour de l'analyse expérimentale du comportement décisionnel. L'étude de testeurs confrontés à des situations de décision simulées

1. <http://www.loustic.net/brest>

doit permettre d'alimenter les réflexions autour des procédures d'agrégation et d'élicitation des préférences. Par ailleurs, les travaux sur la modélisation du processus d'aide à la décision, pourront certainement être alimentés par les résultats de ces expériences. Des réflexions dans ce sens ont été entrepris dans le cadre de la thèse de doctorat de [Deparis \(2012\)](#).

Certains membres de DECIDE se sont déjà déclarés intéressés par ce projet, et nous pensons que les résultats qui pourraient découler de ces expérimentations pourront également alimenter des débats intéressants dans la communauté d'AMCD.

*

De manière générale, nos recherches futures s'inscrivent dans le mouvement, fédérateur, d'*opérationnalisation* de l'AMCD. Comme nous l'avons déjà évoqué, nous pensons en effet que nos travaux contribueront à l'amélioration des outils de soutien (au sens large) du processus d'AMCD. En amont de l'utilisation de ces procédures, nous nous intéresserons à leur validité, et à produire des recommandations sur leur mise en œuvre en pratique. Plus en aval, nous continuerons à produire des outils efficaces et pérennes qui participent à la diffusion de l'AMCD et permettent de résoudre plus facilement des problèmes d'aide à la décision. Il est intéressant de noter que, dans un contexte où les volumes de données à traiter sont en forte croissance, cette opérationnalisation s'accompagne évidemment du développement de méthodes de résolution et d'élicitation approchées, sujet qui nous tient plus particulièrement à cœur pour le moment.

Bibliographie

- M. R. Anderberg : *Cluster Analysis for Applications*. Academic Press, 1973. (page 13).
- K. Arrow : *Social choice and individual values*. J. Wiley, New York, 1951. 2nd edition, 1963. (page 2).
- T. Aven : *Foundations of Risk Analysis*. Wiley, 2003. (page 77).
- T. Aven : *Risk Analysis : assessing uncertainties beyond expected values and probabilities*. Wiley, 2008. (page 77).
- T. Aven : *Misconceptions of Risk*. Wiley, 2010. (page 77).
- C. Bana e Costa, L. Ensslin, Émerson C. Cornêa et J.-C. Vansnick : Decision support systems in action : Integrated application in a multicriteria decision aid process. *European Journal of Operational Research*, 113(2):315 – 335, 1999. ISSN 0377-2217. (page 56).
- C. Bana e Costa et J. Vansnick : Preference relations in MCDM. Dans T. Gal et T. H. T. Steward, éd. : *MultiCriteria Decision Making : Advances in MCDM models, algorithms, theory and applications*. Kluwer, 1999. (page 72).
- R. Baroudi et N. Safia : Towards multicriteria analysis : A new clustering approach. Dans *Proceedings of the 2010 International Conference on Machine and Web Intelligence*, p. 126–131, 2010. (page 13).
- V. Belton et T. Stewart : *Multiple Criteria Decision Analysis : An Integrated Approach*. Kluwer Academic Publishers, 2001. (page 56).
- S. Bigaret, V. Chiprianov, P. Meyer et J. Simonin : Towards a formalisation of the MCDA process. Dans *75th meeting of the European Working Group in Multiple Criteria Decision Aid*, 2012. (page 57).
- S. Bigaret et P. Meyer : ws-RXMCD, 2009-2010. <http://github.com/paterijk/ws-RXMCD>. (page 44).
- R. Bisdorff : Electre-like clustering from a pairwise fuzzy proximity index. *European Journal of Operational Research*, 138(2):320–331, 2002a. (page 13).
- R. Bisdorff : Logical foundation of multicriteria preference aggregation. Dans D. B. et al., éd. : *Essay in Aiding Decisions with Multiple Criteria*, p. 379–403. Kluwer Academic Publishers, 2002b. (page 29).
- R. Bisdorff : Concordant outranking with multiple criteria of ordinal significance. *4OR, Quarterly Journal of the Belgian, French and Italian Operations Research Societies, Springer-Verlag*, 2 :4:293–308, 2004. (page 30).
- R. Bisdorff, P. Meyer et A.-L. Olteanu : A clustering approach using weighted similarity majority margins. Dans J. Tang, I. King, L. Chen et J. Wang, éd. : *Proceedings of the 7th International Conference on Advanced Data Mining and Applications*, vol. 7120 de *Lecture Notes in Computer Science*, p. 15–28. Springer, 2011. ISBN 978-3-642-25852-7. (pages xviii, 16, 20, 88).

- R. Bisdorff, P. Meyer et M. Roubens : Rubis : a bipolar-valued outranking method for the best choice decision problem. *4OR, Quarterly Journal of the Belgian, French and Italian Operations Research Societies*, 6:143–165, 2008. (pages 4, 29).
- R. Bisdorff : The Python digraphs module for Rubis, 2007. <http://ernst-schroeder.uni.lu/Digraph/doc/>. (page 44).
- R. Bisdorff, P. Meyer et T. Veneziano : Inverse analysis from a Condorcet robustness denotation of valued outranking relations. *Lecture notes in computer science*, 5783:180 – 191, 2009. Algorithmic decision theory. (page 8).
- R. Bisdorff, P. Meyer et T. Veneziano : Elicitation of criteria weights maximising the stability of pairwise outranking statements. *Journal of Multi-Criteria Decision Analysis*, 2013. accepted, <http://public.telecom-bretagne.eu/~pmeyer/articles/pdf/bisdorffMeyerVeneziano2013.pdf>. (pages 8, 30, 31, 32).
- G. Bous, P. Fortemps, F. Glineur et M. Pirlot : A novel method for eliciting additive value functions on the basis of holistic preference statements. *European Journal of Operational Research*, 206:435–444, 2010. (page 51).
- D. Bouyssou, T. Marchant, M. Pirlot, P. Perny, A. Tsoukias et P. Vincke : *Evaluation and decision models, A critical Perspective*. Kluwer’s International Series. Kluwer, Massachusetts, 2000. (pages 17, 40, 51, 61).
- D. Bouyssou, T. Marchant, M. Pirlot, A. Tsoukias et P. Vincke : *Evaluation and decision models with multiple criteria, Stepping stones for the analyst*. Springer’s International Series. Springer, New York, 2006. (pages 41, 62).
- D. Bouyssou et M. Pirlot : Preferences for multi-attributed alternatives : Traces, dominance, and numerical representations. *J. of Mathematical Psychology*, 48:167–185, 2004. (pages 8, 9).
- D. Bouyssou et T. Marchant : An axiomatic approach to noncompensatory sorting methods in mcdm, i : The case of two categories. *European Journal of Operational Research*, 178(1):217–245, April 2007a. (page 5).
- D. Bouyssou et T. Marchant : An axiomatic approach to noncompensatory sorting methods in MCDM, I : The case of two categories. *European Journal of Operational Research*, 178(1):217–245, 2007b. ISSN 0377-2217. (pages 20, 21, 78).
- D. Bouyssou et T. Marchant : An axiomatic approach to noncompensatory sorting methods in MCDM, II : more than two categories. *European Journal of Operational Research*, 178(1):246–276, 2007c. ISSN 0377-2217. (pages 20, 21, 78).
- J.-P. Brans et P. Vincke : A preference ranking organization method. *Management Science*, 31(6):647–656, 1985. (pages 45, 51).
- J. Brans et B. Mareschal : *PROMETHEE-GAIA. Une Méthodologie d’Aide à la Décision en Présence de Critères Multiples*. Ellipses, Paris, France, 2002. (page 4).
- C. Bron et J. Kerbosch : Algorithm 457 : finding all cliques of an undirected graph. *Communications of the ACM*, 16(9):575–577, 1973. (page 16).
- R. Brown : Toward a prescriptive science and technology of decision aiding. *Annals of Operations Research*, 19:465–483, 1989. ISSN 0254-5330. URL <http://dx.doi.org/10.1007/BF02283535>. (page 56).
- O. Cailloux, C. Lamboray et P. Nemery : A taxonomy of clustering procedures. *Dans Proceedings of the 66th Meeting of the European Working Group on MCDA*, 2007. (page 13).
- O. Cailloux : J-MCDA, 2010. <http://sourceforge.net/projects/j-mcda/>. (pages 44,

- 45).
- O. Cailloux, B. Mayag, P. Meyer et V. Mousseau : Operational tools to build a multicriteria territorial risk scale with multiple experts. *Reliability Engineering & System Safety*, (120):88–97, 2013. (page 67).
- O. Cailloux, P. Meyer et V. Mousseau : Eliciting ELECTRE TRI category limits for a group of decision makers. *European Journal of Operational Research*, 233(1):133 – 140, November 2012. (pages 8, 21, 22, 23, 26, 27).
- V. Chiprianov, P. Meyer et J. Simonin : Towards a model-based multiple criteria decision aid process. Technical report, Telecom Bretagne, 2012. URL <http://public.telecom-bretagne.eu/~pmeyer/files/ProcessMdeMcdav02.pdf>. (pages 57, 59).
- G. Choquet : Theory of capacities. *Annales de l'Institut Fourier*, 5:131–295, 1953. (page 69).
- S. Damart, L. Dias et V. Mousseau : Supporting groups in sorting decisions : Methodology and use of a multi-criteria aggregation/disaggregation DSS. *Decision Support Systems*, 43(4):1464–1475, août 2007. (pages 22, 26, 79, 81, 83).
- P. Dasgupta : *An Inquiry into Well-Being and Destitution*. Oxford University Press, 1993. (page 68).
- Y. De Smet et S. Eppe : Relational multicriteria clustering : The case of binary outranking matrices. Dans M. Ehrgott, éd. : *Proceedings of the 5th International Conference on Evolutionary Multi-Criterion Optimization*, vol. 5467 de *Lecture Notes in Computer Science*, p. 380–392. Springer Berlin, 2009. (page 13).
- Y. De Smet et L. Guzman : Towards multicriteria clustering : an extension of the k-means algorithm. *European Journal of Operational Research*, 158(2):390–398, 2004. (page 13).
- Decision Deck Consortium : The Decision Deck project. <http://www.decision-deck.org/>, 2009a. (page 38).
- Decision Deck Consortium : Strategic manifesto of the Decision Deck project, 2009b. <http://www.decision-deck.org/manifesto.html>. (page 38).
- S. Deparis : Etude de l'effet du conflit multicritère sur l'expression des préférences : Une approche empirique, 2012. (page 89).
- L. C. Dias et J. Clímaco : On computing ELECTRE's credibility indices under partial information. *Journal of Multi-Criteria Decision Analysis*, 8(2):74–92, 1999. (page 22).
- L. Dias et J. Clímaco : ELECTRE TRI for groups with imprecise information on parameter values. *Group Decision and Negotiation*, 9(5):355–377, sept. 2000. (page 22).
- L. Dias, V. Mousseau, J. Figueira et J. Clímaco : An aggregation/disaggregation approach to obtain robust conclusions with ELECTRE TRI. *European Journal of Operational Research*, 138(2):332–348, avr. 2002. (page 22).
- L. C. Dias et V. Mousseau : Inferring electre's veto-related parameters from outranking examples. *European Journal of Operational Research*, 170(1):172–191, April 2006. (page 28).
- B. S. Everitt, S. Landau et M. Leese : *Cluster Analysis*. Hodder Arnold, 2001. ISBN 9780340761199. (page 13).
- E. Fernandez, J. Navarro et S. Bernal : Handling multicriteria preferences in cluster analysis. *European Journal of Operational Research*, 202(3):819–827, 2010. ISSN 0377-2217. (page 13).
- J. Figueira, V. Mousseau et B. Roy : ELECTRE methods. Dans J. Figueira, S. Greco et M. Ehrgott, édés : *Multiple Criteria Decision Analysis : State of the Art Surveys*, p.

- 133–162. Springer Verlag, Boston, Dordrecht, London, 2005. Chapter 4. (page 20).
- J. R. Figueira, Y. De Smet et J.-P. Brans : MCDA methods for sorting and clustering problems : PROMETHEE TRI and PROMETHEE CLUSTER. Rap. tech. TR/SMG/2004-002, SMG, Université Libre de Bruxelles, 2004. (page 13).
- J. Figueira, S. Greco, B. Roy et R. Słowiński : ELECTRE methods : Main features and recent developments. *Handbook of Multicriteria Analysis*, p. 51–89, 2010. (page 29).
- P. Fishburn : *Utility Theory for Decision Making*. Wiley, New York, 1970. (pages 3, 4).
- L. A. Franco et G. Montibeller : Facilitated modelling in operational research. *European Journal of Operational Research*, 205:489–500, 2010. (page 5).
- F. Glover : Tabu search - part I. *INFORMS Journal on Computing*, 1(3):190–206, 1989. (page 17).
- F. Glover : Tabu search - part II. *INFORMS Journal on Computing*, 2(1):4–32, 1990. (page 17).
- M. Grabisch : The application of fuzzy integrals in multicriteria decision making. *European Journal of Operational Research*, 89:445–456, 1996. (pages 5, 8).
- M. Grabisch : k -order additive discrete fuzzy measure and their representation. *Fuzzy Sets and Systems*, 92:131–295, 1997. (page 73).
- M. Grabisch, I. Kojadinovic et P. Meyer : A review of capacity identification methods for Choquet integral based multi-attribute utility theory ; Applications of the Kappalab R package. *European Journal of Operational Research*, 186(2):766–785, 2008. (pages 10, 11, 45, 69, 73).
- M. Grabisch et C. Labreuche : Fuzzy measures and integrals in MCDA. Dans J. Figueira, S. Greco et M. Ehrgott, édés : *Multiple Criteria Decision Analysis*, p. 563–608. Springer, 2004. (pages 8, 12).
- S. Greco, B. Matarazzo et R. Słowiński : Rough sets methodology for sorting problems in presence of multiple attributes and criteria. *European Journal of Operational Research*, 138(2):247–259, 2002. (page 20).
- S. Greco, V. Mousseau et R. Słowiński : Ordinal regression revisited : multiple criteria ranking using a set of additive value functions. *European Journal of Operational Research*, 191(2):415–435, December 2008. (pages 10, 28, 72).
- S. Greco, M. Kadzinski, V. Mousseau et R. Słowiński : ELECTRE GKMS : Robust ordinal regression for outranking methods. *European Journal of Operational Research*, 214(1):118–135, 2011. (page 28).
- M. Hall, E. Frank, G. Holmes, B. Pfahringer, P. Reutemann et I. H. Witten : The WEKA data mining software : an update. *SIGKDD Explorations*, 11(1):10–18, 2009. (page 6).
- E. Jacquet-Lagrèze et Y. Siskos : Assessing a set of additive utility functions for multicriteria decision making : the UTA method. *European Journal of Operational Research*, 10:151–164, 1982. (pages 10, 28, 42, 45, 51, 63).
- F. Jensen : *Introduction to Bayesian Networks*. Springer, 1997. (page 77).
- R. L. Keeney et H. Raiffa : *Decision with multiple objectives*. Wiley, New-York, 1976. (pages 3, 8, 69).
- S. Kirkpatrick, C. D. Gelatt et M. P. Vecchi : Optimization by simulated annealing. *Science*, 220(4598):671–680, 1983. (page 17).
- I. Koch : Enumerating all connected maximal common subgraphs in two graphs. *Theoretical Computer Science*, 250(1–2):1–30, 2001. (page 16).

- I. Kojadinovic : Minimum variance capacity identification. *European Journal of Operational Research*, 177(1):498–514, 2007. (page 73).
- D. Krantz, R. Luce, P. Suppes et A. Tversky : *Foundations of measurement, volume 1 : Additive and polynomial representations*. Academic Press, 1971. (page 5).
- C. Labreuche et M. Grabisch : The Choquet integral for the aggregation of interval scales in multicriteria decision making. *Fuzzy Sets and Systems*, 137:11–16, 2003. (page 72).
- Y. Le Bras, P. Meyer, S. Lallich et P. Lenca : A robustness measure of association rules. Dans Springer, éd. : *The European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases*, vol. 6322, Lecture Notes in Computer Science, p. 227 – 242, 2010a. (page xviii).
- Y. Le Bras, P. Meyer, P. Lenca et S. Lallich : Mesure de la robustesse de règles d’association. Dans *QDC 2010 : atelier Qualité des Données et des Connaissances, en conjonction avec Extraction et gestion des connaissances*, p. 27 – 38, 2010b. (page xviii).
- Y. Le Bras, P. Meyer, P. Lenca et S. Lallich : Mesure formelle de la robustesse des règles d’association. *Revue des nouvelles technologies de l’information (RNTI)*, E-22:71 – 88, 2011. (page xviii).
- B. Leistedt : UTAR, 2010. <http://github.com/ixkael/RMCDA>. (page 45).
- P. Lenca, P. Meyer, B. Vaillant et S. Lallich : On selecting interestingness measures for association rules : user oriented description and multiple criteria decision aid. *European journal of operational research*, 184(2):610 – 626, january 2008. (page 88).
- R. Luce et H. Raiffa : *Games and Decisions*. J. Wiley, New York, 1957. (page 3).
- P. Meyer : Progressive methods in multiple criteria decision analysis, 2007. Thèse de doctorat. (page xvii).
- P. Meyer et S. Bigaret : diviz : a software for modeling, processing and sharing algorithmic workflows in MCDA. *Intelligent Decision Technologies*, 6(4):283–296, 2012. doi :10.3233/IDT-2012-0144. (pages 38, 55).
- P. Meyer et S. Bigaret : XMCD A : a standard XML encoding of MCDA data. Dans R. Bisdorff, L. Dias, P. Meyer, V. Mousseau et M. Pirlot, eds : *Evaluation and Decision Models with Multiple Criteria : Case Studies*. Springer, 2013. accepted. (pages 37, 39, 40).
- P. Meyer et A. Olteanu : Formalizing and solving the problem of clustering in MCDA. *European Journal of Operational Research*, 227(3):494–502, 2013. (pages 8, 13, 15, 17).
- P. Meyer et M. Pirlot : On the expressiveness of the additive value function and the Choquet integral models. Dans *From Multiple Criteria Decision Aid to Preference Learning (DA2PL’2012)*, p. 48–56. University of Mons, 15-16 Novembre 2012. (pages 8, 9, 11, 12).
- P. Meyer et G. Ponthière : Eliciting preferences on multiattribute societies with a Choquet integral. *Computational economics*, 37(2):133 – 168, 2011. (pages 67, 73, 75).
- G. Miller : The magical number seven, plus or minus two : Some limits on our capacity for processing information. *The Psychological Review*, 63:81–97, 1956. (page 71).
- V. Molak : *Fundamentals of Risk Analysis and Risk Management*. Lewis Publishers, 1996. (page 77).
- J. Moon et L. Moser : On cliques in graphs. *Israel Journal of Mathematics*, 3(1):23–28, 1965. (page 16).
- V. Mousseau : Elicitation des préférences pour l’aide multicritère à la décision. Mémoire

- présenté en vue de l'obtention de l'habilitation à diriger des recherches, Université Paris-Dauphine, 2003. (page 4).
- V. Mousseau, L. Dias et J. Figueira : Dealing with inconsistent judgments in multiple criteria sorting models. *4OR*, 4(3):145–158, 2006. (pages 22, 78).
- V. Mousseau, L. Dias, J. Figueira, C. Gomes et J. Clímaco : Resolving inconsistencies among constraints on the parameters of an MCDA model. *European Journal of Operational Research*, 147(1):72–93, 2003. (pages 22, 78).
- V. Mousseau, J. Figueira et J. Naux : Using assignment examples to infer weights for ELECTRE TRI method : Some experimental results. *European Journal of Operational Research*, 130(2):263–275, avr. 2001. (pages 22, 28).
- V. Mousseau et R. Słowiński : Inferring an ELECTRE TRI model from assignment examples. *Journal of Global Optimization*, 12(2):157–174, 1998. (pages 22, 28).
- V. Mousseau, R. Słowiński et P. Zielniewicz : A user-oriented implementation of the ELECTRE TRI method integrating preference elicitation support. *Computers & Operations Research*, 27(7-8):757–777, 2000. (page 20).
- T. Murofushi et S. Soneda : Techniques for reading fuzzy measures (iii) : Interaction index. *Dans 9th Fuzzy System Symposium*, p. 693–696, Sapporo, Japan, 1993. (page 74).
- P. Narayan, P. Meyer et D. Campbell : Embedding Human Expert Cognition into Autonomous UAS Trajectory Planning. *IEEE Journal of Systems, Man, and Cybernetics, Part B : Cybernetics*, 43(2):530–543, 2013. doi :10.1109/TSMCB.2012.2211349. (page 67).
- P. Nemery et Y. De Smet : Multicriteria ordered clustering. Rap. tech. TR/SMG/2005-003, Université Libre de Bruxelles/SMG, 2005. (page 13).
- OMG : Software & Systems Process Engineering Meta-Model Specification (SPEM) Version 2.0, 2008. (page 57).
- L. Osberg et A. Sharpe : An index of economic well-being for selected OECD countries. *Review of Income and Wealth*, 48(3):291–316, 2002. (page 68).
- L. Osberg et A. Sharpe : How should we measure the economic aspects of well-being? *Review of Income and Wealth*, 51(2):311–336, 2005. (page 68).
- P. Perny : Multicriteria filtering methods based on Concordance/Non-Discordance principles. *Annals of Operations Research*, 80:137–167, 1998. (page 20).
- M. Pidd : From problem-structuring to implementation. *The Journal of the Operational Research Society*, 39(2):pp. 115–121, 1988. ISSN 01605682. (page 56).
- M. Pirlot, H. Schmitz et P. Meyer : An empirical comparison of the expressiveness of the additive value function and the Choquet integral models for representing rankings. *Dans 25th Mini-EURO Conference Uncertainty and Robustness in Planning and Decision Making (URPDM 2010)*. INESC Coimbra, 2010. (pages 8, 9).
- PPRT : Le plan de prévention des risques technologiques (PPRT), Guide méthodologique. <http://www.developpement-durable.gouv.fr/Maitrise-de-l-urbanisation-PPRT,12775.html>, 2011. Ministère de l'écologie, du développement et de l'aménagement durables. (page 79).
- Quantum GIS Development Team : *Quantum GIS Geographic Information System*. Open Source Geospatial Foundation, 2009. URL <http://qgis.osgeo.org>. (page 47).
- R Development Core Team : *R : A language and environment for statistical computing*. R Foundation for Statistical Computing, Vienna, Austria, 2005. URL <http://www.R-project.org>. ISBN 3-900051-00-3. (pages 6, 10).

- R. Ramsin et R. F. Paige : Process-centered review of object oriented software development methodologies. *ACM Comput. Surv.*, 40:3 :1–3 :89, 2008. ISSN 0360-0300. URL <http://doi.acm.org.gate6.inist.fr/10.1145/1322432.1322435>. (page 57).
- G. Retali, P. Meyer et M. L. Goff-Pronost : Implementation of remotely monitored medical dialysis units : dealing with multiple criteria and multiple decision makers. *Dans 10th Annual International Conference on Health Economics, Management and Policy*, 2011a. (page 68).
- G. Retali, P. Meyer et M. L. Goff-Pronost : Taking into account of patients and physicians preferences in the implementation of remotely monitored medical dialysis units. *Dans 8th World Congress International Health Economics Association : Transforming Health & Economics*, 2011b. (page 68).
- C. Rocha, L. C. Dias et I. Dimas : Multicriteria classification with unknown categories : A clustering—sorting approach and an application to conflict management. *Journal of Multi-Criteria Decision Analysis*, 2012. ISSN 1099-1360. (page 13).
- J. Rothenberg : The nature of modeling. *Dans* L. William, K. Loparo et N. Nelson, éd. : *Artificial Intelligence, Simulation, and Modeling*, p. 75–92. New York, John Wiley and Sons, Inc., 1989. (page 57).
- B. Roy : Classement et choix en présence de points de vue multiples (la méthode ELECTRE). *Revue française d'informatique et de recherche operationelle (RIRO)*, 2:57–75, 1968. (pages 3, 45).
- B. Roy : ELECTRE III : Un algorithme de classement fondé sur une représentation floue des préférences en présence de critères multiples. *Cahiers du CERO*, 20(1):3–24, 1978. (page 17).
- B. Roy : *Méthodologie multicritère d'aide à la décision*. Ed. Economica, collection Gestion, 1985. (page 3).
- B. Roy : The outranking approach and the foundations of ELECTRE methods. *Theory and Decision*, 31:49–73, 1991. (page 20).
- B. Roy et D. Bouyssou : *Aide multicritère à la décision : Méthodes et cas*. Economica, Paris, 1993. (pages 3, 4, 29).
- B. Roy et D. Vanderpooten : The European school of MCDA : Emergence, basic features and current works. *Journal of Multi-Criteria Decision Analysis*, 5:22–38, 1996. (page 4).
- P. Samuelson : A note on the pure theory of consumer's behavior. *Economica*, 5(17):61 – 71, 1938. (page 2).
- D. Schmeidler : Integral representation without additivity. *Proceedings of the American Mathematical Society*, 97:255–261, 1986. (page 8).
- L. S. Shapley : A value for n -person games. *Dans Contributions to the theory of games, vol. 2*, Annals of Mathematics Studies, no. 28, p. 307–317. Princeton University Press, Princeton, N. J., 1953. (page 74).
- H. Simon : A behavioural model of rational choice. *Dans* H. Simon, éd. : *Models of man : social and rational ; mathematical essays on rational human behavior in a social setting*, p. 241–260. J. Wiley, New York, 1957. (page 3).
- N. Singpurwalla : *Reliability and risk. A bayesian perspective*. Willey, 2006. (page 77).
- O. Sobrie : Intégration d'outils d'aide à la décision dans un système d'information géographique. Travail de Fin d'Études présenté en vue de l'obtention du grade de Master Ingénieur civil en Informatique et Gestion, 2011. (page 47).

- A. N. The et V. Mousseau : Using assignment examples to infer category limits for the ELECTRE TRI method. *JMCDA*, 11(1):29–43, nov. 2002. (page 22).
- E. Triantaphyllou et A. Sanchez : Sensitivity analysis approach for some deterministic multi-criteria decision-making methods. *Decision Sciences*, 28(1):151–194, 1997. (page 28).
- A. Tsoukiàs : On the concept of decision aiding process : an operational perspective. *Annals OR*, 154(1):3–27, 2007. (pages 5, 56, 57).
- United Nations Development Program (UNDP) : *Human Development Report*. Oxford University Press, New York, 1990. (page 68).
- T. Veneziano : ws-PyXMCD, 2010. <http://github.com/quiewbee/ws-PyXMCD>. (page 45).
- P. Vincke : *L'aide multicritère à la décision*. Ellipses, Paris, 1989. (pages 3, 13).
- P. Vincke : *Multicriteria Decision-aid*. Wiley, 1992. (page 17).
- J. von Neumann et O. Morgenstern : *Theory of games and economic behavior*. Princeton University Press, Princeton, 1944. Second edition in 1947, third in 1954. (pages 2, 77).
- D. von Winterfeldt et W. Edwards : *Decision Analysis and Behavioral Research*. Cambridge University Press, Cambridge, 1986. (page 5).
- P. Wakker : *Prospect Theory for Risk and Ambiguity*. Cambridge University Press, Cambridge, UK, 2012. (page 77).
- P. Wakker : *Additive Representations of Preferences : A new Foundation of Decision Analysis*. Kluwer Academic Publishers, Dordrecht, Boston, London, 1989. (page 9).

Annexe A : Curriculum vitæ et activités de recherche

Sommaire

A.1 Curriculum vitæ	99
A.1.1 État civil	99
A.1.2 Formation	100
A.1.3 Fonctions assurées	100
A.1.4 Compétences	101
A.2 Thèmes de recherche	101
A.2.1 Contributions au processus d'AMCD	101
A.2.2 Fouille de données	102
A.3 Responsabilités scientifiques et animation de la recherche	103
A.3.1 Responsabilités	103
A.3.2 Encadrement de thèses de doctorat	103
A.3.3 Encadrement de stages de Master Recherche	104
A.4 Collaborations récentes nationales et internationales	104
A.5 Participation à des projets	104
A.6 Affiliations	105
A.7 Travaux de rédacteur	105
A.8 Participation à des comités de lecture	106
A.9 Membre de comités de programme de conférences	106
A.10 Membre de comités d'organisation de conférences scientifiques	106
A.11 Séjours de recherche	107

Cette première annexe contient mon curriculum vitæ ainsi qu'un résumé de mes principales activités scientifiques, depuis mes débuts en recherche en 1999.

A.1 Curriculum vitæ

A.1.1 État civil

Patrick Meyer

Institut TELECOM, TELECOM Bretagne

Département LUSSI

UMR CNRS 6285 Lab-STICC

Technopôle Brest-Iroise CS 83818

29238 Brest Cedex 3, France

patrick.meyer@telecom-bretagne.eu

<http://public.telecom-bretagne.eu/~pmeyer/>

Date et lieu de naissance : 2 mai 1976, Luxembourg.

A.1.2 Formation

- ◊ *Octobre 2007 : Doctorat en mathématiques appliquées et en sciences de l'ingénieur*
Université du Luxembourg & Faculté Polytechnique de Mons (co-tutelle)
 - *Titre de la thèse* : Progressive Methods in Multiple Criteria Decision Analysis.
 - *Soutenance* : le 8 octobre 2007 à l'Université du Luxembourg devant le jury composé de
 - Raymond Bisdorff, Professeur, Université du Luxembourg (directeur de thèse) ;
 - Marc Pirlot, Professeur, Faculté Polytechnique de Mons (directeur de thèse) ;
 - Philippe Fortemps, Professeur, Faculté Polytechnique de Mons (membre du comité d'encadrement) ;
 - Jean-Luc Marichal, Assistant Professeur, Université du Luxembourg (président du jury et membre du comité d'encadrement) ;
 - Vincent Mousseau, Professeur, Université Paris Dauphine ;
 - Theodor Stewart, Professeur, University of Cape Town.
 - *Mention* : excellent
- ◊ *Septembre 2003 : DEA (Master recherche) en mathématiques*
Faculté Polytechnique de Mons
 - *Titre du mémoire* : Two Aspects of User Centred Data Analysis and Decision Aid : Quality Measures for Association Rules and Multiple Criteria Sorting.
 - *Encadrant* : Marc Roubens, Université de Liège (Belgique).
- ◊ *Juillet 1999 : Licence (Bac +4) en mathématiques*
Université de Liège
 - *Titre du mémoire* : Problème de tournées de distribution avec fenêtres de temps et application.
 - *Encadrant* : Marc Roubens, Université de Liège (Belgique).

A.1.3 Fonctions assurées

- ◊ *Depuis mars 2008 : Maître de conférences*
TELECOM Bretagne, Département LUSSI
 - Recherche en Aide Multicritère à la Décision (AMCD) et en fouille de données (voir section A.2.1).
 - Enseignements en AMCD, en algorithmique et en recherche opérationnelle.
- ◊ *Septembre 2004 - février 2008 : Assistant - doctorant*
Université du Luxembourg, Service de Mathématiques Appliquées (SMA)
 - Recherche en AMCD (voir Section A.2.1).
 - Enseignements en mathématiques, statistique et informatique.
 - Organisation de plusieurs conférences et de séminaires (voir section A.10).

- Développement et maintenance de sites web (SMA, conférences).
- ◊ *Juillet 1999 - décembre 2003 : Chercheur “Région Wallonne”, Belgique*
Université de Liège
 - Participation au projet CELOFA (Centre d’Étude de la Logique Floue et de ses Applications) : développement d’un logiciel de gestion de portefeuilles financiers à l’aide de systèmes d’inférence floue (voir section A.5).
 - Supervision de travaux d’étudiants et de mémoires de *licence*.
 - Enseignements ponctuels de cours en recherche opérationnelle et en AMCD.

A.1.4 Compétences

- ◊ **Systèmes d’exploitation et langages de programmation**
 - Linux (Ubuntu, Debian), Windows (2000, XP).
 - R, HTML, $\text{\LaTeX}$, GNU MathProg.
- ◊ **Langues**
 - Luxembourgeois, français, hongrois : langues maternelles.
 - Allemand, anglais : courant.

A.2 Thèmes de recherche

J’ai mené mes activités de recherche de juillet 1999 à décembre 2003 au sein de l’*Institut de Mathématiques* de l’Université de Liège, sous la direction du Prof. Marc Roubens. Ensuite, de septembre 2004 à février 2008, j’ai effectué ces activités dans le cadre d’une thèse de doctorat sous la direction du Prof. Raymond Bisdorff au *Service de Mathématiques Appliquées* de l’Université du Luxembourg. Depuis mars 2008, je mène ces travaux au département LUSI de Télécom Bretagne dans le laboratoire Lab-STICC (UMR CNRS 6285).

Elles portent principalement sur deux thèmes :

- Contributions au processus d’AMCD ;
- Fouille de données.

A.2.1 Contributions au processus d’AMCD

Comme je l’ai montré dans ce mémoire, mes contributions à ce thème se situent dans les 2 principaux courants méthodologiques d’AMCD : l’école européenne autour des méthodes de surclassement et l’école américaine autour de la théorie de la valeur multiattribut (MAVT).

Dans le cadre de nombreuses collaborations nationales et internationales (voir section A.4), mes contributions principales à ce thème sont :

- l’extension de l’*intégrale de Choquet* à des *nombres flous* ;
- la proposition d’une *méthode d’identification de la capacité* modélisant l’importance des sous-ensembles de critères dans le cadre de problèmes de MAVT se basant sur l’intégrale de Choquet ;
- la proposition de méthodes de *tri*, de *classement* et de *choix* en MAVT se basant sur l’intégrale de Choquet ;

- l'étude de l'*expressivité* de modèles d'agrégation (voir section 2.1) ;
- la proposition d'une méthode *progressive* de résolution de la *problématique du choix* via une relation de *surclassement bipolaire* évaluée ;
- l'extension de la problématique du choix à celle du *k-choix* (détermination de $k > 1$ alternatives *les meilleures*) dans un contexte de surclassement ;
- des travaux en *désagrégation* de relations de surclassement et l'identification des paramètres préférentiels d'un décideur (analyse inverse) ;
- l'élicitation de *préférences de groupe* pour la méthode de tri multicritère Electre Tri (voir section 2.3) ;
- l'étude de la *stabilité* d'une relation de surclassement par rapport à des variations des poids, et l'élicitation stable de ces poids (voir section 2.4) ;
- l'étude du concept de *progressivité* en AMCD ;
- la définition de la problématique de la classification non supervisée en AMCD (*clustering*) et la proposition de techniques de résolution (voir section 2.2) ;
- la participation au développement d'une boîte à outils, nommée *Kappalab*, permettant l'utilisation de mesures et d'intégrales non additives dans le cadre de problèmes d'AMCD ;
- la proposition du standard de données XMCDa pour stocker des informations issues du domaine de l'AMCD (<http://www.decision-deck.org/xmcd>) dans le cadre du projet Decision Deck (voir section 3.1) ;
- le développement de plusieurs *services-web* implémentant des algorithmes d'AMCD dans le cadre du projet Decision Deck (voir section 3.2) ;
- le développement de l'outil *diviz* qui permet de créer et de partager des méthodes d'AMCD (<http://www.diviz.org>) dans le cadre du projet Decision Deck (voir section 3.3) ;
- le développement d'un modèle du *processus* d'AMCD à l'aide de techniques d'ingénierie dirigée par les modèles (voir section 3.4) ;
- une analyse multicritère des comportements d'*achat* et de *vente* des *analystes financiers* de la salle des marchés de CBC Banque & Assurance, Bruxelles, Belgique ;
- l'utilisation de l'intégrale de Choquet et de théories de l'AMCD dans la perception du *bien-être* en Economie (voir section 4.1) ;
- une proposition de méthodologie pour la construction d'*échelles de risque* territorial (voir section 4.2) ;
- l'utilisation de l'AMCD pour la mise en place de *systèmes de télémédecine* en prenant en compte les préférences des patients et des médecins ;
- l'inclusion de modèles de préférences multicritère dans le contrôle automatique de *drones*.

A.2.2 Fouille de données

- la définition du concept de *robustesse* de règles d'association ;
- une analyse multicritère des *indices de qualité de règles d'association* en fouille de données ;
- l'analyse de données financières accessibles librement sur internet et la confirmation empirique de plusieurs hypothèses, projet MIKADO (Market Investigation and

Knowledge Acquisition through Data Observation) ;

A.3 Responsabilités scientifiques et animation de la recherche

A.3.1 Responsabilités

- ◊ **Co-responsable** de l'équipe DECIDE de l'UMR CNRS 6285 Lab-STICC (Laboratoire des Sciences et Techniques de l'Information, de la Communication et de la Connaissance) (avec Philippe Lenca) ;

Le projet scientifique du Lab-STICC peut se résumer facilement dans le titre “des capteurs à la connaissance : communiquer et décider”. Le laboratoire est constitué de 3 pôles qui permettent de décliner de façon concrète les différents points de vue de ce projet. Le pôle CID traite plus particulièrement de la connaissance, de l'information et de la décision, et à l'intérieur de ce pôle, l'ambition de l'équipe DECIDE (DECision aId and knowleDge discovEry) est de développer des systèmes d'aide à la décision fiables, de qualité et robustes intégrant notamment les différents acteurs du processus de décision. L'équipe s'intéresse principalement aux aspects méthodologiques, algorithmiques, et de robustesse des modèles en aide à la décision et en fouille de données. À ce jour elle est composée de 27 membres, dont 11 doctorants et post-doctorants.

Mon objectif en tant que responsable de cette équipe est de stimuler les interactions entre les différents chercheurs qui la composent, et surtout de tenter de créer un pont entre les différentes disciplines qui y sont représentées : la fouille de données, l'AMCD, les grands systèmes et l'ingénierie des modèles.

- ◊ **Responsable** du séminaire mensuel de l'équipe DECIDE ;

Le séminaire de l'équipe DECIDE a lieu au moins une fois par mois, et permet aux chercheurs de présenter leurs travaux récents, et de discuter librement de nouvelles pistes de recherche. Il permet également à des personnes extérieures à l'équipe de s'y exprimer, en vue de monter de nouvelles collaborations.

Son programme est disponible à sur le site web du Lab-STICC <http://labsticc.fr/>.

- ◊ Participation au groupe de travail sur l'élaboration de la **stratégie de la recherche** de Télécom Bretagne (2011 - 2012).

A.3.2 Encadrement de thèses de doctorat

- ◊ **Thomas Veneziano**, *On the stability of outranking relations : Theoretical and practical aspects*, thèse en co-tutelle entre l'Université du Luxembourg (UL) et Télécom Bretagne (TB), soutenue le 10 septembre 2012, directeurs de thèse : Raymond Bisdorff (UL) et Philippe Lenca (TB) ;
- ◊ **Alexandru Olteanu**, *On clustering in multi-criteria decision aid : Theory and applications*, thèse en co-tutelle entre l'Université du Luxembourg (UL) et Télécom Bretagne (TB), en cours, soutenance prévue le 24 juin 2013, directeurs de thèse : Raymond Bisdorff (UL) et Philippe Lenca (TB).

A.3.3 Encadrement de stages de Master Recherche

- **Sarra Kacem** (Université de Rennes 1 - INSA de Rennes, Master 2 en statistique et économétrie), *Elicitation de préférences en aide multicritère à la décision : application au cas de sociétés multidimensionnelles par l'approche Macbeth* (2009-2010)

A.4 Collaborations récentes nationales et internationales

- ◊ **Vincent Mousseau**, Professeur, Ecole Centrale Paris
XMCD, Elicitation de préférences pour une méthode de tri multicritère
- ◊ **Marc Pirlot**, Professeur, Université de Mons, Belgique
Expressivité des modèles d'agrégation
- ◊ **Raymond Bisdorff**, Professeur, Université du Luxembourg
Développement de la méthode RUBIS dans le cadre de la problématique du choix en AMCD, désagrégation de relations de surclassement bipolaires valuées, stabilité des relations de surclassement, XMCD
- ◊ **Philippe Lenca**, Maître de Conférences, Télécom Bretagne, Brest
Fouille de données, AMCD appliquée aux indices de qualité de règles d'association, robustesse de règles d'association
- ◊ **Pritesh Narayan**, chercheur, University of the West of England, Grande Bretagne, et **Duncan Campbell**, Professeur, Queensland University of Technology, Australie
Inclusion de préférences d'experts dans le contrôle automatique de drones
- ◊ **Grégory Ponthière**, Maître de Conférences, Ecole Normale Supérieure, Paris
Elicitation de préférences à partir de sociétés multicritères
- ◊ **Sébastien Bigaret**, ingénieur de recherche, Télécom Bretagne, Brest
Développement de l'outil diviz et de services-web XMCD

A.5 Participation à des projets

- ◊ **Decision Deck** (2007 -)
Luxembourg - France - Belgique - Portugal - Pologne - Pays Bas
 - *Responsable* : Decision Deck Consortium.
 - *Membres* : Chercheurs européens de l'Université du Luxembourg, de l'Université Paris Dauphine, de l'Université de Mons, de l'Ecole Centrale Paris, de Télécom Bretagne, ...
 - *Objectifs* : Développement d'outils d'AMCD.
 - *Contributions personnelles* : élaboration et maintenance du standard de données XMCD, développement de services-web XMCD, participation au développement de l'outil diviz de création, d'exécution et de partage de workflows d'AMCD.
- ◊ **D2-Decision Deck** (2007 - 2009)
Luxembourg - France - Belgique

- *Financement* : Université du Luxembourg, Projet num. R1F207L07.
- *Responsable* : Raymond Bisdorff.
- *Membres* : Raymond Bisdorff, Patrick Meyer, Jean-Luc Marichal.
- *Objectifs* : Développement de systèmes d'aide à la décision pour l'entreprise : élicitation de préférences, analyse de robustesse et outils logiciels.
- *Contributions personnelles* : participation à l'élaboration du standard de données **XMCD**A, développement de services-web issus de la librairie Kappalab.

- ◊ **Recherches méthodologiques et mathématiques en aide à la décision et classification (2005 - 2007)**
Luxembourg
 - *Financement* : Université du Luxembourg, Projet num R1F205L01.
 - *Responsables* : Raymond Bisdorff et Jean-Luc Marichal.
 - *Membres* : Raymond Bisdorff, Jean-Luc Marichal, Patrick Meyer.
 - *Objectifs* : Etudes méthodologiques et mathématiques en AMCD.

- ◊ **Kappalab (2006 - 2007)**
France - Luxembourg
 - *Financement* : CNRS (GDR Recherche Opérationnelle num 3002).
 - *Responsable* : Ivan Kojadinovic.
 - *Membres* : Michel Grabisch, Christophe Labreuche, Fabien Lange, Jean-Luc Marichal, Patrick Meyer.
 - *Objectifs* : Amélioration du package Kappalab pour la manipulation de capacités et d'intégrales discrètes.

- ◊ **CELOFA (1999 - 2003)**
Belgique
 - *Financement* : Région Wallonne (Belgique) (num 9813803).
 - *Responsables* : Jacques Teghem et Marc Roubens.
 - *Objectifs* : Développement d'un logiciel de gestion de portefeuilles à l'aide de systèmes d'inférence floue.
 - *Partenaire industriel* : CBC Banque & Assurance, Bruxelles, Belgique.

A.6 Affiliations

- ◊ Membre de l'UMR CNRS 6285 Lab-STICC (2008-);
- ◊ Membre et trésorier du *Decision Deck Consortium*, organe de gestion du projet *Decision Deck* (2007-);
- ◊ Membre de l'*EURO Working Group "Multiple Criteria Decision Aiding"*.

A.7 Travaux de rédacteur

- ◊ **Membre du comité de rédaction** de l'International Journal of Multicriteria Decision Making (IJMCDM)

- ◊ **Rédacteur invité** du numéro spécial "EURO 2004" du European Journal of Operational Research (EJOR) *Avec Raymond Bisdorff et Yannis Siskos*

A.8 Participation à des comités de lecture

◇ Relectures pour des journaux internationaux

- *Fuzzy Sets and Systems*;
- *European Journal of Operational Research*;
- *A Quarterly Journal of Operations Research (4 OR)*;
- *EURO Journal on Decision Processes*;
- *Journal of Multicriteria Decision Analysis*;
- *International Journal of Multicriteria Decision Making*.

◇ Relectures pour des revues nationales

- *Revue Nationale des Technologies de l'Information*.

◇ Relectures pour des conférences internationales

- *Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU)*;
- *The International Conference on Data Mining (DMIN)*;
- *Knowledge Discovery and Business Intelligence workshop (KDBI)*.

A.9 Membre de comités de programme de conférences

- ◇ *From Decision Aid to Preference Learning workshop (DA2PL'12)*;
- ◇ *Decision Deck workshop (bi-annuels)*;
- ◇ *Ateliers Qualité des Données et des Connaissances* (en conjonction avec EGC), France;
- ◇ *Quality issues, measures of interestingness and evaluation of data mining models workshop (QIMIE)*;
- ◇ *Knowledge Discovery and Business Intelligence workshop (KDBI)* (Portuguese Conference on Artificial Intelligence, EPIA);
- ◇ *International Conference on Data Mining (DMIN'10)*;
- ◇ 4èmes journées d'Extraction et de Gestion des Connaissances (EGC 2004).

A.10 Membre de comités d'organisation de conférences scientifiques

- ◇ Organisateur du *5th Decision Deck Workshop* (17-18 sep. 2008, Brest)
- ◇ Organisateur adjoint du *38th Meeting of the European Mathematical Psychology Group* (10-13 sep. 2007, Luxembourg)
- ◇ *1st Decision Deck Workshop* (8-9 mar. 2007, Luxembourg)
- ◇ *21st annual meeting of the Belgian Operational Research Society* (18-19 jan. 2007, Luxembourg)
- ◇ *54e et 61e réunions du groupe de travail EURO MCDA* (oct. 2001, Durbuy, Belgique et 10-11 mar. 2005, Luxembourg)

A.11 Séjours de recherche

- ◇ **Université du Luxembourg**, Interdisciplinary Lab for Intelligent and Adaptive Systems
Luxembourg (*plusieurs séjours d'une semaine, 2008-2013*)
Collaborations avec *Raymond Bisdorff* et *Jean-Luc Marichal*, encadrement des thèses de doctorat de *Thomas Veneziano* et d'*Alexandru Olteanu*.
- ◇ **Tampere University of Technology**, Department of Mathematics
Finlande (*24 avr. 2006 - 29 avr. 2006*)
Collaboration avec *Stephan Foldes* et séminaire sur la méthode Rubis.
- ◇ **University of Economics, Prague**, Department of Information and Knowledge Engineering
République Tchèque (*9 mar. 2003 - 15 mar. 2003*)
Collaboration avec *Jan Rauch*, travaux sur les règles d'association et la méthode Guha / 4ft-miner (Cost Action 274, TARSKI).
- ◇ **ENST Bretagne, Brest**, Département IASC
France (*1 semaine en nov. 2002*)
Collaboration avec *Philippe Lenca* et travaux sur les mesures de qualité de règles d'association.
- ◇ **ENST Bretagne, Brest**, Département IASC
France (*1 semaine en juin 2002*)
Collaboration avec *Philippe Lenca* et séminaires sur : TOMASO, la qualité de règles de décision, introduction aux mesures de qualité de règles d'association.

Annexe B : Publications

Sommaire

B.1	Articles dans des revues internationales avec comité de lecture	109
B.2	Articles dans des ouvrages collectifs avec comité de lecture . .	110
B.3	Articles dans des revues françaises avec comité de lecture . . .	111
B.4	Rédacteur invité	111
B.5	Articles dans des actes de conférences avec comité de lecture .	111
B.6	Articles dans des bulletins de sociétés scientifiques	113
B.7	Communications à des séminaires et des conférences sans actes	113

Dans cette deuxième annexe, je présente la liste de mes publications scientifiques, ainsi que mes participations à divers colloques ou conférences. Suivant leur type, ces travaux de dissémination sont regroupés en diverses catégories afin de permettre au lecteur de plus facilement s’orienter dans cette liste.

B.1 Articles dans des revues internationales avec comité de lecture

1. O. Cailloux, B. Mayag, P. Meyer, V. Mousseau, Operational tools to build a multicriteria territorial risk scale with multiple stakeholders, *Reliability Engineering & System Safety*, 120, 88-97, 2013, Elsevier (doi :10.1016/j.ress.2013.06.004).
2. R. Bisdorff, P. Meyer, T. Veneziano, Elicitation of criteria weights maximising the stability of pairwise outranking statements, *Journal of Multi-Criteria Decision Analysis*, 2013, Wiley, available online (doi :10.1016/j.ejor.2013.01.016).
3. P. Meyer, A. Olteanu, Formalizing and solving the problem of clustering in MCDA, *European Journal of Operational Research*, 227 (3), 494 - 502, 2013, Elsevier, (10.1016/j.ejor.2013.01.016).
4. P. Narayan, P. Meyer, D. Campbell, Embedding Human Expert Cognition into Autonomous UAS Trajectory Planning, *IEEE Journal of Systems, Man, and Cybernetics, Part B : Cybernetics*, 43 (2), 530-543, 2013, (doi : 10.1109/TSMCB.2012.2211349).
5. P. Meyer, S. Bigaret, diviz : a software for modeling, processing and sharing algorithmic workflows in MCDA, *Intelligent Decision Technologies*, 6 (4), 283-296, 2012, IOS Press (doi :10.3233/IDT-2012-0144).
6. O. Cailloux, P. Meyer, V. Mousseau, Eliciting Electre Tri category limits for a group of decision makers, *European Journal of Operational Research*, 223 (1), 133 - 140, 2012, Elsevier, (doi :10.1016/j.ejor.2012.05.032).

7. P. Meyer, G. Ponthière, Eliciting preferences on multiattribute societies with a Choquet integral, *Computational Economics*, 37(2), 133-168, 2011, Springer, (doi :10.1007/s10614-009-9196-0).
8. P. Meyer, Progressive Methods in Multiple Criteria Decision Analysis, *4OR*, 7 (2), 191-194, June 2009, Springer (doi :10.1007/s10288-008-0069-5).
9. R. Bisdorff, P. Meyer, M. Roubens, Rubis : a bipolar-valued outranking method for the best choice decision problem, *4OR*, 6 (2), June 2008, Springer, (doi :10.1007/s10288-007-0045-5).
10. M. Grabisch, I. Kojadinovic, P. Meyer, A review of capacity identification methods for Choquet integral based multi-attribute utility theory, Applications of the Kappalab R package, *European Journal of Operational Research*, 186 (2), 766-785, 2008, Elsevier, (doi :10.1016/j.ejor.2007.02.025).
11. P. Lenca, P. Meyer, B. Vaillant, S. Lallich, Useful characterization and multicriteria decision aid : A two step approach to interestingness measure selection, *European Journal of Operational Research*, 184 (2), 610-626, 2008, Elsevier, (doi :10.1016/j.ejor.2006.10.059).
12. P. Meyer, M. Roubens, On the use of the Choquet integral with fuzzy numbers in Multiple Criteria Decision Aiding, *Fuzzy Sets and Systems*, 157 (7), 927-938, 2006, Elsevier, (doi :10.1016/j.fss.2005.11.014).
13. J.-L. Marichal, P. Meyer, and M. Roubens, Sorting multiattribute alternatives : The TOMASO method, *Computers & Operations Research*, 32 (4), 861-877, 2005, Elsevier, (doi :10.1016/j.cor.2003.09.002).

B.2 Articles dans des ouvrages collectifs avec comité de lecture

1. P. Meyer, S. Bigaret, XMCDA : a standard XML encoding of MCDA data. In R. Bisdorff, L. Dias, V. Mousseau, M. Pirlot, editors, *Evaluation and Decision Models with Multiple Criteria : Case Studies*. Springer Handbook, forthcoming 2013.
2. P. Lenca, B. Vaillant, P. Meyer and S. Lallich, Association rule interestingness measures : experimental and theoretical studies. In F. Guillet and H.J. Hamilton, editors, *Quality Measures in Data Mining*, Chapter 3, 51-76, 2007, Springer, ISBN 3540449116.
3. P. Meyer, Use of an ordinal sorting technique (TOMASO) in stock selection. In J.-P. Barthélemy and P. Lenca, editors, *Advances in MultiCriteria Decision Aid*, 2005, ENST Bretagne, ISBN 2-9523875-0-8.
4. P. Meyer, M. Roubens, Choice, Ranking and Sorting in Fuzzy Multiple Criteria Decision Aid. In J. Figueira, S. Greco, and M. Ehrgott, editors, *Multiple Criteria Decision Analysis : State of the Art Surveys*, 471-506, Springer Verlag, Boston, Dordrecht, London, 2005, Springer, ISBN 038723067X.

B.3 Articles dans des revues françaises avec comité de lecture

1. Y. Le Bras, P. Meyer, P. Lenca, S. Lallich, Mesure formelle de la robustesse des règles d'association, *Revue des nouvelles technologies de l'information*, vol. E-22, pp. 71-88, 2011.
2. B. Vaillant, P. Meyer, E. Prudhomme, S. Lallich, P. Lenca, S. Bigaret, Mesurer l'intérêt des règles d'association, *Revue des Nouvelles Technologies de l'Information* (special issue : "Extraction et gestion des connaissances : Etat et perspectives", invited editors : F. Cloppet, J.-M. Petit, N. Vincent), RNTI-E5, 421-426, 2006, Cépaduès.
3. P. Lenca, P. Meyer, B. Vaillant, P. Picouet, S. Lallich, Evaluation et analyse multicritère des mesures de qualité des règles d'association, *Revue des Nouvelles Technologies de l'Information* (special issue "Mesures de qualité pour la fouille de données", invited editors : H. Briand, M. Sebag, R. Gras, F. Guillet), RNTI-E1, 219-246, 2004, Cépaduès.
4. P. Lenca, P. Meyer, P. Picouet, B. Vaillant, S. Lallich, Critères d'évaluation des mesures de qualité en ECD, *Revue des Nouvelles Technologies de l'Information* (special issue : Entreposage et fouille de données, invited editors : O. Boussaid and S. Lallich), no 1, 123-134, 2003, Cépaduès.

B.4 Rédacteur invité

1. R. Bisdorff, P. Meyer and Y. Siskos, editors, OR and the management of electronic services, Feature issue EURO'2004, *European Journal of Operational Research*, 17 articles, in press, Elsevier.

B.5 Articles dans des actes de conférences avec comité de lecture

1. P. Meyer, A. Olteanu, Preferentially ordered clustering. MDAI 2012 : 9th Modeling Decisions for Artificial Intelligence Conference, 21-23 november 2012, Girona, Spain, 2012, vol. Modeling Decisions for Artificial Intelligence, pp. 87-98, ISBN 978-84-695-6132-4.
2. P. Meyer, M. Pirlot, On the expressiveness of the additive value function and the Choquet integral models, *DA2PL'2012 : from Multiple Criteria Decision Aid to Preference Learning*, 15-16 november 2012, Mons, Belgium, 2012, pp. 48-56.
3. R. Bisdorff, P. Meyer, A.L. Olteanu, A clustering approach using weighted similarity majority margins. *ADMA 2011 : 7th International Conference on Advanced Data Mining and Applications*, Heidelberg : Springer, LNAI 7120, 17-19 december 2011, Beijing, China, 2011, vol. 7120, pp. 15-28, ISBN 978-3-642-25852-7.
4. R. Bisdorff, P. Meyer, A.L. Olteanu, A clustering approach using weighted similarity majority margins. *The 23rd Benelux Conference on Artificial Intelligence*, 3-4 november 2011, Ghent, Belgium, 2011, pp. 15-28 (doi :10.1007/978-3-642-25853-4_2).

5. Y. Le Bras, P. Meyer, S. Lallich, P. Lenca, A robustness measure of association rules. *The European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases (ECML-PKDD)*, Springer, 20-24 september 2010, Barcelona, Spain, 2010, vol. 6322, Lecture Notes in Computer Science, pp. 227-242, ISBN 978-3-642-15882-7.
6. S. Bigaret, P. Meyer, Diviz : an MCDA workflow design, execution and sharing tool. *25th Mini-EURO Conference Uncertainty and Robustness in Planning and Decision Making (URPDM 2010)*, Coimbra : 15-17 april 2010, Coimbra, Portugal, 2010, ISBN 978-989-95055-3-7.
7. M. Pirlot, H. Schmitz, P. Meyer, An empirical comparison of the expressiveness of the additive value function and the Choquet integral models for representing rankings. *25th Mini-EURO Conference Uncertainty and Robustness in Planning and Decision Making (URPDM 2010)*, 15-17 april 2010, Coimbra, Portugal, 2010, ISBN 978-989-95055-3-7.
8. T. Veneziano, P. Meyer, R. Bisdorff, Analyse inverse robuste à partir d'informations préférentielles partielles. *ROADEF 2010 : 11ème congrès annuel de la Société française de recherche opérationnelle et d'aide à la décision*, 24-26 février 2010, Toulouse, France, 2010, pp. 75-88, ISBN 2-7238-0414-3.
9. Y. Le Bras, P. Meyer, P. Lenca, S. Lallich, Mesure de la robustesse de règles d'association, *QDC 2010 : atelier Qualité des Données et des Connaissances*, in conjunction with Extraction et gestion des connaissances, 26 January 2010, Hammamet, Tunisia, 27-38, 2010.
10. R. Bisdorff, P. Meyer, T. Veneziano, Inverse analysis from a Condorcet robustness denotation of valued outranking relations. In F. Rossi and A. Tsoukiás (Eds.), *Algorithmic Decision Theory*, Lecture Notes in Computer Science 5783, Springer Verlag Berlin Heidelberg, 180 - 191, (doi :10.1007/978-3-642-04428-1_16).
11. P. Meyer, J.-L. Marichal, R. Bisdorff, Disaggregation of bipolar-valued outranking relations, In L.T. Hoai An, P. Bouvry, P.D. Tao, editors, *Modelling, Computation and Optimization in Information Systems and Management Sciences*, Communications in Computer and Information Science, Springer, Proc. of MCO'08 conference, Metz, France, 8-10 September 2008, 204 - 213, 2008 (doi : 10.1007/978-3-540-87477-5).
12. P. Meyer, R. Bisdorff, Exploitation of a bipolar-valued outranking relation for the choice of k best alternatives, In proc. of *FRANCORO V / ROADEF 2007 : Conférence scientifique conjointe en Recherche Opérationnelle et Aide à la Décision*, Grenoble, France, 20 - 23 February 2007, 193 - 206, 2007.
13. M. Grabisch, I. Kojadinovic, P. Meyer, Using the Kappalab R package for capacity identification in Choquet integral based MAUT, In proc. of the *11th International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU)*, Paris, France, 2-7 July 2006, 1702 - 1709, 2006.
14. R. Bisdorff, P. Meyer, M. Roubens, The RuBy method for the recommendation of a best choice from a bipolar outranking relation, In proc. of the *27th Linz Seminar on Fuzzy Set Theory, Preferences, Games and Decisions*, Linz, Austria, 7-11 February 2006, 40-43, 2006.
15. P. Meyer, M. Roubens, On the use of the Choquet Integral with fuzzy numbers in Multiple Criteria Decision Aiding, In proc. of *EUSFLAT-LFA 2005*, Barcelona, Spain, 6-9 September 2005, 210 - 215, 2005.

16. B. Vaillant, P. Meyer, E. Prudhomme, S. Lallich, P. Lenca, S. Bigaret, Mesurer l'intérêt des règles d'association, In proc. of the *Atelier Qualité des Données et des Connaissances (associated to the conference Extraction et Gestion des Connaissances (EGC) 2005)*, Paris, France, 18 January 2005, 69 - 78, 2005.
17. P. Meyer, M. Roubens, Ordinal Sorting in the Presence of Interacting Points of View : TOMASO. In proc. of the *10th International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU)*, Perugia, Italy, 4-9 July 2004, 217-224, 2004.
18. P. Meyer, M. Roubens, Ordinal Sorting in the Presence of Interacting Points of View :TOMASO. In proc. of the *25th Linz Seminar on Fuzzy Set Theory, Mathematics of Fuzzy Systems*, Linz, February 3-7 2004, 144-152, 2004.
19. P. Lenca, P. Meyer, P. Picouet, B. Vaillant, Aide multicritère à la décision pour évaluer les indices de qualité des connaissances. In proc. of *Extraction et Gestion des Connaissances (EGC) 2003*, Lyon, France, 22-24 January 2003, RSTI-RIA 17/2003, 271-282, 2003.
20. F. Boniver, P. Meyer, MIKADO (Market Investigation and Knowledge Acquisition through Data Observation) : First results and perspectives, In proc. of *Human Centered Processes (HCP) 2003* (R. Bisdorff, editor), Luxembourg, Luxembourg, 5-7 May 2003, 181-185, 2003.
21. P. Lenca, P. Meyer, B. Vaillant, P. Picouet, S. Lallich, Critères d'évaluation des mesures de qualité en ECD, In proc. of *XXXV Journées de Statistique*, Lyon, France, 2 - 6 June 2003, 647-650, 2003.

B.6 Articles dans des bulletins de sociétés scientifiques

1. S. Bigaret, P. Meyer, diviz : an MCDA workflow design, execution and sharing tool, *Newsletter of the EURO Working Group Multicriteria Aid for Decisions*, Series 3, Number 21, 10-13, spring 2010.
2. M. Grabisch, I. Kojadinovic, P. Meyer, 'Kappalab', a GNU R package for capacities and non-additive integral manipulation, *EUSFLAT Newsletter*, Vol. 3, No. 1, April, 2007.
3. M. Grabisch, I. Kojadinovic, P. Meyer, Overview of 'Kappalab', a toolbox for capacities and non-additive integral manipulation, *Newsletter of the EURO Working Group Multicriteria Aid for Decisions*, Series 3, Number 12, 16-19, 2005.
4. J.-L. Marichal, P. Meyer, M. Roubens, Multiple Criteria Sorting : TOMASO, A Solution in the Presence of Interacting Points of View, *Newsletter of the EURO Working Group Multicriteria Aid for Decisions*, Series 3, Number 9, 10-12, 2004.

B.7 Communications à des séminaires et des conférences sans actes

1. S. Bigaret, P. Meyer, D. Kbaier, Innovative software for MCDA. *10th Decision Deck Workshop*, 11 April 2012, Tarragona, Spain, 2012.
2. S. Bigaret, V. Chiprianov, P. Meyer, J. Simonin, On the Formalization and Executability of the Decision Aid Process with Service Oriented Architecture. *75th*

- Meeting of the Euro Working Group MCDA*, 12-14 April 2012, Tarragona, Spain, 2012.
3. P. Meyer, S. Bigaret, T. Veneziano, diviz : un outil innovant pour une nouvelle méthodologie de travail en aide multicritère à la décision. *13e congrès annuel de la société française de recherche opérationnelle et d'aide à la décision*, 11-13 avril 2012, Angers, France, 2012.
 4. A.L. Olteanu, R. Bisdorff, P. Meyer, On multicriteria clustering. *10th Decision Deck Workshop*, 11 April 2012, Tarragona, Spain, 2012.
 5. T. Veneziano, P. Meyer, R. Bisdorff, Stabilité des affectations d'une méthode de tri en aide multicritère à la décision. *13e congrès annuel de la société française de recherche opérationnelle et d'aide à la décision*, 11-13 avril 2012, Angers, France, 2012.
 6. G. Retali, P. Meyer, M. Le Goff-Pronost, Implementation of remotely monitored medical dialysis units : dealing with multiple criteria and multiple decision makers. *10th Annual International Conference on Health Economics, Management and Policy*, 27-30 June 2011, Athens, Greece, 2011.
 7. G. Retali, P. Meyer, M. Le Goff-Pronost, Taking into account of patients and physicians preferences in the implementation of remotely monitored medical dialysis units. *8th World Congress International Health Economics Association : Transforming Health & Economics*, 10-13 July 2011, Toronto, Canada, 2011.
 8. O. Cailloux, P. Meyer, V. Mousseau, Eliciting ELECTRE TRI category limits for a group of decision makers. *The 21st International Conference on Multiple Criteria Decision Making*, 13-17 June 2011, Jyväskylä, Finland, 2011.
 9. B. Gourvennec, P. Meyer, G. Retali, *Elicitation de préférences d'utilisateurs potentiels de la TV 3D interactive*. 9e séminaire Marsouin, 26-27 Mai 2011, Bénodet, France, 2011.
 10. P. Meyer, S. Bigaret, Innovative methodological tools for research, teaching and consulting activities in MCDA. *8th Decision Deck Workshop*, 13 April 2011, Corte, France, 2011.
 11. A.L. Olteanu, R. Bisdorff, P. Meyer, A contribution to the multiple criteria preordering problematique . *74th Meeting of the European Working Group Multiple Criteria Decision Aiding*, 06-08 October 2011, Yverdon-Les-Bains , Switzerland, 2011.
 12. G. Retali, P. Meyer, Implementation of remotely monitored medical dialysis units : dealing with multiple criteria and multiple decision makers. *8th Decision Deck Workshop*, 13 avril 2011, Corte, France, 2011.
 13. T. Veneziano, P. Meyer, R. Bisdorff, Robust elicitation of criteria weights and thresholds in a multicriteria decision aid context. *73rd Meeting of the EURO Working Group MCDA*, 14-15 April 2011, Corte, France, 2011
 14. T. Veneziano, P. Meyer, R. Bisdorff, Elicitation robuste de paramètres en aide à la décision face à un décideur non-expert. *ROADEF 2011 (Recherche opérationnelle et d'aide à la décision)*, 02-04 march 2011, Saint-Etienne, France, pp. II-71, 2011.
 15. P. Meyer, S. Bigaret, XMCDA web services - latest available algorithms and their use via diviz. *7th Decision Deck Workshop*, 06 October 2010, Paris, France, 2010.
 16. P. Meyer, S. Bigaret, Designing, executing and sharing MCDA workflows via the diviz software and the XMCDA web services, *72nd meeting of the European Working Group "Multiple Criteria Decision Aiding"*, 07-09 October 2010, Paris, France, 2010.

17. P. Meyer, T. Veneziano, Decision aid protocols incorporating inverse decision analysis. *ILIAS Decision Aid Systems Seminar*, 25 March 2010, Luxembourg, Luxembourg, 2010.
18. S. Bigaret, P. Meyer, Tutorial : how to create and submit a web service proposal and integration into diviz. *5th Decision Deck Workshop*, 17-18 September 2009, Brest, France, 2009.
19. S. Bigaret, P. Meyer, Latest developments on diviz. *5th Decision Deck Workshop*, 17-18 September 2009, Brest, France, 2009.
20. R. Bisdorff, P. Meyer, T. Veneziano, XMCDa : a standard XML encoding of MCDA data. *EURO XXIII : European conference on Operational Research*, 05-08 July 2009, Bonn, Germany, 2009, pp. 53-53
21. T. Veneziano, S. Bigaret, P. Meyer, Diviz : an MCDA components workflow execution engine. *EURO XXIII : 23rd European Conference on Operational Research*, 05-08 July 2009, Bonn, Germany, 2009
22. P. Meyer, S. Bigaret, V. Mousseau, The Decision Project : Towards Open Source Software Tools Implementing Multiple Criteria Decision Aid. *JOPT 2009 : Journées de l'optimisation, Optimization days*, 4-5 mai, Montréal, 2009.
23. R. Bisdorff, P. Meyer, V. Mousseau, M. Pirlot, Latest News On The Decision Deck Project. *MCDA 69 : 69th Meeting of the Euro Working Group "Multiple Criteria Decision Aiding"*, April 2-3, Brussels, Belgium, Belgium, pp. 25, 2009.
24. P. Meyer, XMCDa 2.0, an XML schema dedicated to MCDA objects and data. Presentation and discussion. *4th Decision Deck Workshop*, 30-31 March 2009, Mons, Belgium, 2009.
25. S. Bigaret, P. Meyer, Diviz, an MCDA components workflow execution engine. *4th Decision Deck Workshop*, 30-31 March 2009, Mons, Belgium, 2009.
26. S. Bigaret, P. Meyer, BioSide : from bioinformatics needs to a generic workflow engine. *2nd Decision Deck Developers Days*, 3-4 December 2008, Paris, France, 2008.
27. G. Dodinet, P. Meyer, A deep dive into D3 web services. *2nd Decision Deck Developers Days*, 03-04 December 2008, Paris, France, 2008.
28. P. Meyer, T. Veneziano, A quick dive into XMCDa 2.0.0. *2nd Decision Deck Developers Days*, 03-04 December 2008, Paris, France, 2008.
29. P. Meyer, G. Ponthière, Eliciting preferences on multiattribute societies : an application of the Kappalab plugin, *Third Decision Deck Workshop*, Coimbra, Portugal, 16-17 June 2008.
30. P. Meyer, R. Bisdorff, J.L. Marichal, Disaggregation of bipolar-valued outranking relations and application to the inference of model parameters. *19th International Conference on Multiple Criteria Decision Making (MCDM 2008)*, Auckland, New Zealand, 7 - 12 January 2008.
31. P. Meyer, R. Bisdorff, J.-L. Marichal, Disaggregation of bipolar-valued outranking relations and application to the inference of model parameters, *38th Meeting of the European Mathematical Psychology Group (EMPG 2007)*, Luxembourg, 10-13 September 2007.
32. R. Bisdorff, P. Meyer, On the "fixation of belief" semantics and the rank of outranking digraphs, *38th Meeting of the European Mathematical Psychology Group (EMPG 2007)*, Luxembourg, 10-13 September 2007.

33. P. Meyer, R. Bisdorff, Solving the k-choice problem in a progressive context, *22nd European Conference on Operational Research*, Prague, Czech Republic, 8-11 July 2007.
34. P. Meyer, R. Bisdorff, M. Roubens, Pragmatic and logical foundations of the RUBY method for the recommendation of a best choice from a bipolar valued outranking relation, *37th Meeting of the European Mathematical Psychology Group (EMPG 2006)*, Brest, France, 11-13 September 2006.
35. P. Meyer, M. Grabisch, I. Kojadinovic, Kappalab : an R package for Choquet Integral based MAUT, *EURO XXI*, Reykjavik, Iceland, 2-5 July 2006.
36. R. Bisdorff, P. Meyer, RuBy : a new methodology for the best choice problematics, *EURO XXI*, Reykjavik, Iceland, 2-5 July 2006.
37. P. Meyer, C. Lamboray, R. Bisdorff, J.-L. Marichal, Decision Mathematics at the University of Luxembourg (Poster), *EURO XXI*, Reykjavik, Iceland, 2-5 July 2006.
38. P. Meyer, R. Bisdorff, M. Roubens, The RuBy method for the recommendation of a best choice from a bipolar valued outranking relation, *Seminar at the Dept. of Mathematics*, Tampere University of Technology, Finland, 17 April 2006.
39. P. Meyer, M. Grabisch, I. Kojadinovic, Kappalab : an R package for Choquet integral based MAUT, *63rd meeting of the EURO Working Group Multicriteria Aid for Decisions*, Porto, Portugal, 30-31 March 2006.
40. R. Bisdorff, P. Meyer, Exploitation en problématique de choix d'une relation de surclassement valuée : La méthode RuBy, *7ème congrès de la Société Française de Recherche Opérationnelle et d'Aide à la Décision (ROADEF)*, Lille, France, 6 - 8 February 2006.
41. R. Bisdorff, P. Meyer, RuBy : A new methodology for the best choice problematics, *ORBEL20*, Gent, Belgium, 19-20 January 2006.
42. P. Meyer, M. Roubens, On a fuzzy extension of the Choquet integral for ranking in Multiple Criteria Decision Aiding, *IFORS Triennial 2005*, Honolulu, Hawaii, 10-15 July 2005.
43. P. Meyer, On the use of the Choquet integral with fuzzy numbers in Multiple Criteria Decision Aiding, *SMA seminar*, University of Luxembourg, Luxembourg, 12 May 2005.
44. P. Meyer, Introduction to Multiple Criteria Decision Aiding, *Les midis de la science*, University of Luxembourg, Luxembourg, 11 & 14 April 2005.
45. P. Meyer, Use of an ordinal sorting technique (TOMASO) in stock selection, *59th meeting of the EURO Working Group Multicriteria Aid for Decisions*, Brest, France, 29-30 April 2004.
46. P. Meyer, Multiple Criteria Decision Aiding for the Selection of Quality Indexes of Association Rules, *Séminaire FNRS Recherche opérationnelle et aide à la décision multicritère*, Brussels, Belgium, 23 May 2003.
47. P. Meyer, TOMASO : a tool for ordered multiattribute sorting and ordering, *Seminar of the Department of Information and Knowledge Engineering*, University of Economics, Prague, Czech Republic, March 2003.
48. P. Meyer, MIKADO : Market Investigation and Knowledge Acquisition through Data Observation, *Seminar of the Department of Information and Knowledge Engineering*, University of Economics, Prague, Czech Republic, March 2003.

- 49. P. Meyer, Introduction to Rough Sets, TOMASO and Quality of Decision Rules, *Seminar of the IASC Department*, ENST Bretagne, Brest, France, June 2002.
- 50. P. Meyer, Qualification d'expertises d'analystes financiers dans le cadre de la gestion de portefeuilles d'actions, *Journée sur la classification : aspects combinatoires et algorithmiques*, Centre Universitaire de Luxembourg, Luxembourg, 29 January 2002.
- 51. P. Meyer, M. Roubens, La méthode TOMASO, *Seminar of the Department for Mathematics, Operational Research, Statistics and Information Systems*, Free University of Brussels (VUB), Brussels, Belgium, 2001.
- 52. P. Meyer, M. Roubens, La méthode et le logiciel TOMASO, *Séminaire FNRS Recherche opérationnelle et aide à la décision multicritère*, Brussels, Belgium, May 2001.
- 53. P. Meyer, M. Roubens, TOMASO : A Software for sorting in the presence of qualitative interacting points of view, *53rd meeting of the EURO Working Group Multicriteria Aid for Decisions*, Athens, Greece, 29-30 March 2001.