



Systèmes d'évaluation et de classification multicritères pour l'aide à la décision : Construction de modèles et procédures d'affectation

Laurent Henriet

► To cite this version:

Laurent Henriet. Systèmes d'évaluation et de classification multicritères pour l'aide à la décision : Construction de modèles et procédures d'affectation. Informatique [cs]. Université Paris Dauphine - Paris IX, 2000. Français. NNT : . tel-00528799

HAL Id: tel-00528799

<https://theses.hal.science/tel-00528799>

Submitted on 22 Oct 2010

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

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



UNIVERSITE PARIS
DAUPHINE

L'Université n'entend donner aucune approbation ni improbation aux opinions émises dans les thèses : ces opinions doivent être considérées comme propres à leurs auteurs.

Sommaire

Introduction	1
 Première Partie	
Problèmes de classification en Aide à la Décision	5
 1 Aide multicritère à la décision	6
1.1 Décision, aide à la décision et acteurs	8
1.2 Modélisation des préférences	9
1.2.1 Le concept d'action	9
1.2.2 Préférence et Indifférence	10
1.2.3 Incomparabilité et Surclassement	12
1.3 Analyse multicritère	13
1.3.1 La notion de critère	13
1.3.2 Différentes approches multicritères	15
1.3.3 Famille de critères	15
1.4 Problématique en aide à la décision	16
 2 Les différentes approches de la classification	18
2.1 Introduction	20
2.2 Classer les méthodes de classification	20
2.2.1 Classification et apprentissage	20
2.2.2 Les domaines de la classification	22
2.2.3 Une taxinomie des méthodes de classification	23
2.2.4 Classification graduelle ou non	24
2.3 Qu'est-ce qu'une catégorie?	26
2.4 Principales différences avec l'approche multicritère de la classification	27
2.5 Procédures multicritères d'affectation	29

2.5.1	Procédure Trichotomique de segmentation	29
2.5.2	nTOMIC	30
2.5.3	Electre Tri	31
2.5.4	Filtrage Flou par Préférence (FFP)	33
2.6	Méthodes non explicatives	34
2.6.1	Méthodes fondées sur l'utilisation de distance	34
2.6.2	Méthodes fondées sur des distributions de probabilité	35
2.6.3	Méthodes neuronales	36
2.7	Méthodes explicatives	38
2.7.1	Réseaux bayésiens	38
2.7.2	Méthodes non paramétriques fondées sur des distances	40
2.7.3	Méthodes de type <i>machine learning</i>	42

Deuxième Partie 46

Méthodes développées

3	Méthodes de Filtrage Flou	47
3.1	Introduction	49
3.2	Quelques limites des méthodes existantes	49
3.3	Construction d'une relation d'indifférence floue	51
3.3.1	Modélisation floue des préférences	51
3.3.2	Précisions sur l'utilisation des sous-ensembles flous	52
3.3.3	Coalitions floues concordantes et discordantes	55
3.3.4	Agrégation multicritère fondée sur un principe de concordance / non-discordance	57
3.4	Affectation multicritère	59
3.4.1	Définition formelle	59
3.4.2	Propriétés Générales d'une Méthode Multicritère d'Affectation par Compara- raisons Binaires	61
3.4.3	Filtrage par préférence	63
3.4.4	Filtrage par indifférence	65
3.4.5	Filtrage combiné	69
3.4.6	Représentation graphique	70
3.5	Méthodes de type Plus Indifférents Prototypes	70
3.5.1	Introduction	70
3.5.2	Filtrage flou par indifférence	71

3.5.3	Amélioration de la rapidité du calcul des plus indifférents prototypes	74
3.6	Explication de l'affectation	74
3.6.1	Structure des explications	74
3.6.2	Les différents types d'explication	76
4	Construction des poids des critères	79
4.1	Un modèle pour la construction des poids des critères	81
4.1.1	Pourquoi construire les poids des critères?	81
4.1.2	Objectif	81
4.1.3	Contraintes	82
4.1.4	Extraction d'information préférentielle et perspective interactive	83
4.2	Exemples de programmes sous forme linéaire	85
4.2.1	Hypothèses sur les opérateurs d'agrégation	85
4.2.2	Indifférence définie sur un principe de concordance	86
4.2.3	Indifférence définie sur le principe de concordance et non-discordance	86
4.3	Exemple	87
4.3.1	Définition du problème	87
4.3.2	Formalisation	87
4.3.3	Résultats	88
4.4	Perspective utilisant une approche connexioniste	90
5	Construction des prototypes et problématique de la typologie	91
5.1	Pourquoi construire des points de référence?	93
5.2	Construction de prototypes par nuées dynamiques	94
5.2.1	Principe	94
5.2.2	Mise en œuvre	95
5.2.3	Exemple	98
5.3	Construction de prototypes par algorithme génétique	99
5.3.1	Principe des algorithmes génétiques	100
5.3.2	Fonction d'évaluation	101
5.3.3	Opérateurs d'évolution utilisés	103
5.3.4	Algorithmes	107
5.3.5	Exemple	107

Troisième Partie	
Mise en œuvre	111
6 Logiciel FFI	112
6.1 Introduction	114
6.2 Méthodes mises en œuvre	114
6.3 Fonctionnalités	116
6.3.1 Démarrage	116
6.3.2 Création des fichiers d’actions à classer et de prototypes	116
6.3.3 Chargement de Fichiers	117
6.3.4 Sauvegarde	117
6.3.5 Gestion des données : actions, catégories et critères	117
6.3.6 Visualisation graphique des actions	117
6.3.7 Affectation des actions	122
6.4 Évolution	122
Conclusion	125
Annexes	128
A Exemples relatifs au chapitre 1	129
B Exemples d’affectation des employés	131
C Méthode de nuées dynamiques	138
Rappel des Notations	141
Table des figures	143
Liste des tableaux	144
Bibliographie	145

Introduction

Le domaine de la recherche opérationnelle et de l'optimisation a connu un essor important au cours de la deuxième moitié du XX^{ème} siècle à la fois par des méthodes numériques permettant de résoudre des problèmes concrets et par la croissance constante de la puissance de calcul des ordinateurs permettant la mise en œuvre de méthodes de plus en plus complexes. Le désir de mieux intégrer des outils de recherche opérationnelle dans le domaine de la décision organisationnelle a conduit à développer une nouvelle conception de l'intervention scientifique : l'*aide à la décision*. Celle-ci a pour but d'aller au delà de la simple prise en compte de données "objective" en tentant d'intégrer les aspects subjectifs qui interviennent dans un processus de décision. Ainsi, l'objectif de l'aide à la décision se définit comme l'accompagnement du décideur dans la *construction* d'une décision satisfaisante. Cela ne passe pas uniquement par l'utilisation d'une méthode mais aussi par une aide à la compréhension et à la structuration du problème, une aide à la communication entre les différents intervenants.

Au cours d'un processus de décision, les aspects à prendre en compte dans toute évaluation (que ce soit pour une affectation à des catégories ou non) sont généralement multiples. Lorsque les données le permettent (données homogènes, échelles facilement mesurables ...) il peut être aisé de les regrouper et de les synthétiser directement en une note (*score*) qui constitue le résultat de l'évaluation. Cependant cette approche devient problématique lorsque les aspects à prendre en compte sont hétérogènes, partiellement contradictoires et mal appréhendés. Une telle situation incite naturellement à considérer plusieurs points de vue que l'on ne peut résumer en un seul.

L'approche *Multicritère* de l'aide à la décision, qui a connu un essor important au cours de ces dernières années, apporte des éléments de réponse au problème de l'évaluation selon de multiples points de vue contradictoires. Elle consiste en la définition et l'étude approfondie de *critères* supposés refléter les différents points de vue entrant en considération dans un processus de décision. Nous parlerons à cet effet d'*Aide Multicritère à la Décision* et c'est l'optique dans laquelle nous nous plaçons tout au long de ce document.

Dans le domaine de l'aide multicritère à la décision, l'effort a principalement porté sur les problèmes de choix ou de rangement, c'est-à-dire comment sélectionner des meilleures solutions ou comment ordonner un ensemble de solutions de la meilleure à la moins bonne. Jusque dans les années 1980, peu de méthodes multicritères de classification (on parle aussi de méthodes de tri) ont été développées (voir chapitre 3) ; à notre connaissance la première fut introduite par Moscorola et Roy [70] et [71] en 1977. La majeure partie des méthodes développées depuis sont des méthodes fondées sur la prise en compte des préférences d'un décideur considérant des catégories ordonnées : par exemple, pour l'évaluation des systèmes de culture respectueux de l'environnement, six catégories ont été considérées correspondant aux différents niveaux de respect de l'environnement ; elles étaient donc ordonnées de la plus respectueuse à la moins respectueuse. Extrêmement peu de méthodes multicritères permettent en revanche de considérer des catégories non ordonnées. On peut citer plusieurs applications s'inscrivant dans ce cadre : -1- le conseil de produits financiers où l'objectif est d'affecter des dossiers de clients à des produits financiers divers (les produits financiers correspondent à des catégories non ordonnées) -2- l'aide à l'orientation de personnel, cela consiste à orienter des individus en fonction de l'évaluation de leurs compétences vers des filières (que ce soit

des étudiants à des filières d'enseignement ou des employés à des services dans une entreprise) -3- la proposition de produits de consommation à des individus effectuant du commerce électronique.

Parallèlement, les statisticiens et les probabilistes, et plus récemment la "communauté intelligence artificielle", se sont largement intéressés aux problèmes et méthodes de classification. Cependant, comme nous le verrons plus précisément au chapitre 2, le terme *classification* peut revêtir plusieurs sens. Afin de mieux cerner ce que l'on entend par classification, nous allons tenter d'en donner une définition intuitive. Voyons comment Le Petit Larousse définit ce terme; il s'agit d'un *distribution par classes, par catégories, selon un certain ordre ou une certaine méthode*. À travers cette définition, nous retrouvons entre autres la notion d'*affectation*. C'est cette idée que nous retiendrons pour définir un problème de classification: nous dirons qu'un *problème de classification va consister à affecter des objets, des candidats, des actions potentielles à des catégories ou des classes prédéfinies*

Les différents exemples qui suivent illustrent assez bien la multitude des problèmes qui s'inscrivent naturellement dans l'optique de la classification.

- l'évaluation de dossiers de crédits [71]

Il s'agit d'évaluer des dossiers de crédits en les affectant à plusieurs classes correspondant aux traitements que l'on doit effectuer sur ceux-ci.

- le diagnostic médical (détection de tumeurs leucémiques) [7]

Il s'agit de détecter différents types de leucémies; les objets sont ici les patients et les catégories, les maladies.

- l'évaluation environnementale [3]

Cette application consiste en l'évaluation de systèmes de culture en fonction de leur impact sur l'eau de profondeur afin dans un premier temps de donner des conseils aux exploitants puis dans un second temps d'envisager une politique de taxation et de subvention; les objets sont les différents systèmes de culture et les catégories correspondent aux différents impacts sur l'eau de profondeur.

- la reconnaissance de la parole [73]

Cette application consiste en la reconnaissance d'une personne par sa voix: les objets sont donc les voix et les catégories les individus auxquelles elles doivent être associées.

- reconnaissance des cafés [39]

Cette application consiste à identifier des mélanges de café et à donner leur composition; les objets sont des mélanges de cafés et les catégories correspondent aux types de café utilisés pour la composition des mélanges.

Le champ de la classification est assez large. Les applications que nous venons de citer, même si elles font toutes référence à la classification, sont néanmoins différentes les unes des autres. Certaines, comme la reconnaissance des cafés ou la reconnaissance de la parole s'inscrivent dans le cadre de la classification et de la décision automatique, d'autres comme l'évaluation de dossiers de crédits et l'évaluation environnementale sont plus proches de l'activité de conseil où la part de subjectivité et les préférences d'un décideur doivent être prises en compte, enfin, la détection de tumeurs, qui correspond à une véritable aide au diagnostic médical peut être vue comme un autre type de conseil où l'objectif n'est pas nécessairement de prendre en compte les préférences d'un décideur. Notre travail se positionne au carrefour de l'*aide multicritère à la décision* et de la *classification*.

L'objectif de notre travail est de proposer une démarche d'aide à la décision s'appuyant sur des techniques de classification issues de l'analyse de donnée et de l'intelligence artificielle. Notre démarche se distingue notamment de celles existantes par :

- l'utilisation d'une modélisation fine des préférences pour décider de l'affectation à une catégorie et favoriser l'élaboration de recommandations auprès du décideur,
- leur capacité à produire une explication de l'affectation effectuée. Cet aspect revêt un intérêt tout particulier en AD où le décideur recherche tout autant à classer des objets que de pouvoir expliquer et argumenter les affectations effectuées,
- l'utilisation de méthodes d'apprentissage numériques issues de l'intelligence artificielle et de l'analyse de données, afin de construire en interaction avec le décideur, les modèles d'affectation que nous proposons.

En outre notre démarche vise à produire de nouvelles méthodes multicritères de classification présentant les caractéristiques techniques suivantes :

- la capacité à effectuer une affectation graduelle à des catégories, c'est-à-dire à donner un degré d'affectation aux catégories. Cela permet d'appréhender la plus ou moins grande appartenance d'un objet à une catégorie. Par exemple, dans le cadre de l'affectation d'automobiles à des catégories, il va être intéressant et utile de pouvoir exprimer qu'une automobile est plus ou moins économique ou plus ou moins sportive. Cette information graduelle s'exprime à travers un degré d'affectation,
- la capacité à considérer indifféremment des catégories ordonnées ou non et des catégories ne formant pas nécessairement une partition. Cela permet d'aborder aussi bien des problèmes complexes de diagnostic où les catégories ne possèdent aucune structure en particulier (ni ordre, ni partition) que des problèmes plus simples de discrimination en deux classes (exemple : les bons et les mauvais clients),
- la capacité à considérer des catégories multiformes (on parlera de catégories multiprofiles), c'est-à-dire auxquelles un objet peut appartenir par différents moyens. Considérant les élèves d'une classe, il peut exister plusieurs manières d'être globalement bon élève : on peut être bon en sciences sans être mauvais en matières littéraires, mais on peut aussi être bon en lettres sans être mauvais en sciences, on peut enfin être bon partout.

De plus, en complément de nos méthodes d'affectation, nous souhaitons aborder la construction de paramètres et de faciliter ainsi la construction des catégories. Abordant les deux vastes thèmes que sont *l'aide multicritère à la décision* et la *classification*, nous commençons par la présentation de ceux-ci. Ainsi, le chapitre 1 introduit l'aide multicritère à la décision à travers les notions de préférence, de décideur, d'actions et de critères. Le chapitre 2 expose quant à lui les différentes notions généralement rattachées à la classification (apprentissage, les domaines de la classification, les types de méthodes et le problème de la classification graduelle). Il clôt cette première partie de présentation en dressant un panorama des méthodes de classification issues de différents domaines en s'intéressant au côté explicatif des méthodes (ce panorama n'a aucune vocation d'exhaustivité, de nombreux ouvrages existant mais n'y suffisant pas).

La deuxième partie de ce document est consacrée aux différentes méthodes que nous avons développées. Le chapitre 3 présente tout d'abord la définition que nous entendons donner des Méthodes Multicritères de Classification par Comparaisons Binaires. Puis nous abordons la méthode de classification multicritère dite de filtrage flou que nous avons conçue ainsi que la possibilité laissée à ce type de méthode de produire une explication de leur résultat (ceci revêtant un intérêt tout particulier en aide à la décision). Nous abordons ensuite au chapitre 4 le problème de la construction

d'un des paramètres utilisés dans l'évaluation des préférences du décideur (les poids des critères). Enfin, le chapitre 5 traite du problème de la construction de la représentation des catégories ; en effet, un décideur n'est pas nécessairement à même de donner une représentation adéquate des classes, et c'est pourquoi nous proposons deux procédures permettant d'aider celui-ci à construire une représentation des catégories.

La troisième et dernière partie est dédiée à la mise en œuvre de nos travaux. Nous présentons au chapitre 6 un prototype du logiciel d'aide à la décision dans lequel est implantée la méthode de classification multicritère que nous avons développée.

PREMIÈRE PARTIE

Problèmes de classification en Aide à la Décision

Chapitre 1

Aide multicritère à la décision

Résumé

Nous présentons dans ce chapitre les principaux concepts autour desquels s'articule l'aide multicritère à la décision. En particulier, la différence fondamentale qui existe entre les sciences de décision et l'aide à la décision. Alors que les sciences de décision vont avoir pour objectif de trouver une décision optimale à partir d'une vision supposée objective de la réalité, l'aide à la décision va s'intéresser à la construction de décisions satisfaisantes en considérant toute la dimension subjective qui peut apparaître au cours d'un processus de décision. Pour cela, une distinction est faite entre l'homme d'étude et le décideur, c'est-à-dire entre le technicien des méthodes d'aide à la décision et la personne, ou le groupe de personnes, en charge de prendre les décisions. L'aide à la décision va avoir pour objet de faire agir conjointement ces deux principaux acteurs afin de faire émerger les décisions. La prise de décision fait généralement intervenir des points de vue différents, voire contradictoire. La tâche de l'homme d'étude va alors être de modéliser les préférences du décideur en faisant émerger les différents points de vue, ou critères, entrant en compte dans le processus de décision.

Mots clés : aide à la décision, multicritère, critère, action, préférence, indifférence, problématique

1.1 Décision, aide à la décision et acteurs

De nombreuses disciplines scientifiques (statistiques, analyse de données, programmation mathématique, sciences économiques . . .), reposent sur l'hypothèse de l'existence d'un *critère objectif d'optimalité*. Ces trois termes revêtent une importance fondamentale en *science de la décision*. En effet, les *sciences de décision* supposent explicitement les trois aspects suivantes :

- On suppose qu'il existe une *meilleure décision* et qu'il est possible, en disposant de suffisamment de temps et de moyens, de l'obtenir,
- On suppose que cette décision optimale peut être obtenue en *optimisant un critère* (ex. fonction de coût, fonction d'utilité),
- On suppose que la décision optimale peut être obtenue quelle que soit la méthode ou la procédure utilisée ; c'est là son caractère *objectif*.

Au contraire, comme le souligne Roy [99] la discipline de l'aide à la décision ne repose pas sur l'existence d'une vérité absolue. Puisque cette vérité n'est pas supposée exister, l'objectif va être de guider et d'éclairer une entité appelée *décideur* tout au long de son processus de décision. On ne cherche pas à trouver "*la meilleure décision*" mais à accompagner le décideur en tentant de faire ressortir les aspects objectifs et ceux qui le sont moins, d'apporter une justification aux décisions en mettant en évidence les conclusions robustes par rapport à celles qui le sont moins. Roy [97] propose la définition suivante :

Définition 1 (aide à la décision) *L'aide à la décision est l'activité de celui qui, prenant appui sur des modèles clairement explicités mais non nécessairement clairement formalisés, aide à obtenir des éléments de réponse aux questions que se pose un intervenant dans un processus de décision, éléments concourants à éclairer la décision et normalement à prescrire, ou simplement à favoriser, un comportement de nature à accroître la cohérence entre l'évolution d'un processus d'une part, les objectifs et le système de valeurs au service desquels cet intervenant se trouve placé d'autre part.*

Au cours des processus d'aide à la décision, il convient de distinguer principalement deux intervenants : *l'homme d'étude* et le *décideur*. Roy et Bouyssou [100] en donnent les définitions suivantes :

Définition 2 (décideur) *Le décideur est l'entité intervenant dans le processus de décision que les modèles mis en œuvre cherchent à éclairer, c'est l'entité pour le nom de qui, ou au compte de qui, l'aide à la décision s'exerce.*

Contrairement à certaines idées reçues le décideur n'est pas nécessairement un individu en particulier. Ce peut être dans certains cas un individu mais ce peut être aussi un groupe d'individus pas nécessairement bien identifié (ex. le gouvernement, les personnes qui le composent peuvent évoluer rapidement sans pour autant empêcher que des décisions soient prises en continu).

Définition 3 (homme d'étude) *L'homme d'étude est celui qui prend en charge l'aide à la décision. Mettant en œuvre des modèles dans le cadre d'un processus de décision, il contribue à l'orienter et à la transformer.*

Le décideur est l'entité qui va prendre la décision en assumant la responsabilité. Il cherche au travers de l'aide à la décision, non pas un moyen de prendre les décisions à sa place, mais un moyen d'éclairer et aussi de justifier ses décisions. On parlera donc de *méthodes d'aide à la décision* plutôt que de *méthodes de prise de décision*. L'homme d'étude va donc être celui qui aide le décideur au

cours du processus d'aide à la décision. C'est lui qui a la connaissance des méthodes d'aide à la décision. En revanche, il n'a aucunement pour fonction de prendre la décision mais a avant tout un rôle de conseil et d'éclaircissement. L'aide à la décision va avoir de multiples rôles, et principalement d'aider à mieux formuler les problèmes à résoudre, d'aider à mieux communiquer afin de que les différents acteurs puissent mieux se comprendre. Ainsi nous voyons que l'aide à la décision apparaît comme une approche beaucoup plus large que celle de la prise de décision.

Il s'agit là d'une démarche *constructive*. Le processus d'aide à la décision peut être défini comme un processus de construction d'une décision satisfaisante et non pas comme la découverte d'une solution existante reflétant une décision optimale objective. Si une décision optimale existe (encore faut-il avoir défini des critères d'optimalité), il est la plupart du temps illusoire, voire impossible, de pouvoir prouver ou montrer le caractère optimal de celle-ci. L'objectif en aide à la décision va donc être construire une décision satisfaisante que le décideur va être capable de justifier, c'est-à-dire de lui donner des raisons suffisamment fortes et objectives pour l'argumenter.

Cependant le domaine d'application des méthodes que nous proposons par la suite ne se limite pas à l'aide à la décision comme nous venons de la présenter. Elles peuvent aussi trouver des champs d'application dans le cadre de la décision automatique. Il s'agit d'un domaine où les décisions doivent être prises de manière répétée. La plupart du temps, les procédures de décision automatique vont avoir pour objectif de reproduire une tâche effectuée par un expert. Dans cette approche, les décisions ne sont pas réellement prises en collaboration avec un décideur. C'est une approche, non pas uniquement constructive comme la précédente, mais aussi en partie *descriptive* dans le sens où l'on cherche à décrire le processus de décision d'un expert. Cependant si l'aspect descriptif est réellement présent, il ne s'agit pas nécessairement uniquement de reproduire un processus effectué par un expert (voir [49] qui donne plusieurs exemples d'applications utilisant des systèmes experts).

Remarque 1 (Décision automatique et CNIL) *À propos du domaine de la décision complètement automatique, il convient de noter que les décisions relevant d'un processus complètement automatique doivent se limiter légalement à des domaines ne touchant pas la personne humaine ainsi la loi numéro 78-17 du 6 janvier 1978 (J.O. 07/01/1978 et 25/01/1978) relative à l'informatique, aux fichiers et aux libertés [19] spécifie que :*

Aucune décision de justice impliquant une appréciation sur un comportement humain ne peut avoir pour fondement un traitement automatisé d'informations donnant une définition du profil ou de la personnalité de l'intéressé.

Aucune décision administrative ou privée impliquant une appréciation sur un comportement humain ne peut avoir pour fondement un traitement automatisé d'informations donnant une définition du profil ou de la personnalité de l'intéressé.

Ainsi, de part cette loi, le champ de la décision automatique est réglementé.

1.2 Modélisation des préférences

1.2.1 Le concept d'action

Toute discipline scientifique est amenée à manipuler et à raisonner sur des objets plus ou moins bien formalisés. Ainsi dans le domaine des statistiques, les objets manipulés vont par exemple être des points décrits par des vecteurs. En aide à la décision, les objets à manipuler vont être les actions potentielles. Vincke [112] définit simplement *l'ensemble des actions noté A , comme l'ensemble des objets, décisions, candidats ... que l'on va explorer dans le processus de décision*. Citons quelques exemples :

- Si l'on recherche un emplacement pour construire une usine dans une certaine région, les

actions vont être les différents lieux envisagés, l'objectif étant d'en choisir un.

- Si l'on considère un problème de voyageur de commerce, l'ensemble des actions sera l'ensemble des tournées possibles pour le voyageur, c'est-à-dire l'ensemble des circuits hamiltoniens du graphe complet qui modélise les problèmes.
- Si dans le domaine bancaire, on cherche à proposer aux clients des produits financiers adaptés à leur profil. L'objectif sera alors d'affecter les différents clients potentiels à des catégories représentant ces produits financiers. Les actions sont ici les clients.

Dans les trois exemples précédents, la définition des actions est assez intuitive. Ce n'est cependant pas nécessairement le cas et l'ensemble des actions ne correspond pas forcément à une vérité objective et unique. Il existe souvent plusieurs manières de modéliser un problème. Ainsi certaines modélisations simples peuvent nécessiter l'utilisation d'une méthode complexe alors que des modélisations moins immédiates permettent parfois une résolution à l'aide de méthodes simples (*cf.* l'exemple 38 des 6 cavaliers [64] en annexe A).

Il nous semble important de préciser la définition du concept d'action par une notion plus précise, celle d'action potentielle, définition (4), introduite par Roy [97].

Définition 4 (action potentielle) *Une action potentielle est une action réelle ou fictive provisoirement jugée réaliste par un acteur au moins, ou présumée telle par l'homme d'étude, en vue de l'aide à la décision ; l'ensemble des actions potentielles sur lequel l'aide à la décision prend appui au cours d'une phase d'étude est noté A .*

Les actions potentielles ont pour objet de délimiter le champ des solutions possibles. On peut faire apparaître deux types d'actions potentielles : les actions réelles et les actions fictives. Les premières correspondent à une réalité susceptible d'être appréhendée par le décideur. Par exemple dans le choix d'une voiture économique, une Renault 4L est une action réelle. On peut néanmoins vouloir considérer des actions qui ne correspondent à aucune réalité existante mais qui permettent quand même d'éclairer les décisions, ce sont les actions fictives. Ces actions fictives vont servir de base pour effectuer des comparaisons, elles sont par exemple utilisées dans les loteries de la Théorie de l'Utilité [91], en programmation mathématique avec le point idéal ou encore dans Prefcalc avec les points idéal et anti-idéal [54]. Par la suite nous utiliserons ce genre d'action dans les méthodes de classification que nous proposons, ces actions fictives sont alors appelées points de référence et servent à décrire les catégories. Si l'on demeure dans le domaine automobile, on peut considérer des actions fictives qui représentent ce qu'est une voiture typiquement familiale ou une voiture typiquement sportive pour un décideur. Dans le domaine de la classification multicritère, les actions à classer vont être comparées à des actions fictives pour évaluer leur appartenance aux classes.

1.2.2 Préférence et Indifférence

L'activité d'aide à la décision passe par la comparaison des actions entre elles. Cela se traduit par l'utilisation de relations de comparaison. Vincke [112] définit le modèle classique, celui-ci ne considère que deux relations : les relations de préférence et d'indifférence, que nous noterons respectivement $\sim$ et $\succ$ (ou I et P).

Préférence : cette relation permet de traduire une situation dans laquelle il existe des raisons claires et suffisantes pour mettre en évidence une préférence entre deux actions a et b . On notera $a \succ b$ (ou aPb) une situation dans laquelle a est préférée à b . De part la sémantique associée à cette relation, il est naturel de considérer cette relation comme étant irreflexive et asymétrique.

Indifférence : cette relation traduit une situation dans laquelle il n'existe pas de raisons suffisamment fortes pour confirmer une préférence dans un sens ou dans l'autre. On notera $a \sim b$ (ou aIb) une situation d'indifférence entre a et b . Cette relation est généralement considérée comme étant réflexive et symétrique.

Ces deux relations apparaissent ainsi comme complémentaires. Lorsque deux actions a et b sont indifférentes, il n'est pas possible d'affirmer une préférence dans un sens ou dans l'autre. De même, lorsqu'il existe une préférence entre a et b , les deux actions ne peuvent être indifférentes.

$$a \sim b \Leftrightarrow \neg(a \succ b \vee b \succ a)$$

Il est souvent admis que les deux relations, $\sim$ et $\succ$, sont transitives. Elles possèdent alors de *bonnes propriétés* qui permettent de considérer que les actions de A forment un préordre complet. Roy [97] parle de *parfaite comparabilité transitive*. Cependant ces *bonnes propriétés* qui s'appliquent parfaitement à l'ensemble des entiers, si l'on considère les relations d'égalité ($=$) et de supériorité ($>$) qui forment le pendant des relations $\sim$ et $\succ$ pour A , ne sont pas nécessairement justifiées dans le cadre de la modélisation des préférences.

S'il est possible de comparer la relation d'égalité à la relation d'indifférence, il ne faut pas admettre qu'elles sont équivalentes et qu'elles possèdent les mêmes propriétés. Le caractère transitif de l'indifférence peut aisément être remis en cause, voyons cela sur l'exemple (1) de Poincaré [86].

Exemple 1 (intransitivité de l'indifférence, Poincaré) *H. Poincaré constate qu'un poids a de masse 10 grammes ne peut être distingué par un homme d'un poids b de masse 11 grammes. Par ailleurs b ne peut être distingué d'un autre poids c de masse 12 grammes. Cependant, il est assez facile de distinguer a de c . Ainsi en considérant un système de préférence où l'on recherche des poids plus légers, on a les trois relations suivantes :*

$$\begin{aligned} a &\sim b \\ b &\sim c \\ a &\succ c \end{aligned}$$

De nombreux autres auteurs se sont penchés sur le problème de l'intransitivité de l'indifférence en construisant à chaque fois des exemples mettant en évidence cette propriété, nous en citons deux parmi plusieurs. Armstrong [2] prend l'exemple de sandwiches au fromage. Luce [66] quant à lui décrit ces phénomènes d'intransitivité à l'aide de la structure de quasi ordre à travers l'exemple de sucres dans une tasse de café (*cf.* annexe A). Nous admettrons par la suite que les relations d'indifférence ne possèdent pas nécessairement de propriété de transitivité.

La relation d'indifférence peut être rangée dans la catégorie des relations de similarité. Ce type de relation, habituellement considéré comme symétrique peut dans certains ne pas l'être. Tversky [109] montre plusieurs exemples mettant en évidence cette caractéristique. Enfin, on peut citer Bouchon-Meunier *et al.* [14] et Rifqi [92] qui présentent une définition générale des relations de comparaisons (similitude, ressemblance inclusion, dissimilarité, ...) et de leurs propriétés.

De même que pour l'indifférence, il est possible de remettre en cause la propriété de transitivité pour la relation de préférence. Pour s'en convaincre, il suffit de considérer une procédure de vote majoritaire pour affirmer une préférence, voir exemple (2).

Exemple 2 (effet Condorcet) *Considérons trois actions a , b et c devant être jugées selon trois points de vue. On utilise la règle suivante pour établir une préférence d'une action sur une autre :*

points de vue	1	2	3
a	15	15	5
b	10	10	10
c	5	17	7

TAB. 1.1 – Évaluations des actions a , b et c

une action sera préférée à une autre si elle est meilleure sur la majorité des points de vue considérés (a est meilleure que b selon un point de vue si l'évaluation de a selon ce point de vue est supérieure à celle de b). Considérons les évaluations du tableau (1.1).

Il apparaît en considérant la règle précédente que :

$$a \succ b$$

$$b \succ c$$

$$c \succ a$$

Ce phénomène de non transitivité des préférences a été étudié par de nombreux auteurs. Le premier à s'y être intéressé est Condorcet [26] dans le cadre de l'étude des procédures de vote (voir Vansnick [111] qui présente les méthodes de Condorcet et de de Borda [12] ainsi que leur application aux principes du choix social et à l'agrégation multicritère). Les difficultés sous-jacentes à l'obtention de préférence globales transitives ont été montrées par Arrow [4] et [5], dont les résultats ont engendré une très riche littérature [60]. On peut enfin citer les travaux de Tversky [108] qui propose plusieurs exemples mettant à jour ce phénomène d'intransitivité.

1.2.3 Incomparabilité et Surclassement

Dans le domaine de l'aide à la décision, tout ne peut pas s'exprimer uniquement en terme de préférence ou d'indifférence. Il peut exister des situations où le décideur ne peut s'exprimer en faveur d'une action ou d'une autre sans pour autant être indifférent : on parlera de situations d'incomparabilité.

Prenons pour exemple la comparaison d'appareils de transport aérien. S'il est aisément possible de mettre en évidence des préférence (relations $\sim$ et $\succ$) entre les différents modèles de Boeing 737, un Airbus A320 et un Fokker 100, il l'est beaucoup moins entre un Concorde, un Boeing 747-400 et ATR 42 : c'est typiquement une situation d'incomparabilité.

Ce type de situation peut se rencontrer lorsque les informations sur les actions sont insuffisantes pour trancher entre une préférence et une indifférence. Une des autres causes que l'on retrouve à l'origine de la situation d'incomparabilité est le fort contraste qui peut exister entre plusieurs actions. En présence d'actions contrastées, c'est-à-dire possédant des points forts et des points faibles mais sur des aspects différents, il sera souvent difficile de mettre une préférence ou une indifférence en évidence, par exemple, le supersonique franco-britannique est un avion extrêmement performant contrairement à l'ATR42, en revanche il est d'un coût exorbitant contrairement au biturbopropulseur français.

Formellement, nous noterons R la relation d'incomparabilité, $a R b$ signifiera que l'action a est incomparable à l'action b . Cette relation est naturellement considérée comme symétrique et irréflexive. Comme pour les autres relations utilisées en modélisation des préférences, nous considérerons que cette relation n'est pas nécessairement transitive. Contrairement à celles-ci il est beaucoup plus facile de s'en persuader. Derrière la notion d'incomparabilité on ne retrouve nullement les notions

appareil	M.M. (kg)	capacité (pers.)	C.U. (kg)	V.M. (km.h ⁻¹)	autonomie (km)
AS 532 A2/U2	9750	2+29	6300	273	1180
AS 532 SC	9000	2+21	4480	251	911
Tigre	5925	1+1	1800	280	800

TAB. 1.2 – *Performances*

d'ordre ou de proximité; les actions incomparables ont d'ailleurs plutôt tendance à être très différentes. Il est donc tout à fait naturel de penser que, si $a R b$ et $b R c$, on puisse quand même établir une préférence (ou une indifférence) entre a et c , voir exemple (3).

Exemple 3 *Considérons trois hélicoptères militaires, les Cougar AS 532 SC, AS 532 A2/U2 et Tigre que l'on cherche à évaluer selon la masse maximale (M.M.), la capacité, la charge utile, la vitesse maximale (V.M.) et l'autonomie. Les performances sont présentées dans le tableau 1.2.*

Il semble évident qu'au vu des performances des trois hélicoptères, le Tigre est visiblement difficilement comparable aux deux autres. Cependant il est assez facile de mettre en évidence une préférence entre l'AS 532 SC et l'AS 532 A2/U2 puisque sur tous les critères l'AS 532 A2/U2 domine l'AS 532 SC. On a donc les trois relations suivantes qui témoignent de la non transitivité de l'incomparabilité :

$$\begin{array}{lll}
\text{AS 532 A2/U2} & R & \text{Tigre} \\
\text{AS 532 SC} & R & \text{Tigre} \\
\text{AS 532 A2/U2} & \Delta & \text{AS 532 SC}
\end{array}$$

avec Δ la relation de dominance.

Parallèlement aux trois relations définies précédemment, nous allons en considérer une quatrième qui correspond au regroupement de $\sim$ et de $\succ$. Il s'agit de la relation de surclassement, notée S , définie par Roy [93], voir aussi Fishburn [34].

Soient deux actions a et b , $a S b$ signifie que a est au moins aussi bon que b , ce qui correspond à une des deux situations suivantes : soit a est indifférent à b , soit a est préféré à b . Il est alors naturel que S soit réflexive et non nécessairement transitive. Cette relation S est liée aux trois autres relations $\sim$, $\succ$ et R vues précédemment par les équations suivantes :

$$\begin{array}{lll}
a S b & \Longleftrightarrow & (a \sim b) \vee (a \succ b) \\
(a S b) \wedge (b S a) & \Longleftrightarrow & a \sim b \\
(a S b) \wedge \neg(b S a) & \Longleftrightarrow & a \succ b \\
\neg(a S b) \wedge \neg(b S a) & \Longleftrightarrow & a R b
\end{array}$$

1.3 Analyse multicritère

1.3.1 La notion de critère

Face à un problème de décision, le décideur et l'homme d'étude vont être amenés à juger et à évaluer les différentes actions potentielles. Pour rendre compte de tels jugements, on fait appel à

la notion de critère. Les critères vont être le moyen utilisé pour décrire les actions. Formellement, on représentera les critères par des fonctions à valeurs réelles. On peut citer Vincke [112] qui donne une définition concise des critères (définition 5) et Fishburn [35] qui propose une définition plus formelle (définition 6).

Définition 5 (critère, Vincke) *Un critère est une fonction g , définie sur l'ensemble A des actions, qui prend ses valeurs dans un ensemble totalement ordonné, et qui représente les préférences du décideur selon un point de vue. Lorsque le problème repose sur la considération de plusieurs critères, nous les notons $g_1, \dots, g_n$: dans la suite nous parlerons indifféremment du critère g_j ou du critère j . L'évaluation d'une action a suivant le critère j est notée $g_j(a)$.*

Définition 6 (critère, Fishburn) *[...] a criterion function usually indicates a real valued function on X that directly reflects the worth or value of the elements in X according to some criterion or objective. [...] Unlike attribute mappings, which usually describe objective characteristics of alternatives or consequences, criterion functions often represent subjective values on a more or less arbitrary scale. [...] In a situation with n criteria ($j = 1, \dots, n$) and corresponding criteria functions $g_j : X \rightarrow \mathbb{R}$, each x in X is mapped into an n -tuple $(g_1(x), \dots, g_n(x))$ of criterion values, scores or utilities. It is then common to associate some notion of preference or value with these n -tuples. It is often assumed, for example, that preference monotonically increases in each g_j*

Un critère peut donc être défini comme le moyen de modéliser un point de vue. Cependant, plusieurs aspects d'une action peuvent concourir à un même point de vue. Par exemple, si l'on s'intéresse au point de vue *confort d'une automobile*, plusieurs aspects doivent être pris en compte comme la suspension, la tenue de route, le niveau sonore, etc. Un critère peut donc résulter d'une agrégation de *sous critères*. Roy [97] définit l'ensemble de ces sous critères comme le sous nuage des conséquences et définit un critère de la manière suivante :

Définition 7 (critère, Roy) *Une fonction g à valeurs réelles définie sur A est, pour un acteur Z , une fonction critère ou un critère appréhendant le sous nuage des conséquences $v_g(A)$ si :*

1. *Le nombre $g(a)$ est déterminé si et seulement si une évaluation $\Gamma_g(a)$ de $v_g(a)$ est disponible ; le modèle $\Gamma_g(A)$ qui fournit cette évaluation est appelé support de la fonction critère g .*
2. *L'acteur Z (ou l'homme d'étude jugeant au nom de Z) reconnaît l'existence d'un axe de signification sur lequel deux actions potentielles quelconques a et a' peuvent être comparées relativement aux seuls aspects des conséquences que recouvrent $v_g(A)$ et il accepte de modéliser conformément à :*

$$g(a') \geq g(a) \implies a' S_g a$$

où S_g désigne une relation de surclassement restreint à l'axe de signification du critère g (faisant en particulier abstraction de tous les aspects des conséquences non modélisées dans le support de g).

Chaque action a de A sera donc représentée dans l'espace des critères, $E = E_1 \times \dots \times E_n$, par un vecteur $(g_1(a), \dots, g_n(a))$, que l'on appelle vecteur de performances. L'ensemble A des actions sera représenté par une matrice appelée matrice de performances.

1.3.2 Différentes approches multicritères

Différentes approches multicritères peuvent être distinguées, Pomerol et Barba-Romero [88] présentent un large panorama des différentes méthodes existantes (méthodes purement ordinales, somme pondérée, méthodes fondées sur l'utilité, méthodes de surclassement, ...) ainsi qu'une revue importante des logiciels existants. Par ailleurs, Roy [97] présentent dans le détail les méthodes de critère unique de synthèse et les méthodes surclassement, et Roy et Bouyssou [100] présentent en plus des exemples d'application.

Nous ne présenterons pas ici toutes les approches existantes, nous distinguerons simplement les approches de critère unique de synthèse et de surclassement. Tout d'abord, le critère unique de synthèse : cette approche consiste à considérer que les différents critères $g_1, \dots, g_n$ peuvent être agrégés en un critère unique $g = f(g_1, \dots, g_n)$, ce qui permet de juger les actions uniquement sur l'évaluation de ce critère unique. L'utilisation de ce type de critère conduit à une structure de préordre sur les actions de A . Il convient cependant de ne pas confondre cette approche avec une approche monocritère. En effet l'approche monocritère, même si elle permet aussi de juger les actions sur un critère unique, n'appréhende pas plusieurs dimensions de préférence.

L'autre approche, appelée approche du surclassement de synthèse [97], consiste à établir des préférence critère par critère (à l'aide de relations de surclassement monocritère) puis à agréger ces relations en une relation de surclassement global. On notera $S_j, \forall j \in \{1, \dots, n\}$ la relation de surclassement restreinte au critère j . Cette approche conduit à une structure où la relation de préférence considérée entre les actions de A n'est pas nécessairement transitive et la structure de préférence qui en découle n'est pas non plus nécessairement complète.

Ces deux approches divergent donc principalement par la prise en compte de préférence critère par critère dans l'approche du surclassement de synthèse. En effet, alors que l'approche du critère unique de synthèse se contente d'agréger directement les performances des actions, l'approche du surclassement de synthèse passe par une étape supplémentaire en édictant des préférence critères par critères. Ce sont ces préférences monocritères qui doivent ensuite être agrégées pour asseoir une comparaison entre actions.

1.3.3 Famille de critères

Dans tout problème multicritère, il convient de considérer un ensemble de critères que l'on nomme *Famille de Critères* et que l'on notera $F = \{g_1, \dots, g_n\}$ (on trouvera aussi la notation $F = \{1, \dots, n\}$). Pour que la famille F constitue une représentation appropriée des points de vue à prendre en compte dans la modélisation des préférences, Roy [97] définit la notion de *Famille Cohérente de Critères* à l'aide des trois propriétés suivantes. Ainsi une famille de critères sera dite cohérente si elle respecte :

L'Exhaustivité : Une famille F de n critères sera dite exhaustive si elle recouvre tous les aspects concourants à l'évaluation des actions. Autrement dit, si deux actions a et b sont indifférentes au sens des n critères, il ne doit pas être possible de faire apparaître des arguments permettant de préférer a à b ou b à a .

La Cohésion : Cette condition concerne la cohésion entre les évaluations critère par critère et les évaluations globales. Soient deux actions a et b indifférentes, si l'on dégrade une performance de a sur un critère pour obtenir l'action a' et que l'on améliore une performance de b sur un autre critère pour obtenir l'action b' , alors la condition de cohésion implique que b' est au moins aussi bonne que a' globalement.

La Non Redondance : Cette condition traduit l'idée qu'il ne doit pas exister de critères superflus au sein d'une famille F . Formellement, un critère g_i sera dit non redondant au sein d'une

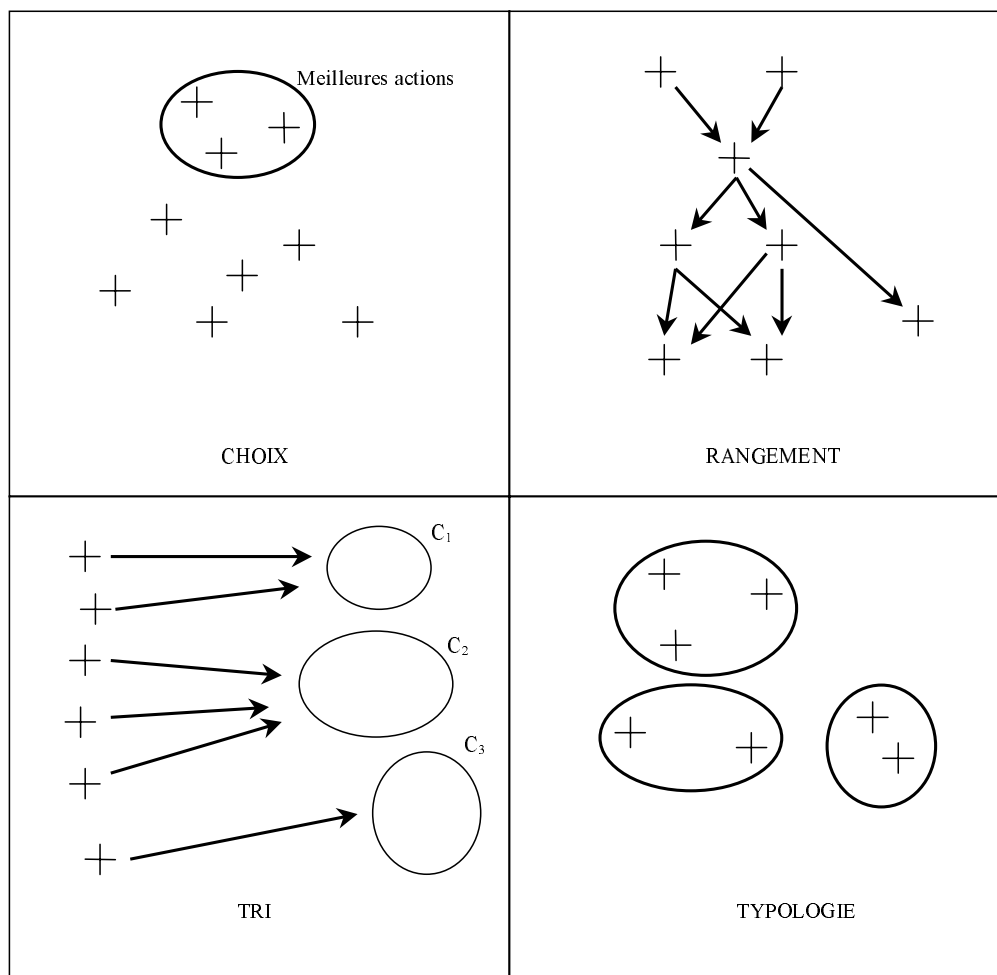


FIG. 1.1 – Les différentes problématique en A.D.

famille F si et seulement si sa suppression implique que la famille $F \setminus \{h\}$ met en défaut une des deux conditions de cohésion ou d'exhaustivité.

1.4 Problématique en aide à la décision

La première étape du processus d'aide à la décision va consister à définir vers quoi la prescription de l'homme d'étude au décideur va s'orienter. Cette étape passe par le choix d'une problématique. Roy [97] définit quatre problématiques de référence, le choix, le tri, le rangement et la description, respectivement notées, $P.\alpha$, $P.\beta$, $P.\gamma$ et $P.\delta$ (voir figure 1.1). A celles-ci il nous semble important, tout particulièrement dans le cadre de notre étude, d'en ajouter une cinquième, la problématique de la typologie.

Problématique du Choix Il s'agit de la problématique la plus classique en aide à la décision. Elle consiste à sélectionner un sous-ensemble aussi restreint que possible d'actions A' (réduit dans le cas le plus favorable à un singleton) d'un ensemble A qui justifie l'élimination des autres actions. Cette problématique généralise la problématique de la recherche opérationnelle. Elle aboutit à la mise au point d'une *procédure de sélection*.

Problématique du Tri Cette problématique consiste à affecter les actions de A à des catégories prédéfinies (caractérisées par exemple par des actions de référence). Contrairement aux autres problématiques, on ne compare pas les actions de A entre elles mais on se fonde uniquement sur les comparaisons des actions de A aux actions de référence, il s'agit d'une évaluation intrinsèque des actions. Ici, on parlera de *procédure d'affectation à des catégories*.

Problématique de Rangement Il s'agit sans doute de la problématique la plus ambitieuse. Elle a pour objectif d'ordonner les actions de A . Cependant, on ne recherche pas nécessairement un ordre complet sur les actions au sens d'une relation de préférence globale ; on cherche plutôt à regrouper les actions en classes d'équivalence, celles-ci étant totalement ou partiellement ordonnées. La procédure recherchée est une *procédure de classement*.

Problématique de la Description Cette problématique fait généralement partie de la première phase d'analyse des problématiques précédentes. Elle consiste juste à éclairer l'analyse des actions en aidant le décideur à appréhender celles-ci. Cela passe par la définition des conséquences élémentaires et des critères, ainsi que par le choix ultérieur d'une autre problématique. Si cette problématique aboutit à une procédure, on parlera de *procédure cognitive*.

Problématique de la Typologie Cette problématique est complémentaire de la problématique du tri. Elle consiste tout comme celle du tri à considérer des catégories qui correspondent à des décisions cependant l'objectif n'est pas d'affecter des actions à des catégories prédéfinies mais de regrouper les actions en sous-ensembles d'actions qui méritent de recevoir une même décision. Nous présentons cette problématique plus précisément au chapitre 5.

L'utilisation de procédures de tri peut présenter certains avantages par rapport à l'utilisation de procédures de rangement ou de choix. Le tri peut permettre dans le cas de catégories ordonnées, d'isoler la classe des meilleures actions. Cette classe possède un caractère certes moins discriminant que l'action (ou le sous ensemble d'actions A') obtenue en problématique du choix cependant l'ensemble complet des actions est ordonné à travers les catégories ce qui constitue une information plus riche que celle disponible en avec une méthode de choix.

Avec une méthode de rangement les actions sont évaluées les unes par rapport aux autres en vue d'obtenir un préordre partiel. Cependant, même si l'on parvient à obtenir un classement très fin des actions, rien ne dit que la meilleure action soit effectivement bonne. Avec une méthode de tri, de part l'évaluation intrinsèque des actions par rapport aux actions de référence, il est possible de savoir si les actions à évaluer sont effectivement bonnes. Une méthode de tri peut alors être considérée comme un bon compromis entre une méthode de choix et de rangement. Elle permet à la fois de sélectionner le sous ensemble des meilleures actions et d'ordonner au moins partiellement les actions à évaluer.

Dans la suite de ce document nous nous intéresserons exclusivement à la problématique du tri, aux problèmes de classification et à une approche fondée sur l'utilisation de relations d'indifférences. Après avoir présenté différentes approches des problèmes de classification issus de plusieurs domaines comme l'analyse de données et l'intelligence artificielle, nous donnons une définition des Méthodes Multicritères d'Affectation par Comparaison Binaire (MMACB). Nous présentons par la suite la méthode d'affectation que nous avons développée, le filtrage flou par indifférence. Enfin nous abordons le problème de la construction des différents paramètres de notre méthode et nous proposons notamment une procédure d'apprentissage des profils des catégories et une procédure interactive de construction de coefficients d'importance des critères.

Chapitre 2

Les différentes approches de la classification

Résumé

Nous présentons dans ce chapitre différents concepts en relation avec les problèmes et méthodes de classification : les notions d'apprentissage supervisé et non supervisé, les différents domaines de la classification (analyse de données, intelligence artificielle et machine learning), certaines propriétés qui permettent de classer les méthodes de classification : méthodes directes et indirectes, et explicites et implicites ainsi que la notion de classification graduelle. Par la suite, nous présentons différentes définitions que l'on peut donner à la notion de catégorie. Il convient ensuite de spécifier en quoi la classification en aide à la décision se différencie des autres domaines. En aide à la décision l'action de classer ne se fait pas uniquement en fonction des données mais aussi en fonction des préférences du décideur. On ajoute là une dimension humaine à la description supposée objective des objets à classer, le point de vue subjectif du décideur. De ce fait il peut exister autant des manières d'affecter des objets qu'il existe de décideurs. De plus, la notion de classification optimale (objective) utilisée pour évaluer un taux de mauvaise classification dans bien des domaines n'a pas de sens en aide à la décision. Les méthodes de classification sont souvent regroupées suivant trois principaux domaines qui sont l'analyse de données, l'intelligence artificielle et le domaine des machine learning. Cependant, considérant l'optique qui est la notre nous avons décidé de les regrouper non pas de part leur provenance mais en s'intéressant à leur capacité explicative. En effet, la capacité d'une méthode à pouvoir expliciter clairement les raisons des affectations effectuées revêt un intérêt particulier en aide à la décision. Nous avons donc regroupé les méthodes en méthodes explicatives et non-explicatives.

Mots clés : apprentissage supervisé et non supervisé, catégorie, état de l'art, explicatif.

2.1 Introduction

Après avoir présenté plusieurs manières de regrouper les méthodes de classification (apprentissage, domaine et propriétés) et afin de mettre en exergue les points communs et les différences fondamentales qui existent entre les approches inhérentes à ces méthodes et celle de l'aide multicritère à la décision dans le cadre de la problématique du tri, nous proposons de les regrouper différemment. Nous allons tout d'abord présenter les méthodes qui sont qualifiées de non explicatives. Ce sont celles qui fournissent un résultat de type *boîte noire* au décideur. Ensuite nous présentons les méthodes que nous qualifions d'explicatives dans le sens où elles sont susceptibles de donner une argumentation compréhensible de leur résultat à un décideur. Il convient d'ajouter que certaines méthodes peuvent être abordées selon différents points de vues, certaines méthodes statistiques sont parfois rattachées à l'intelligence artificielle. On observe de plus en plus en un décloisonnement des domaines de recherche, on peut par exemple citer [87], [62] et [84] qui abordent les relations entre les problèmes de décision et le domaine de l'intelligence artificielle.

On trouve dans Diday *et al.* [29] une présentation des méthodes statistiques, dans Weiss et Kulikowski [114] une présentation et une comparaison des méthodes de type statistique, neuronale, *machine learning* et systèmes experts et dans Michie *et al.* [69] une comparaison de ces mêmes méthodes fondée sur des résultats provenant de diverses applications. Par ailleurs, on trouve aussi dans [41] une brève présentation axée autour du *data mining* et du *data warehouse*, des arbres de décision, réseaux de neurones, méthodes bayésiennes, méthodes statistiques et méthodes de type *k*-voisins et surtout un panorama d'une vingtaine de logiciels avec les méthodes implantées ainsi que les sites <http> où certains logiciels sont téléchargeables en évaluation.

2.2 Classer les méthodes de classification

Le terme *classification* est un terme large qui peut revêtir plusieurs significations et décrire plusieurs types de problèmes. Nous allons tout d'abord aborder les liens étroits qui peuvent exister entre apprentissage et classification avec les notions d'apprentissage supervisé et non supervisé. Nous nous intéresserons ensuite aux différents domaines scientifiques dont sont issues les méthodes de classification. Enfin nous tenterons de donner une taxinomie des méthodes de classification en nous intéressant aux caractères de celles-ci.

2.2.1 Classification et apprentissage

Le terme classification est souvent associé à la notion d'apprentissage. En effet on associe souvent l'action de classer à celle d'apprendre à classer. Selon Michie *et al.* [69], qui ont un point de vue axé apprentissage, le terme classification peut avoir deux significations bien distinctes. La classification peut être l'action de regrouper en différentes catégories des objets ayant certains points communs ou faisant partie d'un même concept, on parle alors de problème de *clustering* et d'apprentissage non supervisé. Cependant la classification peut aussi correspondre à l'action d'affecter d'objets à des catégories prédéfinies, on parle dans ce cas de problème d'affectation et d'apprentissage supervisé. L'apprentissage non supervisé consiste donc à découvrir des catégories d'objets (ou *clusters*) alors que l'apprentissage supervisé consiste à mettre à jour une fonction d'affectation d'objets à des catégories et à l'utiliser en vue de d'affecter de nouveaux objets. Il convient de noter que la notion d'apprentissage n'est pas nécessairement présente dans les méthodes de classification. Nous parlerons d'apprentissage lorsque des points d'apprentissage sont utilisés pour construire soit une fonction d'affectation soit un des catégories.

Clustering et apprentissage non supervisé : ce problème consiste à regrouper des objets afin de construire des catégories. On utilise pour cela un ensemble d'exemples, appelé ensemble

d'apprentissage, constitués d'objets dont **on ne connaît pas l'appartenance aux catégories** (aspect non supervisé). On parle aussi de problème de regroupement ou de *clustering*. Dans ce type de problème les catégories sont des ensembles d'objets. Les méthodes résolvant ce type de problème peuvent permettre de construire/découvrir les catégories, elles sont à rapprocher de la problématique de la typologie. On utilisera ce type de méthode dans deux cas : tout d'abord lorsque le décideur n'est pas capable de spécifier les catégories, on pourra vouloir trouver les catégories *naturelles* que forment les points d'apprentissage. Il sera aussi possible d'utiliser de telles méthodes pour typer chaque catégorie par des points de référence.

Formellement, on peut modéliser le problème de la manière suivante :

Données :

$F = \{1, \dots, n\}$ une famille de critères,

Z un ensemble de points d'apprentissage non classés de l'espace des critères.

Objectif :

Construire $C_1, \dots, C_q$ un ensemble de q catégories.

Classification et apprentissage supervisé : ce problème consiste en la mise à jour d'une fonction d'affectation à des catégories prédéfinies. Pour cela on utilise un ensemble d'apprentissage constitué d'objets dont **l'affectation est connue** (donnée par un superviseur). Les catégories sont alors de étiquettes que l'on peut mettre sur les objets à classer. On utilisera des méthodes résolvant ce type de problème lorsque le décideur disposera d'un ensemble d'apprentissage d'objets déjà affectés. Elles seront notamment utiles pour les problèmes d'évaluation subjective [18] où le décideur appréhende mal le lien entre les données objectives et la perception subjective qui leur est associée. Il est alors difficile pour celui-ci de fixer les paramètres d'un modèle tels des seuils ou des coefficients d'importance de critères.

Données :

$F = \{1, \dots, n\}$ une famille de critères,

$C_1, \dots, C_q$ un ensemble de q catégories,

Z un ensemble de points d'apprentissage de l'espace des critères dont l'affectation est connue.

Objectif :

Construire une fonction d'affectation d'un point aux catégories en se fondant sur Z .

Classification et affectation : il s'agit de l'action de calculer l'appartenance des objets aux catégories, il n'y a pas d'apprentissage à proprement parlé, il y a juste construction d'une fonction d'affectation. Les méthodes multicritères de classification que nous proposons se placent dans cette catégorie

Données :

$F = \{1, \dots, n\}$ une famille de critères,

$C_1, \dots, C_q$ un ensemble de q catégories,

μ une fonction d'affectation.

Objectif :

Affecter un ensemble de nouveaux objets aux catégories en utilisant la fonction μ .

Le problème général de la classification multicritère d'objets consiste à affecter des objets à des catégories prédéfinies. B. Roy [97] donne la définition (8) de la classification multicritère à travers la problématique du tri :

Définition 8 (Problématique du Tri, P. β) *Elle consiste à poser le problème en terme de tri des actions par catégories, celles-ci étant conçues relativement à la suite à donner aux actions qu'elles sont destinées à recevoir, c'est-à-dire à orienter l'investigation vers une mise en évidence d'une affectation des actions de A à ces catégories en fonction de normes portant sur la valeur*

intrinsèque de ces actions et ce compte tenu du caractère révisable et/ou transitoire de A ; cette problématique prépare une forme de prescription ou de simple participation visant :

- soit à préconiser l'acceptation ou le rejet pour certaines actions, d'autres pouvant donner lieu à des recommandations plus complexes compte tenu de la conception des catégories ;*
- soit à proposer l'adoption d'une méthodologie fondée sur une procédure d'affectation à des catégories de toutes les actions convenant à une éventuelle utilisation répétitive et/ou automatisée.*

Selon la définition de Roy, cette problématique se place dans le cadre de la section intitulée classification et affectation. En effet, il n'y a pas à proprement parler d'apprentissage, on recherche uniquement à évaluer l'appartenance d'objets à des catégories. Cependant il convient de replacer cette définition dans son contexte. A l'époque peu de méthodes de classification multicritère existaient (segmentation trichotomique [70] et [71] en 1977) et cette problématique se définissait surtout en opposition aux autres problématiques : les actions ne sont pas comparées entre elles mais à des normes prédéfinies. Les catégories considérées correspondent à des décisions ou à des recommandations dont la définition doit être indépendante de l'ensemble A , c'est là le caractère intrinsèque des normes des catégories. Lorsque Roy évoque *l'acceptation ou le rejet pour certaines actions* il s'agit d'une référence directe à la procédure de segmentation trichotomique.

Nous pouvons toutefois rapprocher les méthodes qui mettent en œuvre cette problématique de la classification supervisée. En effet, si l'on cherche à définir la fonction d'affectation multicritère à l'aide d'actions déjà affectées on se place dans le cas de l'apprentissage supervisé. La méthode multicritère d'affectation peut alors être définie en deux phases, l'apprentissage et l'affectation.

En revanche les méthodes de tri ne cadrent pas avec la classification non supervisée. Celle-ci a pour objectif la construction des catégories dans ce sens elle se rapproche plus de la problématique de la typologie qui consiste à regrouper des actions entre elles. Nous aborderons des méthodes de ce type dans le chapitre 5 sur la construction des points de référence des catégories.

2.2.2 Les domaines de la classification

Les méthodes de classification sont souvent regroupées suivant trois grands domaines : l'approche statistique, l'approche neuronale et l'approche *machine learning* [69] et [114].

Les méthodes statistiques sont issues de l'analyse de données. Elles supposent l'existence d'un modèle probabiliste décrivant les données, l'objectif des méthodes est ainsi de caractériser ce modèle. On distinguera deux groupes de méthodes statistiques, les méthodes paramétriques et non paramétriques. Les méthodes paramétriques font une hypothèse sur le modèle statistique (par exemple en supposant un modèle gaussien) et vont chercher à estimer les paramètres du modèle (espérance et écart type par exemple). Au contraire les méthodes non paramétriques ne font pas d'hypothèse sur le modèle probabiliste, elles utilisent directement des séries de données exemple pour typer le modèle probabiliste recherché.

Les méthodes connexioniste (réseaux de neurones) ne font, en revanche, aucune hypothèse sur la structure des données. Ce sont des méthodes d'apprentissage à partir d'exemples, elles permettent d'effectuer une approximation d'une fonction complexe par agrégation de plusieurs fonctions simples. La fonction approchée est alors utilisée pour affecter de nouveaux objets. On reproche souvent à ces méthodes d'être de type "*boîtes noires*" car le processus de classification demeure complètement opaque et ne possède aucun caractère explicatif pour un utilisateur.

Les méthodes de type *machine learning* (apprentissage automatique) répondent au double objectif de l'apprentissage automatique à partir d'exemples et de la définition d'un nombre

minimal de règles, compréhensibles par un opérateur humain, susceptibles d'effectuer l'affectation d'un nouvel objet. Ce sont des méthodes explicatives contrairement aux méthodes précédentes.

Il convient d'ajouter que certaines méthodes peuvent être abordées selon différents points de vues, certaines méthodes statistiques sont parfois rattachées à l'intelligence artificielle. On observe de plus en plus en un décloisonnement des domaines de recherche, on peut par exemple citer [87] et [62] qui abordent les problèmes de décision et le domaine de l'intelligence artificielle.

On trouve dans Diday *et al.* [29] une présentation des méthodes statistiques, dans Weiss [114] une présentation et une comparaison des méthodes de type statistique, neuronale, *machine learning* et systèmes experts et dans Michie *et al.* [69] une comparaison de ces mêmes méthodes fondée sur des résultats provenant de diverses applications.

2.2.3 Une taxinomie des méthodes de classification

Les méthodes de classification peuvent être regroupées suivant leur domaine d'application comme à la section précédente, cependant il est aussi possible de les regrouper par leur propriétés intrinsèques. Comme nous avons vu précédemment, deux principaux types de méthodes peuvent être distinguées en considérant la dimension apprentissage : les méthodes supervisées et non supervisées. Certains auteurs étendent ces deux notions au delà de l'apprentissage et parlent alors directement de méthodes de classification supervisées et non supervisées même quand la notion d'apprentissage est absente. Les méthodes de classification non supervisée correspondent alors aux méthodes de *clustering*, c'est-à-dire au cas où les catégories ne sont pas connues et les méthodes de classification supervisées aux méthodes d'affectation à des catégories connues. Nous nous intéressons ici uniquement aux méthodes de classification supervisées.

Parmi les méthodes supervisées, nous distinguerons les méthodes explicites des méthodes implicites et au sein de celles-ci les méthodes directes des méthodes indirectes. Parallèlement, nous distinguerons aussi les méthodes graduelles des méthodes booléennes, ce point sera abordé à la section suivante.

– Méthodes explicites :

Les méthodes explicites sont celles qui utilisent une expertise précise. Les classes sont décrites uniquement à l'aide d'une expertise et non par des exemples. C'est le cas lorsque l'on utilise une base de règles et une base de faits. L'affectation est effectuée à partir d'une description formalisée par des règles généralement recueillies auprès d'un expert.

– Méthodes implicites :

Les méthodes implicites sont celles qui n'utilisent pas directement des règles de décisions provenant des experts mais des exemples d'objets déjà affectés aux classes. Il y a alors un véritable apprentissage effectué à partir de ces exemples. Ces méthodes sollicitent beaucoup moins les experts que les méthodes explicites : le problème de la définition de règles par l'expert disparaît. En effet, souvent si l'expert est capable de prendre une décision pour faire de la classification, il n'est pas nécessairement capable de décrire le processus qui lui a permis de prendre cette décision, c'est-à-dire de décomposer celui-ci afin de l'exprimer clairement. Les méthodes implicites permettent donc de contourner une partie du problème de recueil de l'expertise en s'intéressant exclusivement à des échantillons d'apprentissage et non à des règles d'experts.

– Méthodes directes :

Les méthodes directes sont celles qui vont directement utiliser les points des échantillons pour effectuer la classification et non des règles d'experts ou déduites par le système.

Les idées sous-jacentes à ces méthodes sont celles de la comparaison et de l'interpolation. L'objet à classer va être comparé à d'autres qui sont déjà classés. La qualité de l'échantillon est d'autant plus essentielle que la taille des échantillons d'apprentissage doit souvent être réduite pour éviter de trop longs temps de calcul lors des comparaisons. Ces méthodes vont en revanche permettre de manière plus aisée un apprentissage de type incrémental, c'est-à-dire si de nouvelles informations sont disponibles (comme de nouveaux points d'apprentissage) il ne va pas être nécessaire de relancer toute la procédure d'apprentissage. Ainsi le nouvel apprentissage peut être fait à partir d'un autre apprentissage effectué précédemment.

– **Méthodes indirectes :**

Ces méthodes ont pour objet de construire une description des classes à partir d'ensemble d'apprentissage. En revanche lorsque l'apprentissage est terminé les points d'apprentissage sont "oubliés" et seule la description construite lors de la phase d'apprentissage est utilisée pour l'affectation.

2.2.4 Classification graduelle ou non

Indépendamment des distinctions faites aux trois sections précédentes sur l'apprentissage, les domaines et les types de méthode, il nous semble important d'en faire une dernière : la capacité à effectuer une affectation graduelle.

De nombreuses méthodes de classification considèrent que les objets à classer appartiennent ou non aux catégories. Les degrés d'affectation aux catégories sont alors égaux à 0 ou à 1, on parle d'affectation nette. Cependant, d'autres méthodes offrent la capacité d'avoir des degrés compris dans l'intervalle $[0, 1]$, auquel cas on parlera d'affectation floue ou d'affectation graduelle (*fuzzy assignment*). Le cas le plus fréquent est de considérer que ces degrés correspondent à des probabilités, cependant ceux-ci peuvent aussi être d'autre nature comme des possibilités, des degrés de confiance etc. La méthode que nous proposerons par la suite possède cette caractéristique.

Le caractère graduel de l'affectation peut avoir plusieurs origines qui correspondent aux interprétations que l'on peut faire du degré d'affectation. Afin d'analyser le caractère flou d'une information, Dubois [31] isole deux principales causes à la mauvaise qualité d'une information. Il distingue l'information incertaine de l'information imprécise. L'information incertaine correspond à une information précise et claire mais dont un doute subsiste sur la véracité. Au contraire, l'information imprécise ne fait pas apparaître ce caractère non déterministe, l'information est certaine mais sa valeur n'est pas précise.

Par ailleurs Roy [98] et Bouyssou [15] distinguent dans le cadre de l'aide à la décision trois sources de mauvaise connaissance. La notion d'imprécision est définie comme provenant d'un défaut de mesurage. C'est l'incapacité à mesurer de manière précise. Quant à l'incertitude, elle correspond à un événement non encore réalisé que l'on cherche à estimer. Enfin, il distingue une troisième source de la mauvaise connaissance : l'indétermination. Cette imperfection de l'information peut prendre la forme l'imprécision et l'incertitude cependant elle est d'une autre nature, elle fait référence à une mauvaise perception ou connaissance de l'individu, à l'information que l'individu a du mal à définir ou à maîtriser.

Afin d'interpréter la nature du flou des degrés d'affectation, nous allons utiliser les notions d'incertitude et d'imprécision en distinguant les catégories incertaines des catégories imprécises. En revanche, nous n'utilisons pas la notion d'indétermination, qui si elle permet de définir un certain type d'information imparfaite, n'apparaît comme pertinente pour interpréter les degrés d'affectation. Par ailleurs nous distinguons un troisième type de flou qui affecte les degrés d'affectation : la notion de mélange.

Affectation incertaine aux catégories : Le degré d'appartenance graduelle peut tout d'abord

correspondre à une incertitude sur l'appartenance de l'objet à la catégorie. C'est le cas si le degré est une probabilité, une possibilité ou une nécessité. La catégorie peut être parfaitement connue mais un doute subsiste sur l'appartenance. Il faut noter que dans ce cas ce n'est pas la catégorie qui est floue mais l'information sur l'appartenance de l'objet à celle-ci.

Exemple 4 *L'appartenance à une catégorie peut être exprimée de manière probabiliste. C'est le cas lorsque l'on cherche à estimer un événement futur à travers les catégories. On peut avoir une information parfaite à instant t mais l'événement se produira à l'instant $t + \alpha$, l'information aura donc pu changer. C'est le cas par exemple lors d'une élection, si l'on cherche à savoir qui va être élu pendant la campagne électorale. Même si l'on connaît parfaitement les intentions de vote de tous les votants, on ne sait pas ce que seront effectivement les votes lors du scrutin ; les votants peuvent par exemple changer d'avis.*

Affectation à des catégories aux contours graduels : Les catégories peuvent être intrinsèquement graduelle de part leur définition, en effet la définition d'une catégorie peut être floue. Cela signifie que la transition entre l'appartenance et la non-appartenance est graduelle. Le degré d'affectation correspond ici à la plus ou moins grande appartenance à la catégorie et non à un doute sur l'appartenance comme dans le cas précédent.

Exemple 5 *C'est le cas si l'on considère la catégorie des personnes de grande taille. Dans ce cas l'information est parfaitement connue (la taille aussi précise soit elle) mais les bornes sont floues : à partir de quelle taille est on grand ? Fixer une limite n'a pas nécessairement de sens. Si l'on est sûr qu'une personne de 1,90 m. est grande et qu'une personne de 1,60 m. ne l'est pas, que dire d'une personne de 1,76 m. ? Le mieux est sans doute de considérer une transition graduelle entre les grands et les non-grands, et de définir un degré d'appartenance à la catégorie grand comme une fonction continue croissante de la taille de l'individu.*

Affectation à des catégories correspondant à un mélange : Enfin les catégories peuvent correspondre à des parties d'un mélange. On considère dans ce cas que les catégories correspondent à des "composants" des objets. Un objet apparaît alors comme un mélange des différentes catégories.

Exemple 6 *C'est le cas de l'application décrite par Foulloy et al. [39] qui consiste à reconnaître différents mélanges de cafés. Un café correspond à un mélange de différentes origines (ex. Kenya, Mozambique, Brésil, ...). Les catégories correspondent alors aux différentes provenances. Le résultat peut alors être interprété comme une estimation du taux de chaque café provenant de chaque catégorie pure.*

Confronté à une affectation graduelle, un décideur peut se sentir désarmé par ce type de résultat. En effet ce qui est attendu par celui-ci est plutôt un résultat de type non graduel permettant de prendre une décision. Il est donc légitime qu'il cherche à obtenir un résultat net à partir du résultat flou (que nous appellerons *crispation*¹ du résultat flou) en utilisant par exemple des seuils de coupe (l'affectation à la catégorie est validée si le degré d'affectation est supérieur à un seuil α).

Cependant, il n'est pas toujours nécessaire d'aboutir à une *crispation* du résultat flou d'un affectation graduelle. L'information fournie par le degré d'affectation est beaucoup plus riche qu'une simple appartenance binaire aux classes. Elle peut par exemple mettre en évidence une hésitation entre des classes. De plus dans certains cas (catégories imprécises et surtout mélange) le résultat obtenu est intrinsèquement graduel, c'est le cas de l'application sur la reconnaissance des cafés où l'information recherchée est graduelle et ne doit surtout pas aboutir à une *crispation*.

1. Cette dénomination est due à M. Pirlot.

2.3 Qu'est-ce qu'une catégorie?

Qu'est ce qu'une catégorie? Suivant les domaines et les méthodes, plusieurs réponses peuvent être données à cette question. Michie [69] à cet effet distingue trois principaux types de catégories :

1. Les catégories correspondent à des familles parfaitement **différenciées** d'objets, elles sont alors **exclusives** et l'appartenance d'un objet à l'une d'elle est **certaine**. L'appartenance à une catégorie résulte d'une évaluation globale qui ne fait pas référence aux critères. Si l'on considère des animaux, les catégories *chien* et *chat* sont de ce type.

Exemple 7 *Le problème que l'on peut se poser est de vouloir expliquer une telle affectation à l'aide de critères. On peut citer le cas d'une application dont l'objectif visait à évaluer la cuisson des biscuits Triangolini [85] et [33]. Les critères n'étaient pas explicitement données par un expert, celui-ci était juste capable d'effectuer une classification subjective. Le but de l'application était d'expliquer l'affectation de l'expert à l'aide de critères que lui même n'appréhendait pas.*

2. Les catégories correspondent à des estimations, dans ce sens elle sont très proches de la notion de variable aléatoire statistique.

Exemple 8 *Un exemple simple est l'estimation de l'évolution de taux d'intérêts en utilisant les classes augmentation et diminution.*

3. Les catégories sont définies à l'aide d'un ensemble de points exemples ou un ensemble de points de référence (prototypes ou bornes) qui correspondent directement un sous-espace des critères ou des variables. L'affectation d'un objet à une catégorie est directement fonction de ses performances. C'est cette fonction que l'on cherche à approcher.

Exemple 9 *C'est le cas si l'on cherche à définir la catégorie des voitures citadines par des voitures comme la Renault Twingo, la Smart, la Peugeot 206 et Volkswagen Polo.*

à ces définitions, nous pouvons aussi en ajouter d'autres

- Les catégories peuvent correspondre à des ensembles d'objets. C'est l'optique qui est adoptée en *clustering*, par exemple en statistique descriptive lorsque l'on cherche à regrouper des objets entre eux : une catégorie est l'ensemble des objets qui la composent.

Exemple 10 *Si l'on cherche à faire une partition $\mathcal{P}$ sur une ensemble $O = \{o_1, o_2, o_3, o_4\}$, les catégories correspondent aux parties obtenues : $\mathcal{P}_1 = \{o_1, o_2, o_3\}$ et $\mathcal{P}_2 = \{o_4\}$.*

- Une catégorie peut aussi être définie comme une région de l'espace des critères (ou des variables suivant le domaine d'application). C'est l'optique adoptée dans les systèmes à base de règles. Il s'agit là d'une approche descriptive.

Exemple 11 *C'est le cas si l'on cherche à définir les catégories C_1 et C_2 des personnes ayant de la fièvre et ceux qui n'en ont pas en considérant le critère température. C_1 à une température supérieure à $37,5^\circ\text{C}$ et C_2 à une température inférieure à $37,5^\circ\text{C}$.*

- Dans l'optique de l'aide à la décision, on peut considérer que les catégories correspondent à des traitements devant être effectuées sur les actions. Les catégories sont alors des décisions. Il s'agit ici d'une attitude constructive, les catégories ne préexistent pas à la décision, ce sont des constructions du décideur.

Exemple 12 *C'est le cas si l'on cherche à conseiller des produits financiers aux clients d'une banque, les catégories vont correspondre aux différents produits financiers.*

L'optique qui nous intéresse est celle de l'aide à la décision, nous nous plaçons donc dans une perspective constructive. Cependant ces différentes définitions des catégories nous amènent à différencier une catégorie de sa description.

Les catégories correspondent à des décisions. Ce sont des étiquettes ou des "*labels*" que l'on met sur les objets à affecter. Ces catégories peuvent être représentées de différentes manières :

- une droite : en discrimination linéaire les catégories correspondent à des parties du plan séparées par des droites.
- des règles : dans les systèmes experts les catégories peuvent être représentées par une base de règles.
- des exemples : dans les méthodes de type k plus proches voisins les catégories sont représentées par de prototypes ou des points exemples.

Par la suite dans les méthodes que nous avons développées, nous représenterons les catégories par des points de référence.

2.4 Principales différences avec l'approche multicritère de la classification

Que ce soit en analyse de données ou en intelligence artificielle, les objets que l'on cherche à affecter aux catégories sont décrits par des points dans des espaces multi-dimensionnels (ou multi-attributs). Les attributs pris en compte résultent alors le plus souvent de mesures effectuées ou directement des caractéristiques des objets. Prenons pour exemple le choix d'ordinateurs pour différentes applications comme les base de données, le calcul scientifique et le multimédia. Il s'agit d'un problème de classification où l'on cherche à affecter différents ordinateurs aux problèmes qu'ils sont les plus aptes à traiter. Les machines vont alors être décrites par des attributs comme le type et la cadence du microprocesseur, la taille du disque, la quantité de mémoire, etc. Ceux-ci peuvent être de nature qualitative ou quantitative.

L'approche multicritère va différer par la prise en compte de dimensions de préférence. En effet, dans le cadre de la classification multicritère les actions sont décrites selon des critères et non selon des attributs. L'objectif est alors, non pas de représenter un objet par une vision absolue qui se voudrait objective (l'espace multi-attribut), mais par la vision qu'en a le décideur ; c'est-à-dire en utilisant des critères. Cela signifie que pour un même objet il peut exister différentes visions qui dépendent des préférences du décideur. Ces différentes visions vont se traduire par des représentations différentes des objets. Reprenons l'exemple des ordinateurs, ce qui va être important pour un développeur en bases de données va être différent de ce qui importe pour un infographiste. Les deux vont avoir des visions totalement différentes de la même machine car ils vont utiliser des critères de choix différents.

Le passage d'un attribut à un critère va donc consister à convertir les données dont on dispose en dimensions de préférence. Chaque décideur va ainsi donner ses préférences et construire ses propres critères à partir des mêmes attributs. Cela permet l'utilisation de relations de préférence pour comparer les objets, ce qui constitue une différence avec les modèles utilisés habituellement en classification.

Ainsi, **la classification multicritère ne se fait pas uniquement en fonction des données. Elle se fait aussi en fonction des préférences d'un décideur.** Il va donc exister autant de manières de classer que de décideur. C'est pourquoi parler d'erreur de classification comme on le fait habituellement en analyse de données n'est pas nécessairement pertinent. On pourra tout

au plus parler de divergence d'appréciation entre une méthode et un décideur. De même, il va sembler difficile de comparer une méthode de classification multicritère à d'autres méthodes de classification plus traditionnelles ne faisant pas intervenir la notion de préférence d'un individu. En effet **l'objectif des méthodes de classification multicritère n'est pas de décrire au mieux des données mais de respecter un ensemble de préférences** qui auront été élicitées auparavant.

Exemple 13 (évaluation de matériel photographique) *Dans cet exemple nous présentons les différences qui peuvent exister entre une évaluation multicritère et une évaluation multiattribut. Nous allons considérer quatre types de photographes pour l'évaluation de boîtiers et d'objectifs. Chaque utilisateur va avoir ses préférences en fonction des types d'images qu'il est amené à réaliser.*

- **photographe animalier** *Il s'agit de celui qui va prendre des photos d'animaux dans leur milieu naturel, au cours de safaris par exemple. Il est amené à travailler dans des conditions souvent difficiles et dans des milieux parfois hostiles. Il a donc une grande exigence pour ce qui est de la robustesse et de la fiabilité (si un appareil est défaillant il va souvent être difficile de le réparer). Pour ce qui est des objectifs, seuls des téléobjectifs vont être utilisés c'est-à-dire des focales nécessairement supérieures à 200mm. De plus, l'ouverture maximale des objectifs va être primordiale car la majeure partie des images est prise soit au lever soit au coucher du soleil.*
- **photographe de reportage** *Il s'agit du photographe qui va prendre des images sur le vif. Son souci principal va être la maniabilité et la discrétion, il va donc s'orienter vers du matériel léger et maniable. Les objectifs utilisés vont être des grands angles et des télé courts, les zooms (de type 20-35 mm et 70-200 mm) et l'autofocus (AF) étant évidemment des plus indéniables. En revanche une attention moins importante sera accordée à la qualité optique et aux caractéristiques techniques des appareils.*
- **photographe de studio** *Au contraire, le photographe de studio va apporter une attention particulière à la technique alors que la maniabilité du matériel ne va pas intervenir dans ses choix. Les objectifs utilisés doivent être dotés d'une qualité optique irréprochable, seuls des focales fixes sont utilisées. L'AF, quant à lui, ne présente quasiment aucun intérêt, enfin le petit format 24×36 va souvent être délaissé au détriment du moyen format ($6 \times 4,5$, 6×6 , $6 \times 7 \dots$).*
- **touriste en voyage** *Il s'agit du seul photographe amateur, celui-ci va avoir des exigences totalement autres, que l'on peut résumer par photographier léger et facile pour le moindre coût. On va donc se diriger vers des zooms de faible ouverture alliant un prix raisonnable à un faible encombrement.*

Voyons maintenant la construction de deux critères à partir de descriptions multiattribut, le critère performance d'un boîtier et le critère distance focale d'un objectif. Pour ce qui est du critère performance, il va résulter de l'agrégat des différents caractères techniques des appareils.

Pour le photographe animalier, les points importants vont être la robustesse du boîtier, la présence et la qualité de l'AF, le temps de pose minimum et la qualité de l'exposition. En revanche des aspects comme la maniabilité de l'appareil vont beaucoup moins importer. Le reporter, quant à lui, va s'intéresser à des boîtiers légers et discrets munis d'un AF, le but premier étant de pouvoir prendre des images le plus rapidement possible. Le photographe de studio, va utiliser, pour construire son critère performance, les possibilités au flash des boîtiers, la gamme d'objectifs disponibles et la qualité de l'obturateur. En revanche, la présence de l'AF, la masse ou la présence d'un flash intégré ne sont quasiment pas pris en compte. Quant à notre touriste, il va principalement être limité par son budget et par ses vertèbres cervicales, ce qui lui importe en fait va être de disposer d'un boîtier facile d'emploi et léger. La construction de ce critère montre bien l'hétérogénéité des points de vue

qui peuvent exister entre différents décideurs. Ce qui se trouve derrière la notion de performance d'un boîtier est totalement différent suivant ce que l'on veut en faire. Le critère va être dépendant à la fois de l'objet considéré et du décideur.

Prenons maintenant le critère distance focale, le problème est comment passe-t-on de l'attribut distance focale au critère? On dispose d'un attribut qui résulte d'une mesure en millimètres sur une échelle qui va de 14 à 1200 mm. Si l'on considère cet attribut indépendamment de tout décideur il va s'agir d'une mesure sur laquelle on ne peut exprimer aucune préférence. L'attribut va devenir un critère à partir du moment où l'on va plaquer dessus une dimension de préférence, c'est-à-dire par exemple dire que l'on préfère les longues focales aux courtes focales. Cela va être le cas du photographe animalier, en dessous de 200 mm les objectifs sont quasiment équivalents (le décideur va être indifférent entre un 20 mm et un 50 mm) car ils ne permettent pas de prendre de clichés d'animaux. Sur ce même attribut, chaque utilisateur va pouvoir plaquer son échelle de préférence.

2.5 Procédures multicritères d'affectation

Parmi les 4 problématiques définies par Roy [97], description, choix, rangement et tri, ces sont les trois première qui ont données lieu au plus grand nombre de méthodes développées. Le champs de la classification (ou du tri) a, en comparaison, été relativement peu exploré. On peut cependant citer les travaux effectués par Moscarola et Roy [71] autour de la procédure trichotomique de segmentation, ceux de Ostanello et Massaglia [67] avec la méthode N-tomic et ceux de Yu Wei [113] pour la mise au point de Electre Tri.

2.5.1 Procédure Trichotomique de segmentation

Les premières études sur le Tri sont principalement dues travaux de Moscarola et Roy [71], [70] et [96] avec la mise au point d'une procédure trichotomique de segmentation. Cette procédure a pour objectif d'effectuer une recommandation sur la base de l'affectation des actions à trois catégories. Les actions ayant des raisons suffisamment importante pour être recommandés à un décideur sont affectées à la catégorie C_1 , celles qui possèdent des raisons suffisamment importantes pour ne pas être recommandés à un décideur sont affectées à C_3 enfin celles pour lesquelles aucune décision ne semble solidement établie sont affectées à C_2 . Cette procédure se présente en fait comme une procédure de choix dans laquelle on introduit un niveau intermédiaire (C_2) entre les bons et les mauvais (C_1 et C_3) à la manière d'une logique trivalente.

Le contenu sémantique associé à la catégorie C_2 peut être soit la mauvaise connaissance soit la mauvaise détermination des performances. Dans le cas de la sélection de dossiers pour l'octroi de crédits où C_1 correspond à la catégorie des bon dossiers et C_3 à celle des mauvais dossier, C_2 peut être interprétée comme la catégorie contenant les différents types de client suivants.

- Les clients dont les performances sont mal connues. Il s'agit d'informations que le décideur (le banquier) ne parvient pas à connaître ou qu'il connaît mal du fait par exemple de rétentions d'information.
- Les clients dont les performances sont soumises à un aléa. Il s'agit de performances qui ne sont pas certaines, qui possèdent par exemple une distribution de probabilité ou qui sont le résultat d'une évaluation (par exemple un investissement non arrivé à terme qui peut engendrer soit des gains soit des pertes). L'information n'est pas connue, elle est évaluée, elle sera connue par la suite.
- Les clients dont les performances sont moyennes. Il s'agit des clients dont les performances sont certaines et connues mais qui sans être mauvais ne parviennent pas à être bons. Leur acceptation dépendra du nombre de clients affectés à C_1 .

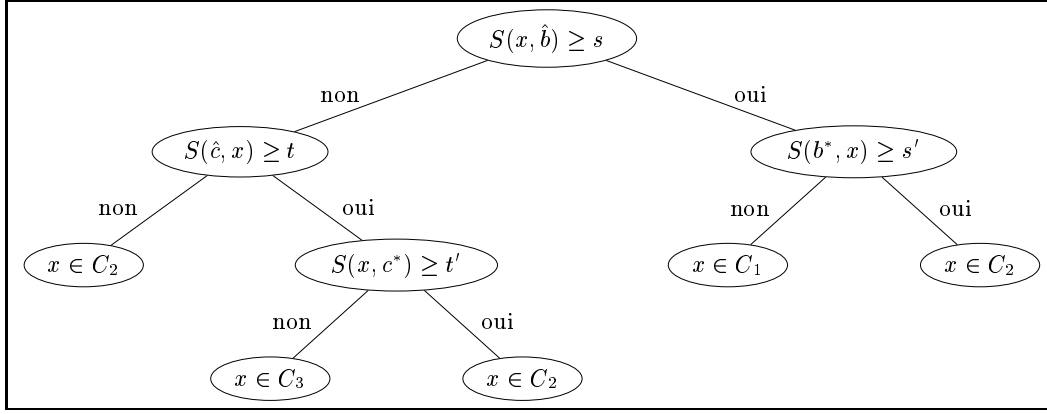


FIG. 2.1 – Arbre de décision de la procédure trichotomique de segmentation

La procédure d'affectation fait appel à une relation de surclassement flou, le surclassement de x sur y est alors évalué par une fonction à valeur dans $[0, 1]$ où 1 exprime un surclassement certain et 0 une absence totale de surclassement. Cette relation floue est évaluée de la même manière que dans Electre III [95], un résultat net est par la suite obtenu en utilisant une α -coupe. Le niveau de la coupe étant fixé par le décideur. Cela pose le problème de la valeur du seuil de coupe, à partir de quelle valeur le surclassement pourra effectivement être validé? D'une part le décideur n'est pas forcément à même d'appréhender une telle valeur et d'autre part la fixation d'un tel seuil est toujours arbitraire. Le surclassement établi sur la base d'un degré de $\alpha + \epsilon$ est il significatif par rapport à celui non établi avec le degré $\alpha - \epsilon$ (avec ϵ aussi petit que l'on veut).

Formellement, les catégories sont représentées par deux familles de profils limites $B = \{b_1, \dots, b_l\}$ et $C = \{c_1, \dots, c_k\}$. Ceux-ci correspondent respectivement aux bornes inférieures de C_1 et supérieures de C_3 . L'affectation d'une action (représentée par le point x) se fonde sur les degrés de surclassement avec les profils de B et C . On note :

$$\begin{aligned}
 S(x, \hat{b}) &= \max_{i=1}^l S(x, b_i) \\
 S(b^*, x) &= \max_{i=1}^l \{S(b_i, x) / b_i \neq \hat{b}\} \\
 S(\hat{c}, x) &= \max_{i=1}^k S(c_i, x) \\
 S(x, c^*) &= \max_{i=1}^k \{S(x, c_i) / c_i \neq \hat{c}\}
 \end{aligned}$$

Les actions à classer sont alors affectées suivant l'arbre de décision de la figure 2.1.

Cependant l'utilisation de catégories multiprofiles peut conduire à des situations difficilement compréhensibles pour le décideur du fait de la non-transitivité de la relation de surclassement. Enfin, un effet Condorcet (voir p. 11) peuvent être aisément mis en évidence même dans un cas monoprofil.

2.5.2 nTOMIC

Cette méthode, due à Ostanello et Massaglia [67], permet de classer des actions suivant des catégories ordonnée. Elle se fonde sur l'utilisation de deux profils fictifs, $b = \{b_1, \dots, b_n\}$ et $c = \{c_1, \dots, c_n\}$. La performance b_j correspond au niveau à partir duquel une action est considérée

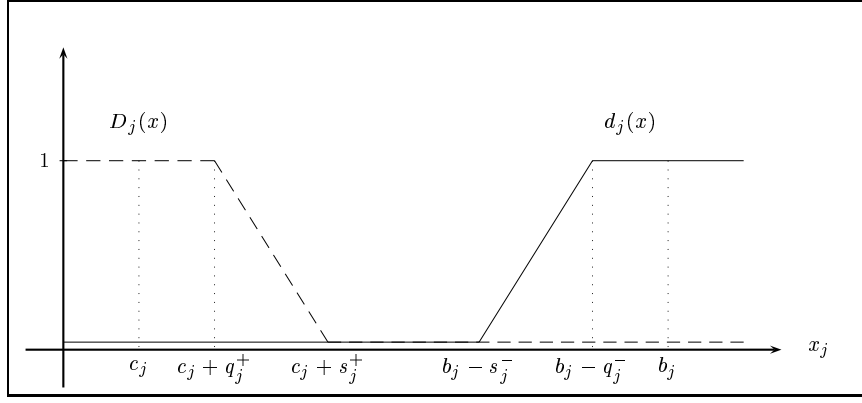


FIG. 2.2 – Fonction de "goodness" et de "badness" dans *nTOMIC*

comme bonne de manière certaine du point de vue du critère j et c_j au niveau en dessous duquel une action est considérée comme mauvaise. Il faut noter que b et c n'ont pas forcément de significations réelles en terme d'action, ce sont des actions fictives qui correspondent respectivement à une action jugée comme bonne et comme mauvaise sur tous les critères simultanément. Pour faire face à des phénomènes de mauvaise détermination et d'incertitude des performances des seuils de discrimination, s , et d'indifférence, q , sont introduits dans le modèle tels que $s \geq q \geq 0$. à chaque performance de b_j , respectivement c_j , sont donc associés un seuil d'indifférence q_j^- , respectivement q_j^+ , et un seuil de discrimination s_j^- , respectivement s_j^+ .

À partir de ces valeurs, sont définis deux sous-ensembles flous de critères : le sous-ensemble flou des critères qui confirment qu'une action est *bonne* et le sous-ensemble de ceux qui confirment qu'une action est *mauvaise*. L'appartenance du critère j à ces sous-ensembles est défini à l'aide des fonctions $d_j(x)$ (indice de *goodness*) et $D_j(x)$ (indice de *badness*), respectivement pour les sous-ensembles *bon* et *mauvais*, voir figure (2.2).

Deux mécanismes d'agrégation sont proposés pour obtenir les indices globaux de *badness*, $D(c, x)$, et de *goodness*, $d(x, b)$, un compensatoire et un non-compensatoire en utilisant le formule de l'indice de crédibilité de Electre III. Contrairement à la segmentation trichotomique, les catégories ne sont pas définies explicitement par des profils mais par une partition du plan (d, D) à l'aide de seuils. La fonction d'affectation se fonde donc uniquement sur les indices $D(c, x)$ et $d(x, b)$ et sur des seuils, voir figure (2.3).

2.5.3 Electre Tri

La méthode Electre Tri [113] est une méthode multicritère d'affectation à des catégories complètement ordonnées. Elle se fonde sur une représentation de chacune d'entre elles par un profil limite supérieur et inférieur. La fonction d'affectation prend appui sur la relation de surclassement flou de Electre III à travers deux procédures de filtrage, les procédures pessimiste et optimistes. L'utilisation de ces deux procédures permet de gérer les situations d'incomparabilité.

Formellement on a q catégories $C_1, \dots, C_q$ ordonnées de la moins bonne à la meilleure. Celles-ci sont définies par les profils limites $b_0, \dots, b_q$ où b_{j-1} correspond au profil inférieur de C_j et b_j au profil supérieur. La procédure d'affectation pessimiste consiste à comparer une action à classer aux profils des catégories en partant du meilleure, b_q , et à affecter dès que l'on trouve un profil b_i qui est préféré à l'action. L'action est alors affecté à C_{i+1} . Tout au contraire la procédure optimiste consiste à comparer l'action aux profils mais en partant du plus mauvais et à affecter dès que l'on en trouve un qui est préféré à cette, on affecte alors à C_i voir figure (2.4).

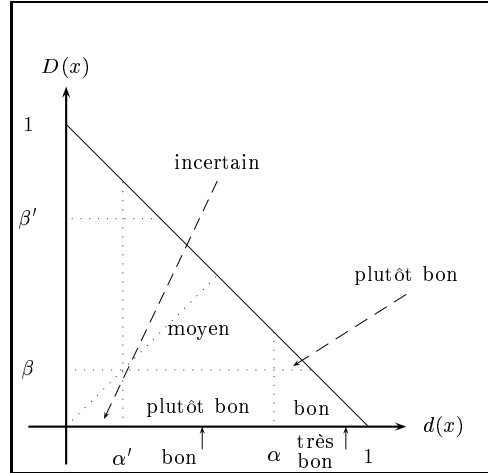


FIG. 2.3 – *Partition de l'espace (d, D)*

procédure pessimiste	procédure optimiste
$i = k + 1$ RÉPÉTER $i - -$ évaluer x S b_i JUSQU'À x S b_i affecter x à C_{i+1}	$i = -1$ RÉPÉTER $i + +$ évaluer b_i S x JUSQU'À b_i S x affecter x à C_i

FIG. 2.4 – *Procédures d'affectation de Electre Tri*

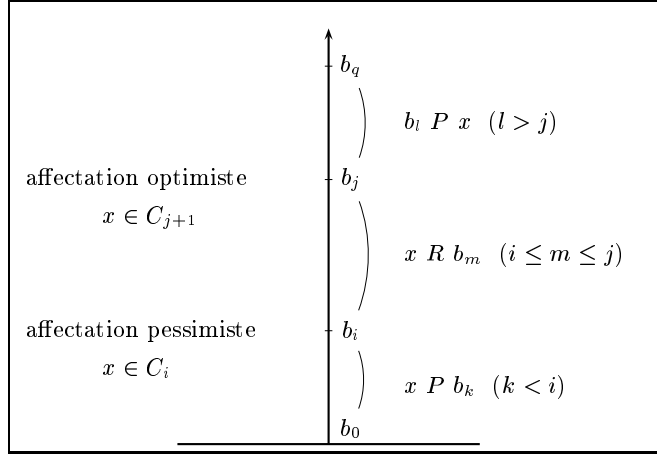


FIG. 2.5 – *Affectation optimiste et pessimiste de Electre Tri*

L'utilisation de ces procédures permet de choisir une attitude vis à vis de l'incomparabilité. Que faire lorsque une action est incomparable à plusieurs bornes successives $b_i, \dots, b_j$ (avec $i < j$) tout en étant meilleure que les bornes b_k ($k < i$) et moins bonne que les bornes b_l ($l > j$)? Il convient de choisir une catégorie entre C_{i+1} et C_j . L'attitude optimiste (logique disjonctive) consiste à affecter à la meilleure des catégories dont la borne inférieure est incomparable alors que l'attitude pessimiste (logique conjonctive) va affecter à la plus mauvaise catégorie dont la borne supérieure est incomparable, voir figure 2.5.

2.5.4 Filtrage Flou par Préférence (FFP)

Cette méthode de classification multicritère est due à Perny [81]. Il s'agit d'une méthode de filtrage, qui, contrairement à ELECTRE Tri, utilise une relation de préférence floue. Les catégories, $C_1, \dots, C_q$ sont supposées ordonnées de manière décroissante et chaque catégorie est représentée par deux frontières supérieure et inférieure. On suppose que la frontière inférieure de C_t est aussi la frontière supérieure de C_{t-1} pour $t \in \{2, \dots, q\}$. Ainsi, les q catégories sont représentées par $q + 1$ frontières $Y_0, y_1, \dots, Y_q$ où Y_{t-1} et Y_t sont respectivement les frontières supérieures et inférieures de C_t . La règle d'affectation d'une action à une catégorie peut alors être exprimée par *a est affecté à C_t si et seulement si a est préféré à au moins un élément de Y_t sans pour autant être préféré à aucun élément de Y_{k-1}* . Ce principe se traduit par l'équation logique suivante :

$$(a \in C_t) \equiv (\exists y \in Y_k, a \succ y) \wedge (\forall y' \in Y_{k-1}, \neg(a \succ y')) \quad (2.1)$$

Si $\succ$ est une relation de préférence floue alors l'équation 2.1 se traduit par l'équation 2.2 pour donner le degré d'affectation :

$$\mu_t(a) = T(\succ(a, Y_k), N(\succ(a, Y_{k-1}))) \quad (2.2)$$

avec T une t-norme, N un opérateur de négation et $\succ(a, Y_t) = \sup_{y \in Y_t} \{\succ(a, y)\}$ pour tout $a \in A$ et $t \in \{1, \dots, q\}$

Exemple 14

$$\mu_t(a) = \min(\succ(a, Y_t), 1 - \succ(a, Y_{t-1})) \quad (2.3)$$

Par ailleurs, afin que les catégories soient *bien formées* (*well designed categories*), Perny [81] spécifie les trois propriétés suivantes qui se doivent d'être respectées.

Catégories Totalelement Ordonnées Afin que les catégories soient complètement ordonnées, les points de référence se doivent de respecter :

$$\forall y \in Y_{k-1}, \forall y' \in Y_k, \forall j \in \{1, \dots, n\}, g_j(y) \geq \quad (2.4)$$

2.6 Méthodes non explicatives

Si l'approche de l'aide à la décision consiste, pour l'homme d'étude, à proposer des solutions au décideur en vue de le laisser effectuer un choix, elle a aussi pour objectif d'aider le décideur à justifier ses choix. Il est souvent aussi important, si ce n'est plus, de pouvoir justifier une décision que de pouvoir l'obtenir.

Comme le disent Roy et Bouyssou [100], *une "bonne" aide à la décision est une aide qui conduit à de "bonnes" décisions, voire à des décisions "optimales". Cependant, la reconnaissance de la complexité des processus de décision, des rationalités multiples qui s'y côtoient, des conflits qui s'y déroulent et des transformations qui s'y opèrent ne peuvent que faire douter de la possibilité de toujours prouver qu'une décision est, ou non, la meilleure : meilleure pour qui ?, selon quels critères ?, à quel moment dans les temps ?, etc. [...] Aucune démarche objective fondée sur la seule raison ne peut démontrer l'optimalité ni même le bien-fondé d'un système de valeurs ou d'un mode d'anticipation de l'avenir.*

Autrement dit, en admettant qu'une décision optimale puisse exister, ce qui est déjà une hypothèse forte, le caractère optimale de cette décision n'est pas forcément démontrable. De plus, l'optimalité par rapport à un modèle et à une méthode d'une décision ne correspond pas nécessairement à l'optimalité de cette décision par rapport à la réalité. C'est pourquoi, on s'intéresse, en aide à la décision, à la capacité d'une méthode à expliquer ses résultats.

Nous avons regroupé dans cette partie un ensemble de méthodes qui fonctionnent par optimisation. Pour chacune d'entre elles un critère à optimiser est utilisé, par exemple une somme de distances à minimiser, une densité de probabilité conditionnelle à maximiser, une fonction d'erreur à minimiser, etc. Les méthodes présentées dans cette section justifient leur résultat dans le sens où le critère pris en compte est effectivement optimisé. Cependant nous les qualifions de non explicatives car elles ne permettent pas (ou mal) à un décideur de justifier explicitement la décision, c'est-à-dire de l'expliquer par des termes autres que mathématiques.

2.6.1 Méthodes fondées sur l'utilisation de distance

Discrimination linéaire

Ces méthodes sont le fruit de travaux de Fisher [36]. L'objectif est de séparer les points d'apprentissage des différentes catégories par des hyperplans. La construction de ces hyperplans peut être effectuée par des méthodes comme la méthode des moindres carrés et la méthode du maximum de vraisemblance, voir [69] p.17-21. L'affectation de nouveaux points est alors effectuée en fonction de leur position par rapport aux hyperplans, voir figure 2.6. Les hyperplans sont construits de manière à minimiser la dispersion des points d'une même catégorie autour du centre de gravité de celle-ci. L'utilisation d'une distance est alors nécessaire pour mesurer cette dispersion.

Il s'agit de méthodes très aisées à mettre en œuvre ; elles présentent l'avantage de pouvoir traiter des données de très grande taille. Cependant, si la discrimination de deux catégories est un cas très favorable, la discrimination d'un nombre supérieur de classes nécessite quelques adaptations.

Les classes non linéairement séparables ne sont pas reconnues (XOR en 2d). Plusieurs séries d'essais, [6] et [24], sur divers jeux de données (ayant des caractères aussi diverses que des données plus ou moins bruitées, des critères aussi bien qualitatifs que quantitatifs, des catégories peu homogènes ...) montrent que d'autres méthodes, souvent plus lourdes à mettre en œuvre comme les méthodes de type neuronal, sont beaucoup plus performantes sur ce type de problème. De même les cas des classes non homogènes et multi-profil posent des problèmes de reconnaissance. Enfin, du fait du principe de la séparation des catégories, les catégories qui ne forment pas une partition ne peuvent être envisagées.

D'autre part, Diday *et al.* [29] montre que cette méthode est équivalente à la règle de décision bayésienne que nous abordons à la section 2.6.2 dans le cas d'égalité des probabilités a priori, d'égalité des coûts, de normalité des densités de probabilité conditionnelle et d'égalité de dispersion des classes.

Discrimination quadratique

Cette méthode diffère de la discrimination linéaire par l'utilisation d'une métrique par catégorie pour mesurer la dispersion de chaque classe. Les classes sont ici regroupées en ellipses et non séparées par des hyperplans. La règle de décision en est plus complexe à calculer, la méthode est donc plus lourde à mettre en œuvre. Les phénomènes de dispersion sont plus importants, il convient alors d'utiliser des échantillons de taille d'autant plus grande que le nombre de variable est important.

Les classes ne formant pas une partition et multiprofiles ne sont pas explicitement gérées. Le problème du à l'utilisation de distances est d'autant plus important que le nombre de distances à considérer est grand.

Discrimination logistique

Cette approche est une extension du modèle de la discrimination linéaire présenté au paragraphe précédent. Elle vise à mettre en relation la description d'un individu, $x = (x_1, \dots, x_n)$ avec la probabilité conditionnelle de la catégorie C_i sachant la description de x , $\pi(x) = Pr(C_i/x)$. Tout comme la discrimination linéaire elle cherche à séparer les catégories par des hyperplans. Cependant, alors que la discrimination de Fisher optimise le carré d'une différence de coûts, la discrimination logistique consiste à maximiser une vraisemblance conditionnelle. Les résultats obtenus sont très proches de ceux obtenus par discrimination linéaire.

2.6.2 Méthodes fondées sur des distributions de probabilité

Règle de décision de Bayes d'erreur minimale

La règle de décision de Bayes consiste à affecter un objet à la catégorie qui minimise le taux moyen d'erreur de classification. Cela revient à affecter un objet o à C_i si $P(C_i/x)$ est maximum avec x l'image de l'objet dans l'espace des variables. La pratique montre que cette probabilité est difficilement évaluable. L'utilisation du théorème de Bayes permet de pallier cette difficulté :

$$P(C_i/x) = \frac{P(C_i) \times f(x/C_i)}{f(x)}$$

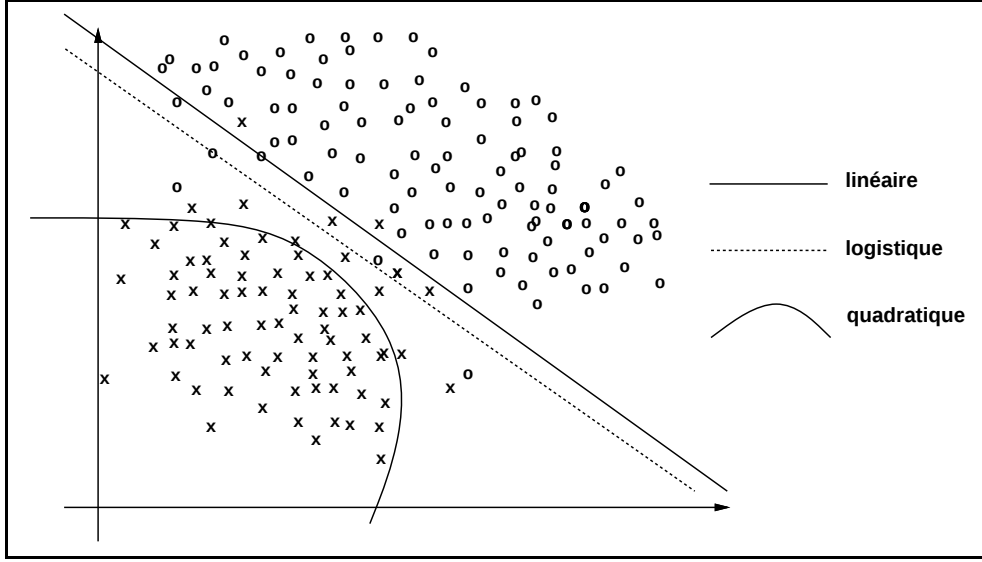


FIG. 2.6 – Frontières construites par les méthodes de discrimination linéaire, logistique et quadratique

avec $f(x)$ la densité de la loi de probabilité de x et $f(x/C_i)$ la densité de probabilité conditionnelle de x . $f(x)$ étant indépendant de C_i , la règle de décision de Bayes revient à maximiser $P(C_i) \times f(x/C_i)$. La méthode suppose donc la connaissance de la loi de probabilité conditionnelle de x . Celle-ci n'est pas forcément connue, elle devra donc être évaluée à l'aide d'une méthode d'estimation, nous présentons dans les paragraphes suivants des méthodes d'estimation paramétriques ou non paramétriques (voir section 2.7.2).

Méthodes paramétriques

Ces méthodes reprennent la règle de décision de Bayes d'erreur minimale. On suppose ici connue la forme générale de la densité conditionnelle de probabilité des objets à leur classe, $f(x/C_i)$. Dans bien des cas, la cette fonction est supposée normale. La forme générale de cette loi, $f(x, \theta)$ est paramétrée par le vecteur $\theta = \{\theta_1, \dots, \theta_q\}$. Le but de la méthode est de trouver le vecteur θ qui va maximiser l'appartenance des points exemples à leur classe. Saporta [102] utilise la méthode du maximum de vraisemblance, on maximise alors la densité de probabilité d'apparition des exemples de l'ensemble d'apprentissage, $f(z_1, \dots, z_{|Z|}, \theta)$.

2.6.3 Méthodes neuronales

Neurone Formel

Les méthodes neuronales ont pour ambition la reproduction des fonctions de raisonnement humain. À leur origine on trouve le modèle du neurone formel de McCulloch et Pitts [68]. Il s'agit d'un modèle à seuil dont la valeur de sortie binaire, y_k est calculée en fonction de la somme pondérée des valeurs d'entrée, $\{x_1, \dots, x_j\}$ et d'un seuil fixé μ_k , formellement on a pour un neurone k :

$$y_k = \begin{cases} 1 & \text{si } \sum_{i=1}^j \omega_{ki} x_i \geq \mu_k \\ 0 & \text{sinon} \end{cases}$$

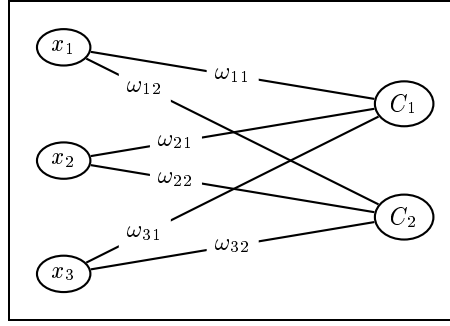


FIG. 2.7 – *Modèle à deux couches*

les paramètres ω_{ki} sont appelés les poids du réseau. De tels réseaux ne permettent de modéliser que des problèmes où les classes sont linéairement séparables.

Le modèle du neurone de McCulloch et Pitts a, par la suite, été généralisé par

$$y_k = f_k \left(\sum_{i=1}^j \omega_{ki} x_i \right)$$

où f_k est appelée fonction d'activation, on utilise souvent une fonction sigmoïde [74].

à l'aide de neurones de ce type il est possible de construire des réseaux simples composés d'une couche d'entrée et d'une couche de sortie. Ces réseaux sont composés d'autant de neurones d'entrée que de critères sont nécessaires et d'autant de neurones de sortie que de classes sont présentes, voir figure 2.7 pour un problème de classification en deux classes avec des objets décrits sur trois critères.

Le Perceptron Multi Couches

Le Perceptron Multi Couches (P.M.C.) est un des réseaux de neurones les plus répandus. Ces réseaux utilisent généralement des fonctions de transfert non linéaires et se différencient des réseaux à deux couches par l'adjonction d'une ou plusieurs cachée(s). Les couches cachées sont des ensembles de neurones qui s'intercalent entre la couche d'entrée et la couche sortie. Le calcul des coefficients de sortie se fait alors par propagation de l'information donnée aux neurones d'entrée à travers le réseau. Il en résulte que le réseau est assimilable à une fonction non linéaire des paramètres d'entrée.

Dans l'exemple de la figure 2.8, le réseau est composé de trois couches. La couche d'entrée qui est instanciée à l'aide des valeurs des performances des points (d'apprentissage et de test), une couche cachée composée d'un nombre de neurones fixé par l'utilisateur et la couche de sortie donnant le coefficient d'appartenance du point aux deux catégories C_1 et C_2 .

Ce réseau peut être vu comme une méthode statistique non paramétrique. Kabrisky *et al.* [55] démontrent à cet effet la convergence statistique de la méthode vers la règle de décision optimale de Bayes (aussi bien dans le cas en deux classes que dans le cas général) lorsque l'algorithme de rétropropagation de gradient [101] est utilisé.

Algorithme d'apprentissage

Deux principaux problèmes se posent à la construction d'un réseau de neurones : la définition du nombre et de la taille des couches cachées, et la définition des poids. L'algorithme de *rétropro-*

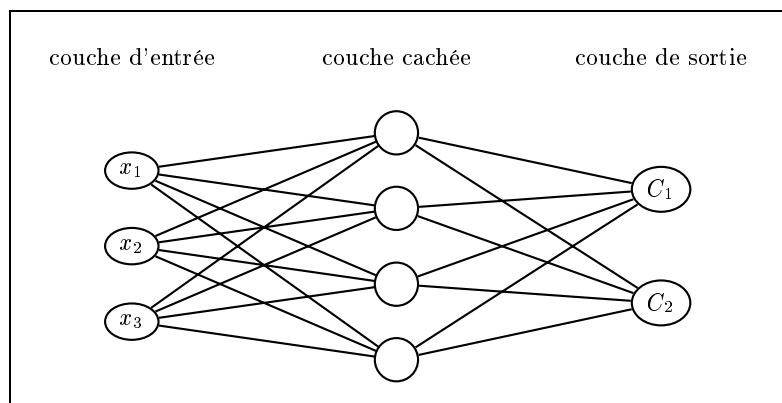


FIG. 2.8 – *Perceptron Multi Couches à une couche cachée*

pagation de gradient [101] permet de trouver de manière optimale les poids d'un réseau. Il s'agit d'un algorithme d'apprentissage supervisé.

Il suppose un réseau dont la structure est fixée (connexions, nombres de couches et nombre de neurones par couche). L'objectif de l'algorithme va être de minimiser une fonction d'erreur entre l'affectation calculée des points d'apprentissage et leur appartenance aux classes, l'erreur quadratique est la plus souvent utilisée.

Les réseaux de neurones sont des méthodes de classification qui doivent impérativement être utilisés en deux phases, une phase d'apprentissage supervisé puis la phase d'affectation. Il s'agit de méthodes de type *boite noire* présentant peu de caractère explicatif. En revanche du fait de la capacité à approcher une fonction d'affectation non linéaire elles permettent d'appréhender des problèmes de forte complexité. Du point de vue du temps de calcul, la phase d'apprentissage peut être très longue ; au contraire, la phase d'affectation peut être effectuée en temps réel.

Un des problèmes souvent évoqués à propos de la phase d'apprentissage des réseaux de neurones est celui du *sur-apprentissage*. Le *sur-apprentissage* intervient lorsque le réseau parvient à classer de manière optimale tous les points d'apprentissage en ayant effectué une trop grande spécialisation. Cela signifie qu'il reconnaît parfaitement tous les points de Z mais uniquement ceux là. Il perd alors toute sa capacité de généralisation, on parle aussi d'apprentissage par cœur.

2.7 Méthodes explicatives

2.7.1 Réseaux bayésiens

La première formulation de ce type de représentation est due à Wright [115] pour l'analyse des problèmes de récoltes agricoles (crop failure). Ses représentations ont été reprises par la suite sous d'autres noms, *causal network* (Cooper [20]), *belief network* (Pearl [78]) et *influence diagram* (Howard et Matheson [52]). On peut citer Heckerman [46] pour une présentation concise ainsi que quelques exemples d'application.

Avant de donner la définition d'un réseau bayésien, introduisons quelques notations. Soit $G = (V, E)$ un graphe orienté sans cycle, à chaque nœud est associé un ensemble fini d'états Ω_v . L'ensemble de toutes les configurations possibles est noté : $\Omega = \times_{v \in V} \Omega_v$. Nous disposons d'une loi de probabilité $P(V)$ sur Ω définie telle que : $P(V) = P\{X_v = x_v, v \in V\}$. On a alors la définition

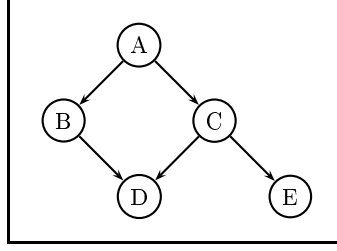


FIG. 2.9 – Graphe de diagnostic de cancer

suivante :

Définition 9 (Réseaux Bayésiens) Soit un graphe orienté sans cycles $G = (V, E)$. Pour tout $v \in V$, $c(v) \subseteq V$ est l'ensemble de tous les prédécesseurs (ou parents) de v et $d(v) \subseteq V$ l'ensemble de tous les successeurs (ou descendants) de v . Pour tout $v \in V$, soit $a(v)$ l'ensemble des variables de V excluant v et ses successeurs. Si pour tout $W \in a(v)$, W et v sont conditionnellement indépendants par rapport à $c(v)$, $P(W|c(v)) = P(W|c(v)) \times P(v|c(v))$, alors $C = (V, E, P)$ est un réseau bayésien.

Les nœuds vont représenter des variables de décision pouvant se trouver dans différents états. Les arcs vont représenter les liens de cause à effet qui existent entre les variables. La distribution de probabilité conditionnelle sert à évaluer la probabilité de l'état d'une variable en fonction de l'état de ses prédécesseurs. Dès que l'état des variables est connu, les probabilités conditionnelles permettent de propager l'information jusqu'au nœud décision.

Exemple 15 Voyons un exemple médical simple de réseau bayésien du à Cooper [20]. Supposons qu'un cancer métastaté (A) puisse être à l'origine d'une augmentation du taux de calcium dans le sang (B) et aussi d'une tumeur cérébrale (C). (B) et (C) peuvent eux même être à l'origine d'un coma (D). De plus (C) peut aussi provoquer un œdème des papilles (E). Associons à chaque nœud les ensembles d'états suivants : $\Omega_A = \{a_1, a_2\}$, $\Omega_B = \{b_1, b_2\}$, $\Omega_C = \{c_1, c_2\}$, $\Omega_D = \{d_1, d_2\}$ et $\Omega_E = \{e_1, e_2\}$ avec les significations suivantes :

a_1	cancer métastaté	a_2	pas de cancer métastaté
b_1	augmentation du taux de calcium	b_2	pas d'augmentation
c_1	tumeur cérébral	c_2	pas de tumeur
d_1	coma	d_2	pas de coma
e_1	œdème des papilles	e_2	pas d'œdème

On peut alors représenter les données connues par le graphe figure 2.9.

Voyons comment utiliser un réseau bayésien en vue de faire de la classification. Nous allons distinguer deux types de variable, celles qui correspondent aux critères, $X = \{x_1, \dots, x_n\}$, et la variable de sortie qui correspond à l'affectation A . Les différents états que A peut prendre correspondent aux différentes catégories. La structure du réseau doit être évaluée à l'aide d'un ensemble d'apprentissage. Pour effectuer une affectation d'un nouvel objet, il suffit alors d'instancier les variables de X avec les performances du point puis de propager l'information jusqu'au nœud de sortie A . Les probabilités des états de A donnent le degré d'affectation aux différentes catégories, il y a donc la possibilité d'effectuer une affectation floue aux catégories. On peut noter que dans

un cadre d'aide à la décision les nœuds peuvent aussi modéliser les conséquences élémentaires des critères.

Il nous semble important d'ajouter que la connaissance de toutes les performances n'est pas nécessaire à la propagation de l'information. Ceci fait des réseaux bayésiens une méthode très attractive dans l'hypothèse de performances manquantes et d'informations incomplètes. En revanche du fait de la nature des coefficients d'appartenance aux catégories, seules les partitions peuvent être appréhendées.

Le problème principal qui se pose est la construction du réseau. Si la structure peut parfois être apprise par une expertise, la définition des probabilités pose plus de difficultés. Les communautés statistique et intelligence artificielle se sont penchées sur ce problème, Pearl [78]. Crawford et Fung [23] présentent *Constructor*, un système de construction de réseaux bayésien et markovien à partir d'exemples fondé sur une recherche heuristique de voisins de chaque nœud dans le réseau. Un autre reproche est souvent fait aux réseaux, il s'agit de leur complexité. Dans la pratique seuls des réseaux de taille réduite peuvent être envisagés, Breeze *et al.* [16] présentent à cet effet une approche de la détection de pannes par réseaux bayésiens (on trouve par exemple, un réseau pour la détection de problèmes d'impressions dans un réseau informatique constitué de 25 nœuds).

2.7.2 Méthodes non paramétriques fondées sur des distances

Ces méthodes ne font pas d'hypothèses a priori sur la distribution de probabilité des objets d'apprentissage. Elles utilisent directement les points de Z (on parle aussi de méthodes directes). Deux règles de décision différentes peuvent leur être associées.

La première, que l'on nomme souvent règle simple de décision, relève d'un principe de vote. Le point à classer est affecté à la classe représentée majoritairement parmi les points d'apprentissage qui lui sont les plus proches.

La seconde consiste à reprendre la règle de Bayes d'erreur minimale, les points de Z sont alors utilisés en vue d'évaluer la densité conditionnelle de probabilité $f(x/C_i)$. Si l'on note V_h le volume d'un hypercube centré autour de x et $k_i(x)$ le nombre de points de Z_i inclus dans l'hypercube alors un estimateur de $f(x/C_i)$ est :

$$\hat{f}(x/C_i) = \frac{k_i(x)/|Z_i|}{V_h} \quad (2.5)$$

Les méthodes que nous présentons dans les deux sections suivantes diffèrent par le paramètre qui est fixé a priori, le nombre k ou le volume de l'hypercube. Ce sont des méthodes d'estimation locale cependant elles possèdent des propriétés de convergence asymptotiques [47].

Les méthodes de type k plus proches voisins (k -ppv)

Il s'agit ici d'une approche tout d'abord introduite par Fix et Hodges [37] et [38]. Puis Cover et Hart [22] ont ensuite montré la convergence de cette approche vers la règle de Bayes d'erreur minimale. En reprenant l'équation 2.5, cette approche consiste à fixer le nombre $k_i(x)$ et à adapter le volume pour contenir exactement le nombre de points de Z recherchés. On affecte alors x à la catégorie majoritaire parmi les k plus proches voisins.

Ici c'est la fixation de la valeur du paramètre k qui va être délicate. En effet un k trop petit va engendrer une forte sensibilité au bruit d'échantillonnage. La méthode est alors faiblement robuste. Au contraire un k trop grand va engendrer un phénomène de non décision qui correspond à une uniformisation des décisions. La dérive est alors la suivante, plus k va être grand plus la méthode

va avoir tendance à choisir la classe la plus représentée au sein de Z . Bezdek *et al.* [10] proposent *an inelegant but pragmatic way* pour choisir une bonne valeur de k , elle consiste à faire varier k puis à choisir le k^* qui minimise le taux d'erreur de classification.

Cette méthode a été étendue au cas de l'appartenance floue à des catégories par différents auteurs Keller *et al.* [58], Bezdek *et al.* [10], Béreau et Dubuisson [8] et, Yang et Chen [118] qui présentent une approche généralisée des trois précédents auteurs. On trouve dans l'article de Keller *et al.* [57] une comparaison des cas flou et non flou à une travers une application de reconnaissance des formes (évaluation de la cuisson de steak).

Le principe est d'utiliser la distance du point à classer à ses k plus proche voisins pour pondérer l'importance de chaque voisin dans la décision finale de classification, plus un voisin va être proche du point à classer plus son poids sera important dans le résultat de l'affectation. Dans l'approche de Keller *et al.* [58] le degré d'affectation d'un point x à la catégorie C_t est donnée par :

$$\mu_t(x) = \frac{\sum_{j=1}^k u_{ij} \left(1/\|x - z_j\|^{\frac{2}{m-1}}\right)}{\sum_{j=1}^k \left(1/\|x - z_j\|^{\frac{2}{m-1}}\right)}$$

avec $\{z_1, \dots, z_k\}$ l'ensemble des k plus proches voisins de x , u_{ij} le degré d'appartenance de z_j à C_i et m un paramètre technique de l'intervalle $]1, +\infty[$ permettant de régler le caractère plus ou moins graduel de l'affectation. m pondère la distance lors du calcul de la contribution de chaque voisin au degré d'affectation. Plus m est élevé plus la partition est floue, et plus il est proche de 1 plus les voisins proches prennent de l'importance et plus l'affectation est nette.

Remarque 2 *Cette extension floue de la méthode induit un changement assez important : la prise en compte de la distance comme une valeur qui n'est plus uniquement ordinale.*

Bezdek *et al.* [10] proposent une définition générale des règles de type *kppv* (*kN.N* en anglais). et introduit le traitement du cas où les points d'apprentissage ne sont pas classés. Il propose pour cela une approche fondée sur les *c-moyennes floues* (Fuzzy *C*-Means, F.C.M.) [11]. Il ressort que l'utilisation de F.C.M.-N.P. (association de F.C.M. et du plus proche prototype, Nearest Prototype) est équivalente à l'utilisation de F.C.M.-(f)*k*.N.N. (association des F.C.M. et des k plus proches voisins flous) en terme de taux de mauvaise classification [9]. On aura donc tendance à conseiller le choix de F.C.M.-N.P. du fait d'une plus grande facilité d'implémentation et d'une plus grande rapidité en temps d'exécution. On peut aussi citer Keller et Krishnapuram [59] qui proposent une approche possibiliste des *c-moyennes*, les *c-moyennes possibilistes* (Possibilistic *C*-Means).

Tout comme pour l'approche de Parzen, la méthode converge asymptotiquement vers la règle de décision de Bayes d'erreur minimale lorsque $|Z|$ augmente et que $k = \left(1/\sqrt{|Z|}\right)$. Les performances des méthodes de type voisin peuvent être grandement augmentées, au prix d'une faible augmentation du temps de calcul, par l'utilisation de plusieurs distances, une par catégorie. L'inconvénient majeur en fait réside dans la difficulté du choix de plusieurs distances.

Sur l'exemple de la figure 2.10, le point à classer sera affecté à C_1 par les deux méthodes. Il convient pour chacune d'entre elles de fixer un paramètre, h pour Parzen et k pour les k -plus proches voisins.

Ces méthodes permettent d'aborder des problèmes de séparabilité aussi bien linéaire que non linéaires. L'information nécessaire est relativement faible, pas de justification de connaissance par un expert par exemple. Il y a aussi la possibilité de réaliser un apprentissage incrémental par ajout et suppression au cours du processus de classification de nouveaux points exemples de Z . Différentes procédures d'amélioration ont été proposées pour améliorer ces méthodes, comme l'utilisation de

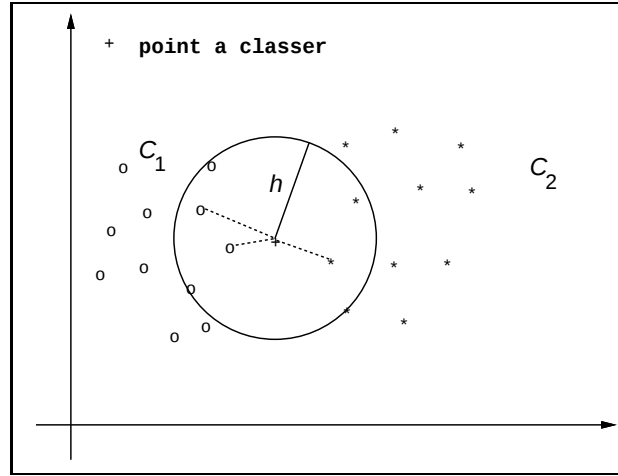


FIG. 2.10 – Affectation d'un point par les 3 plus proches voisins et par la fenêtre de Parzen de rayon h

F.C.M. ou la décomposition de l'espace en régions pour limiter la recherche du k -ppv. Les problèmes inhérents au choix des distances sont fortement présents ainsi que celui du choix de k .

Les noyaux ou fenêtres de Parzen

Cette approche est due à Parzen [75]. Alors que la méthode des k -ppv consiste en la recherche d'un nombre fixe d'individus dans un espace non contraint, la méthode des fenêtres de Parzen consiste en la recherche d'un nombre maximum d'individus dans un espace fixé. Le nombre de voisins n'est pas fixé, c'est l'espace de recherche qui l'est.

On ne s'expose plus au risque d'accepter un individu très éloigné comme voisin du fait de la limitation de l'espace de recherche. En revanche on s'expose à ne pas trouver suffisamment de voisins (ou même pas du tout). Hand [47] démontre que l'estimation de la densité conditionnelle converge vers $f(x/C_i)$ quand la dimension de la fenêtre diminue suivant une loi en $O\left(1/\sqrt{|Z_i|}\right)$ quand $|Z_i|$ augmente.

2.7.3 Méthodes de type *machine learning*

Les méthodes de type *machine learning* - on parle aussi d'apprentissage automatique - sont des méthodes d'apprentissage et de classification à partir d'exemples. On peut distinguer deux phases, la phase d'apprentissage supervisé, qui correspond à la construction à partir des exemples du modèle de décision, et la phase d'affectation de nouveaux objets.

Michie *et al.* [69] distinguent deux approches possibles pour l'apprentissage. La première, *du particulier au général*, consiste à considérer un ensemble maximal de règles de classification, puis à réduire cet ensemble de règles à un sous-ensemble qui le résume au mieux. La seconde approche, *du général au particulier*, consiste à construire les règles une à une jusqu'à obtenir une bonne description de l'ensemble d'apprentissage. Nous présentons ces approches à travers respectivement les tables de décision et la construction d'arbres de décision.

Patients	Symptômes			Diagnostic
	<i>maux de tête</i>	<i>douleurs musculaires</i>	<i>temp.</i>	<i>grippe</i>
e_1	oui	oui	normale	non
e_2	oui	oui	élevée	oui
e_3	oui	oui	très élevée	oui
e_4	non	oui	normale	non
e_5	non	non	élevé	non
e_6	non	oui	très élevée	oui
e_7	non	oui	élevée	oui
e_8	non	oui	très élevée	non
e_9	oui	oui	élevée	oui

TAB. 2.1 – Table de décision pour le diagnostic de la grippe

Tables de décision

Ces méthodes ont pour objectif la production de règles (proches de celles utilisées par les systèmes experts). Le schéma général consiste en la création d’une table dite complète, c’est-à-dire une table constituée de tous les exemples disponibles, chaque exemple étant assimilé à une règle. Cet ensemble important de règles ainsi engendré n’est pas utilisable dans la pratique (nombre important de règles, règles trop spécifiques et règles contradictoires) la méthode va donc constituer en la réduction de cette table. La réduction peut s’effectuer de deux manières, par suppression de règles en cas de redondance et par simplification de règles en cas d’attributs non pertinents.

Voyons un petit exemple de table de décision, table 2.1, constituée de 9 objets d’apprentissage décrits sur 3 critères qualitatifs et affectés à deux classes. L’objectif est de diagnostiquer la grippe chez des patients présentant divers symptômes.

On voit ici que les règles engendrées par les objets e_2 et e_9 sont redondantes, il conviendra donc d’en éliminer une. On peut aussi déduire une règle simplifiée qui n’utilise qu’un attribut (ou variable) : **SI temp. = normale ALORS pas de grippe** car quels que soient la valeur des autres attributs la décision sera toujours la même.

L’exemple du tableau 2.1 met à jour un des problèmes auxquels ces méthodes sont confrontées, la prise en compte de règles contradictoires, on parle d’inconsistance de la table de décision. Ici, il s’agit des règles e_6 et e_8 . Cette situation peut avoir deux causes principales. Tout d’abord une erreur peut s’être glissée dans les données, cela est une éventualité à ne pas écarter. La taille importante des fichiers peut en être la raison aussi bien que le caractère subjectif de certaines affectations effectuées par des opérateurs humains. La seconde cause est de nature plus structurelle, elle vient directement du modélisateur et intervient lorsque les attributs (ou critères ou variables) ne forment pas une *famille cohérente* au sens de Roy [97]. C’est en fait l’exigence d’exhaustivité qui n’est pas respectée (deux objets ayant les mêmes performances se doivent d’être indifférent et donc induisent la même décision).

La théorie des ensembles approximatifs (*rough sets*) proposée par Pawlak [76] permet de donner une réponse à ce problème. Elle consiste à édicter deux types de règle, les règles déterministes qui concluent de manière certaine à l’appartenance ou à la non appartenance à une catégorie et les règles non déterministes qui permettent de décrire une part des objets d’apprentissage mais qui sont contredites par quelques uns d’entre eux. Remarquons que cette approche ne supprime pas l’inconsistance mais la répercute à travers les règles induites. De ce fait il n’y a pas de parti pris vis à vis de l’inconsistance, celle-ci n’est pas corrigée [45]. Par ailleurs la théorie des ensembles approximatifs est propice à la classification, voir Pawlak [77]. On trouve aussi dans Słowiński [106] diverses applications des *rough sets* dont certaines concernent directement les problèmes de classification (chap. 1.6, 1.8, 2.2, 2.3 et 3.3).

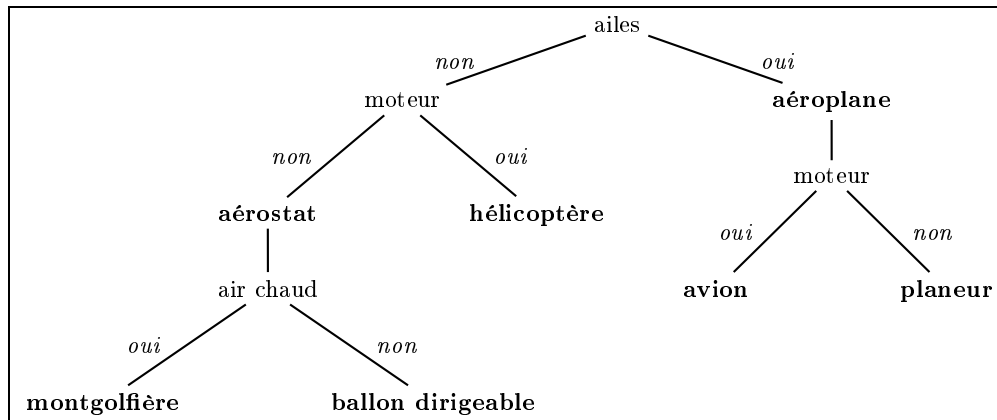


FIG. 2.11 – *Arbre de décision pour la classification d'objets volants*

Arbres de décision

Les arbres de décision sont des structures de raisonnement. La fonction d'affectation y est représentée par un arbre, elle modélise des enchaînements de tests qui aboutissent aux affectations. Chaque feuille représente une décision possible, dans le cadre de la classification il s'agira de la classe d'affectation. Chaque nœud représente un test sur un critère, les réponses possibles correspondent alors aux branches issues du nœud. Voyons sur la figure 2.11, un exemple d'arbres de décision permettant de détecter la nature d'un objet volant.

Il s'agit ici d'un arbre binaire. Cet arbre permet de modéliser la série de règles suivantes :

R1 *SI ailes ALORS aéroplane*

R2 *SI aéroplane ET moteur ALORS avion SINON planeur*

R3 *SI non ailes ET moteur ALORS hélicoptère SINON aérostat*

R4 *SI aérostat ET air chaud ALORS montgolfière SINON ballon dirigeable*

Ici quatre tests sont effectués. Les catégories qui se trouvent sur les feuilles forment une partition. Il est aussi possible d'envisager des résultats intermédiaires avec des catégories au niveau des nœuds. Ces catégories sont en fait des regroupements de plusieurs autres, par exemple la catégorie des aéroplanes correspond à l'union des catégories avions et planeurs. Cela permet d'appréhender le problème des valeurs manquantes aisément. Comme il n'est pas nécessaire d'arriver au niveau des feuilles pour prendre une décision, un objet qui ne répond pas à un test pourra quand même être affecté. Par exemple si l'information sur l'air chaud n'est pas disponible il sera quand même possible de l'affecter à la catégorie aérostat.

Si la construction de l'arbre de décision de la figure 2.11 peut être effectuée par entretien auprès d'un expert ce n'est, dans la plupart des cas, pas possible. Il va donc falloir utiliser des méthodes de construction de ces arbres.

Méthodes de construction d'arbres de décision

Ce sont des méthodes d'apprentissage non supervisé, elles consistent en une segmentation récursive d'un ensemble de points d'apprentissage. Pour cela un critère évaluant la séparation de deux sous-ensembles de points est nécessaire. La variable la plus discriminante est isolée par optimisation

du critère de segmentation, deux sous populations sont alors obtenues ainsi que la racine de l'arbre. La procédure est ensuite appelée de manière récursive sur chaque sous population jusqu'à obtention de sous populations constituée de points de la même catégorie. On peut citer ID3 et C4.5 de Quinlan respectivement [89] et [90]. Ces deux méthodes permettent la construction d'arbres de décision *nets*. Plus récemment, Bouchon-Meunier *et al.* [13] ont proposé un nouveau critère de segmentation, l'*entropy star* qui est une généralisation floue de l'entropie de Shannon. L'utilisation de ce critère permet la construction d'arbre de décision flous par une méthodes dérivées de ID3. Des probabilités floues sont définies au niveau de chaque branche de l'arbre, l'affectation à une catégorie est alors calculée à l'aide d'une probabilité conditionnelle floue.

Remarquons que les tests construits à chaque nœuds ne font intervenir que les variables de décision (et non des paramètres du modèle n'ayant pas de signification concrète comme les poids dans les réseaux neuronaux), de ce fait le parcours de l'arbre en vue d'une affectation présente un caractère explicatif important pour un décideur (le parcours de l'arbre peut être exprimé par une règle au sens des systèmes experts). De plus le principe de segmentation permet aussi d'éliminer des variables qui ne seraient pas pertinentes. Cela peut être du au fait qu'elles n'interviennent pas dans le processus de décision ou à une redondance d'information.

Il s'agit là du principe général de construction d'un arbre de décision. Les méthodes vont différer par leur capacité à considérer des arbres n -aires (avec dans la pratique n pas trop grand) ou bien des critères d'arrêt plus souples que l'obtention d'une sous population parfaitement homogène (par exemple en acceptant un certain pourcentage de mauvaise classification). En effet considérer comme critère d'arrêt l'obtention de sous population pure conduit souvent à des arbres de taille trop importante pour pouvoir être utilisés. C'est pourquoi des coûts de mauvaise classification peuvent par exemple être utilisés. Cependant la technique de l'élagage est la plus couramment utilisée, elle consiste à engendrer un arbre suffisamment important (voir complet comme NewID [69]) puis à effectuer des coupes dans celui-ci (*back-track pruning*).

A priori les arbres de décision permettent de résoudre des problèmes d'affectation où les catégories forment une partition. Cependant il est possible de construire une fonction d'affectation à des catégories qui forment un regroupement. Il faut pour cela construire un arbre pour chaque catégorie en considérant pour chacune d'entre elles un problème de segmentation en deux classes.

DEUXIÈME PARTIE

Méthodes développées

Chapitre 3

Méthodes de Filtrage Flou

Résumé

Nous présentons dans ce chapitre la méthode de classification multicritère que nous avons développée. Il s'agit d'une méthode de filtrage flou exploitant une relation d'indifférence floue. Après avoir évoqué les différences principales entre notre méthode et les méthodes existantes, nous présentons la construction d'une relation floue d'indifférence fondée sur le principe de concordance et non discordance. Nous présentons ensuite les deux variantes de la méthode de filtrage flou par indifférence. Enfin nous terminons ce chapitre en proposant un moyen d'engendrer une explication du résultat de l'affectation.

Mots clés : indifférence, concordance et non-discordance, filtrage flou par indifférence, catégories multiprofiles, prototype, profils centraux

3.1 Introduction

Nous présentons dans ce chapitre les méthodes multicritères de classification que nous avons développé. Ce type de méthode repose sur l'utilisation d'une relation d'indifférence floue. Nous présentons donc dans un premier temps la construction d'une relation d'indifférence floue qui diffère quelque peu de la définition de l'indifférence donnée par Roy [97] comme une fonction du surclassement. Nous présentons ensuite une définition des méthodes multicritères de tri par comparaisons binaires et des principales propriétés que ce genre de méthode peut revêtir en insistant sur les différences entre les méthodes fondées sur l'utilisation d'une relation de préférence et celles fondées sur l'utilisation d'une relation d'indifférence. Enfin nous présentons une famille de méthodes, les méthodes de type Plus Indifférents Prototypes ainsi que la capacité qu'elles présentent à pouvoir fournir une explication de leur résultat.

3.2 Quelques limites des méthodes existantes

Les méthodes multicritères de classification que nous avons présentées ne s'inscrivent pas exactement dans l'esprit la classification au sens auquel les statisticiens et les spécialistes de l'intelligence artificielle l'entendent. Elles sont plus proches de l'idée de rangement que de la reconnaissance des formes. L'objectif est en fait d'obtenir un préordre sur les actions en ne les comparant pas entre elles, et ce pour les raisons suivantes :

- Chaque action doit être jugée intrinsèquement généralement en vue d'une acceptation ou d'un rejet, ou d'une évaluation du type *très bon*, *bon*, *moyen*, *mauvais* et *très mauvais*. Ces évaluations sont effectuées en fonction de norme prédéfinies qui servent par la suite de justification au décideur. Chaque objet se voit donc attribuer sa propre décision.
- L'ensemble des actions n'est pas donné a priori, celles-ci apparaissent la plupart du temps au fur et à mesure et ne peuvent donc pas entrer en compétition les unes avec les autres.
- Le nombre d'actions peut être très important et la comparaison de toutes les actions entre elles ne peut être envisagée en un temps raisonnable.

Ces méthodes se placent toutes dans le cadre de la Segmentation Multicritère (*cf.* [113]) et supposent des catégories ordonnées, elles sont donc inutilisables avec des catégories non ordonnées. Des problèmes avec des telles catégories se posent par exemple dans le domaine bancaire quand il s'agit de proposer différents produits financiers à des clients, dans le domaine de l'enseignement si l'on désire faire des proposition pour l'orientation des étudiants vers différentes filières, ...

La définition des catégories n'est pas forcément une chose figée dans le temps, il apparaît donc utile de voir s'il est possible de faire aisément évoluer la définition des catégories. La méthode Moscarolla et Roy s'y prête bien car il suffit de changer un des profils en respectant la condition de cohérence. En revanche nTOMIC pose plus de problèmes du fait de la définition des catégories par des seuils sur les indices d et D . S'il est facile de modifier ces seuils, il est beaucoup plus difficile de savoir comment les modifier et de combien pour que cela corresponde aux aspirations réelles du décideur. Quant à Electre Tri, la modification des catégories passe par la modification des profils ce qui est aisé à faire et possède une signification concrète pour le décideur.

Les modèles d'affectation que nous avons présentés ne possèdent pas tous le même pouvoir explicatif. La procédure de segmentation trichotomique et Electre Tri possèdent une procédure d'affectation qui se fonde uniquement sur les résultats des comparaisons de l'objet à classer aux profils, cela permet de présenter au décideur une explication du type *x a été affecté à C_i parce que'il est meilleur que b_{i-1} et moins bon que b_i* . Comme les profils des catégories ont été fixé par

Propriétés	segmentation trichotomique	n-TOMIC	ELECTRE TRI	FFI	FFP
nb. catégories	3	3 à 12	quelconque	quelconque	quelconque
affectation	nette	nette	nette	floue	floue
relation binaire	S	S	S	$\sim$	$\prec$
structure des catégories	ordre	ordre	ordre	indifférente	ordre
définition des catégories	multiprofil limites	intervalles d'indices de surclassement	monoprofil limites	multiprofil centraux	multiprofil limites
structure sur les objets	préordre	préordre	préordre	aucune en particulier	préordre
caractère explicatif	fort	faible	fort	fort	fort

TAB. 3.1 – *Tableau comparatif des méthodes de tri multicritère*

lui cette information est riche et lui permet de justifier son choix s'il décide de suivre la recommandation. Au contraire, nTOMIC ne possède pas cette richesse explicative. L'explication que l'on peut en tirer est beaucoup trop proche du modèle mathématique pour pouvoir être exploitée aisément par le décideur puisque l'information sémantique que l'on peut associer à une règle d'affectation est du type *x est affecté à la catégorie C_i car $\alpha' \leq d(x) \leq \alpha$ et $\beta \leq D(x) \leq \beta'$* .

Nous présentons dans la section suivante une méthode d'agrégation multicritère non-totalement compensatoire permettant de comparer des objets deux à deux. Cette méthode utilise une relation floue d'indifférence. Nous montrons alors comment cette relation floue peut être utilisée pour réaliser l'affectation graduelle d'objets à des catégories prédéfinies en s'appuyant sur une technique de filtrage flou (voir [80]).

Le principe d'affectation que nous proposons consiste à examiner si une action représentée par un point x dans l'espace des critères *est proche* d'un prototype y figurant un élément typique d'une catégorie. Il nous faut donc évaluer la proximité entre deux points x et y . Pour cela, nous pourrions utiliser une métrique du type $d(x, y) = \|x - y\|$. Cependant, le désir de prendre en compte des seuils perceptifs en deçà desquels une différence n'est pas significative, la volonté de traduire des phénomènes de non-linéarité des préférences et enfin le désir d'utiliser des mécanisme d'agrégation non-totalement compensatoires nous a conduit à envisager une relation d'indifférence.

La méthode que nous proposons se caractérise et se différencie principalement des méthodes existantes par les caractères suivants :

- La capacité à effectuer une affectation graduelle des actions aux catégories c'est-à-dire à calculer un degré d'appartenance (ou d'affectation),
- Le caractère multiprofil des catégories,
- La représentation des catégories par des profils centraux (prototypes ou points exemples *cf.* §3.5.2),
- La capacité à construire les prototypes des catégories à l'aide d'exemples (*cf.* chap. 5),
- La capacité à gérer des catégories ordonnées et non-ordonnées.

3.3 Construction d'une relation d'indifférence floue

3.3.1 Modélisation floue des préférences

Dès lors que l'on cherche à comparer deux actions a et b sur un critère j afin d'asseoir une préférence ou une indifférence sur j (notées respectivement $\succ_j$ et $\sim_j$), la manière la plus simple consiste à considérer :

$$a \succ_j b \text{ si } g_j(a) > g_j(b)$$

$$a \sim_j b \text{ si } g_j(a) = g_j(b)$$

Cette modélisation correspond au modèle du *vrai critère*. Cependant cette modélisation conduit à une discrimination absolue. La moindre différence de performance entre deux actions, aussi infime soit elle, va entraîner une préférence d'une action sur l'autre. Ainsi une automobile qui roule à 226 km.h⁻¹ serait irrémédiablement préférée à une autre qui roulerait "seulement" à 225 km.h⁻¹. Ainsi afin d'éliminer ce phénomène de discrimination absolue, il est possible d'introduire la notion de seuils de discrimination. Nous avons montré au chapitre 1 à la section consacrée à la transitivité de la relation d'indifférence, que de faibles différences de performances entre deux actions n'étaient pas incompatibles avec l'établissement d'une indifférence entre ces actions sur un critère. Cette faible différence de performances pouvant être interprétée de différentes manières : par exemple comme une réelle différence de performances mais indiscernable (ex. Luce [66] et la tasse de café sucrée, les poids de Poincaré ex. 1), comme une faible différence pouvant résulter d'une part d'arbitraire qui affecte les données ou encore d'une faible différence de performances pouvant résulter de l'imprécision des mesures. Afin de mieux prendre en compte ces divers phénomènes, il est possible d'introduire un seuil de discrimination appelé seuil d'indifférence et noté q_j qui va avoir pour objectif de modéliser la plus grande différence de performance sur le critère j compatible avec une indifférence, ainsi on obtient le modèle du *quasi critère* :

$$a \succ_j b \text{ si } g_j(a) > g_j(b) + q_j$$

$$a \sim_j b \text{ si } |g_j(a) - g_j(b)| \leq q_j$$

Ce modèle, s'il semble plus réaliste que le précédent, n'a en fait fait que déplacer le problème du passage brutal de l'indifférence à la préférence. En effet il va toujours être possible par une très faible variation de passer de l'indifférence à la préférence. Si l'on reprend l'exemple automobile du paragraphe précédent mais en fixant maintenant un seuil d'indifférence $q_j = 15$ km.h⁻¹. Deux voitures a et b roulant respectivement à 220 km.h⁻¹ et à 235 km.h⁻¹ vont être considérées comme indifférentes alors que si b roulait à 236 km.h⁻¹, soit une différence de 1 km.h⁻¹ (4 à 5 plus lent que la marche d'un être humain), b serait alors préférée à a . Cette constatation amène à considérer une véritable transition entre les situations de préférence et d'indifférence. Elle se traduit par l'adjonction d'une nouvelle relation de comparaison, la préférence faible $\succsim_j$. Cette relation constitue une situation intermédiaire entre l'indifférence et la préférence, la préférence $\succ_j$ est alors nommée préférence stricte. La préférence faible peut être définie à l'aide d'un second seuil, le seuil de préférence p_j qui va avoir pour objectif de modéliser la plus petite différence de performance compatible avec l'établissement d'une préférence entre deux actions. Ce modèle est appelé *pseudo-critère* et les trois relations de comparaisons sont définies par :

$$a \sim_j b \text{ si } |g_j(a) - g_j(b)| \leq q_j$$

$$a \succ_j b \text{ si } g_j(a) \geq g_j(b) + p_j$$

$$a \succsim_j b \text{ si } q_j < |g_j(a) - g_j(b)| < p_j$$

L'utilisation de la relation de préférence faible $\succsim_j$ permet de remédier au passage brutal d'une situation d'indifférence à une situation de préférence par une situation intermédiaire. Cependant le phénomène de seuil évoqué précédemment, même si il est atténué par la relation de préférence faible,

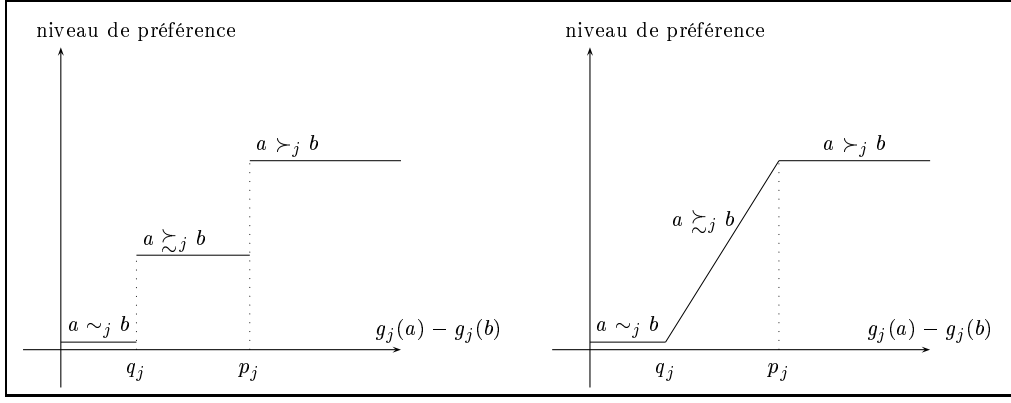


FIG. 3.1 – Préférence faible stricte et graduelle

demeure: une infime différence de performance suffit pour passer d'une situation d'indifférence à une situation de préférence faible, et il en est de même pour passer d'une situation de préférence faible à une situation de préférence stricte. Si l'on reprend l'exemple précédent en fixant le seuil d'indifférence $q_j = 10 \text{ km.h}^{-1}$ et le seuil de préférence $p_j = 20 \text{ km.h}^{-1}$, un voiture a roulant à 220 km.h^{-1} sera indifférente à une voiture b roulant à 210 km.h^{-1} alors qu'elle lui aurait été faiblement préférée si b avait roulé 209 km.h^{-1} . De même a est strictement préférée à c qui roule à 200 km.h^{-1} alors qu'elle ne lui serait que faiblement préférée si c roulait à 201 km.h^{-1} . Nous pourrions encore atténuer ce phénomène de seuil par l'adjonction de nouveaux échelons entre la préférence et l'indifférence, cependant cela deviendrait rapidement ingérable de part le nombre de seuils à considérer et le nombre de relations résultantes. Le désir de prendre en compte ce phénomène de seuil nous amène donc à envisager une autre voie.

Nous pouvons modéliser cette zone de passage de l'indifférence à la préférence, non pas par un (ou plusieurs) échelons, mais par une transition graduelle. Au lieu de définir des relations de comparaisons nettes il est possible de définir des relations de comparaisons graduelles. Un décideur ne sera plus indifférent ou pas entre deux actions, il sera plus ou moins indifférent. De même, une action a ne sera pas nécessairement préférée ou pas préférée à une action b mais plus ou moins préférée à b . Les relations de préférence et d'indifférence $\succ_j$ et $\sim_j$ sont alors des relations binaires valuées à valeur dans $[0, 1]$ où cette valeur exprime la force de la relation considérée (voir figure 3.1).

Par la suite nous nous intéresserons exclusivement à la construction de relations d'indifférence. Ainsi une situation telle que $\sim_j(a, b) = 1$ signifie que a et b sont totalement indifférentes, $\sim_j(a, b) = 0$ signifie que a et b ne peuvent en aucun cas être considérées comme indifférentes et $0 < \sim_j(a, b) < 1$ signifie que a est d'autant plus indifférent à b que la valeur de $\sim_j(a, b)$ est grande. On peut par exemple représenter l'indifférence à l'action a roulant à 220 km.h^{-1} sur le critère vitesse en considérant les deux seuils $q_j = 10 \text{ km.h}^{-1}$ et $p_j = 20 \text{ km.h}^{-1}$ par la figure 3.2.

3.3.2 Précisions sur l'utilisation des sous-ensembles flous

Nous présentons brièvement dans cette section les concepts nécessaires à l'utilisation d'une logique multivalente pour l'agrégation des préférences dans le cadre qui nous intéresse. Pour plus de précisions, nous pouvons nous rapporter pour l'utilisation de la logique floue en agrégation des préférences à [79] et [84], à [107] pour l'utilisation du flou en aide à la décision et à [56] pour une introduction à la théorie des sous-ensembles flous.

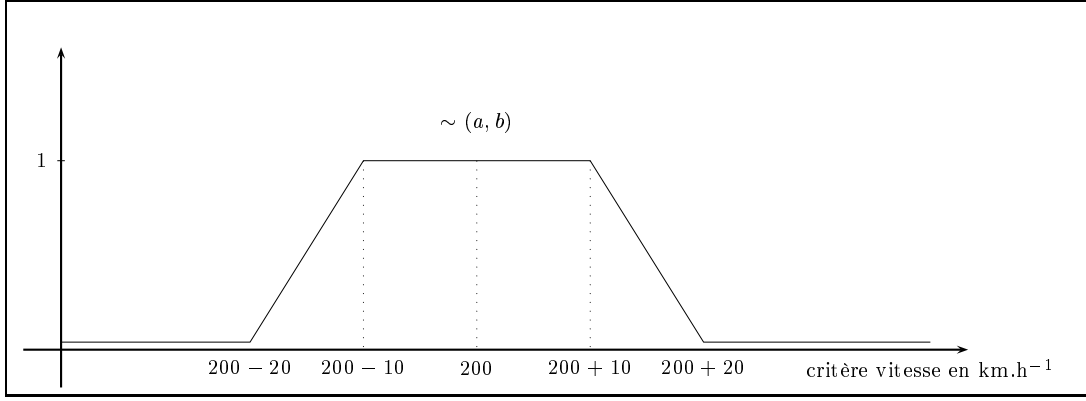


FIG. 3.2 – Valeurs de l'indifférence sur le critère vitesse

Nous définissons tout d'abord la notion de sous-ensemble flou :

Définition 10 (sous-ensemble flou) Soit X un ensemble d'objets, un sous-ensemble flou E construit sur X est un ensemble de couples $\{(x, \mu_E(x)) \mid x \in X\}$ où $\mu_E(x)$ est appelée la fonction d'appartenance de x à E . La fonction μ_E est une fonction définie sur X à valeurs dans $[0, 1]$. L'ensemble X est appelé le support de E .

L'utilisation des sous-ensembles flous va permettre de considérer, de prendre en compte des informations vagues (imprécises, incertaines ou indéterminées) et de gérer des prédicats vagues. Les prédicats vagues vont pouvoir être représentés à l'aide de fonctions d'appartenance. La fonction d'appartenance d'un objet x peut alors être interprétée comme un degré d'adéquation de l'objet au prédicat, voir [17]. Ainsi si l'on considère le sous-ensemble flou des personnes de grande taille, le degré d'appartenance d'une personne à ce sous-ensemble peut être interprété comme le caractère plus ou moins grand de la personne. Le support du sous-ensemble flou est alors l'ensemble de toutes les personnes.

Nous pouvons maintenant définir la notion d'inclusion, pour cela voyons un exemple. Si l'on considère, l'ensemble des personnes de grande taille et l'ensemble des géants. Il semble évident de considérer que l'ensemble des géants est inclus dans l'ensemble des personnes de grande taille, de ce fait il semble naturel d'affirmer que le degré d'appartenance d'un individu à la catégorie des personnes de grande taille est au moins aussi grand que celui à la catégorie des géants, voir figure 3.3. Nous avons alors la définition suivante selon [119] :

Définition 11 (inclusion [119]) Soient A et B deux sous-ensembles flous de support X , A est inclus dans B et noté $(A \subset B)$ si et seulement si : $\forall x \in X, \mu_A(x) \geq \mu_B(x)$.

Nous venons ici de présenter une manière de représenter des prédicats vagues à partir de sous-ensembles flous. Tous en conservant cette représentation ensembliste, il nous faut maintenant définir une manière permettant de représenter des informations plus complexes constituées de plusieurs prédicats. Nous allons à cet effet définir les opérateurs logiques classiques (négation $\neg$, et $\wedge$ et ou $\vee$) ainsi que leurs liens avec les opérations ensemblistes. Nous avons pour toute paire (A, B) de sous-ensembles flous de support X :

$$\overline{A} = \{x \in X \mid \neg(x \in A)\} \quad (3.1)$$

$$A \cap B = \{x \in X \mid (x \in A) \wedge (x \in B)\} \quad (3.2)$$

$$A \cup B = \{x \in X \mid (x \in A) \vee (x \in B)\} \quad (3.3)$$

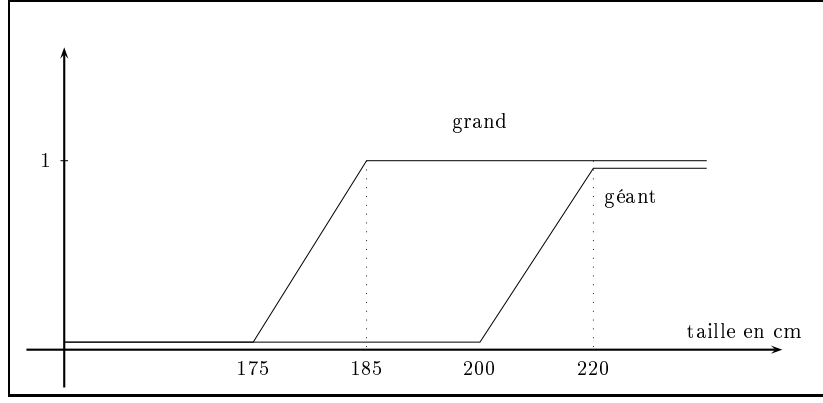


FIG. 3.3 – Appartenance aux grands et aux géants

Nous pouvons maintenant définir les connecteurs logiques, ainsi que les quantificateurs *il existe* ($\exists$) et *quel que soit* ($\forall$) en considérant que $\mathcal{P}$ et $\mathcal{Q}$ sont des propositions (des formules bien formées, cf. [79] p.123) et que $v(\mathcal{P}) \in [0, 1]$ est la valeur de vérité de $\mathcal{P}$:

$$v(\neg \mathcal{P}) = N(v(\mathcal{P})) \quad (3.4)$$

$$v(\mathcal{P} \wedge \mathcal{Q}) = T(v(\mathcal{P}), v(\mathcal{Q})) \quad (3.5)$$

$$v(\mathcal{P} \vee \mathcal{Q}) = S(v(\mathcal{P}), v(\mathcal{Q})) \quad (3.6)$$

$$v(\exists x \mathcal{P}(x)) = S_x \{v(\mathcal{P}(x))\} \quad (3.7)$$

$$v(\forall x \mathcal{P}(x)) = T_x \{v(\mathcal{P}(x))\} \quad (3.8)$$

avec N , T et S respectivement des opérateurs de types négation, t-norme et co-norme définis tels que :

Définition 12 (négation stricte) *Un opérateur de négation stricte N est une fonction strictement décroissante de $[0, 1]$ dans $[0, 1]$ définie telle que :*

$$N(0) = 1$$

$$N(1) = 0$$

$$N(N(x)) = x$$

Exemple 16 (négations strictes)

$$N(x) = 1 - x$$

$$N(x) = (1 - x^2)^{\frac{1}{2}}$$

$$N(x) = (1 - x^{\frac{1}{2}})^2$$

Définition 13 (t-norme) *Une t-norme T est un opérateur d'agrégation défini comme une fonction de $[0, 1]^2$ dans $[0, 1]$ respectant les propriétés suivantes :*

$$\forall x \in [0, 1], T(x, 1) = x \text{ (1 est élément neutre)}$$

$$\forall (x, y) \in [0, 1]^2, T(x, y) = T(y, x) \text{ (symétrie)}$$

$$\forall (x, y, z) \in [0, 1]^3, T(x, T(y, z)) = T(T(x, y), z) \text{ (associativité)}$$

$$\forall (x, y, z, t) \in [0, 1]^4, T(x, y) \leq T(z, t) \text{ avec } x \leq z \text{ et } y \leq t \text{ (monotonie)}$$

On généralise la t-norme T à n arguments par : $T(x_1, \dots, x_n) = T(T(x_1, \dots, x_{n-1}), x_n)$ avec $n \geq 2$.

Exemple 17 (t-normes)

$$T(x, y) = \min(x, y)$$

$$T(x, y) = x \times y$$

$$T(x, y) = \max(x + y - 1, 0)$$

Définition 14 (co-norme) Une co-norme S est un opérateur d'agrégation défini comme une fonction de $[0, 1]^2$ dans $[0, 1]$ respectant les propriétés suivantes :

$$\forall x \in [0, 1], S(x, 0) = x \text{ (0 est élément neutre)}$$

$$\forall (x, y) \in [0, 1]^2, S(x, y) = S(y, x) \text{ (symétrie)}$$

$$\forall (x, y, z) \in [0, 1]^3, S(x, S(y, z)) = S(S(x, y), z) \text{ (associativité)}$$

$$\forall (x, y, z, t) \in [0, 1]^4, S(x, y) \leq S(z, t) \text{ avec } x \leq z \text{ et } y \leq t \text{ (monotonie)}$$

De même que pour la t-norme, on généralise la co-norme S à n arguments par : $S(x_1, \dots, x_n) = S(S(x_1, \dots, x_{n-1}), x_n)$ avec $n \geq 2$.

Exemple 18 (co-normes)

$$S(x, y) = \max(x, y)$$

$$S(x, y) = x + y - x \times y$$

$$S(x, y) = \min(x + y, 1)$$

Par la suite nous utiliserons les opérateurs que nous venons de définir pour représenter les préférences graduelles et pour traduire les principe d'affectation. La t-norme est un opérateur conjonctif permettant de traduire le **et** logique ($\wedge$) et le quantificateur **quel que soit** ($\forall$) et la co-norme est un opérateur disjonctif permettant de traduire le **ou** logique ($\vee$) et le quantificateur **il existe** ($\exists$).

Remarque 3 Nous pouvons noter que contrairement à [84] et [79] nous utilisons la t-norme (respectivement la co-norme) pour le quantificateur $\forall$ (respectivement $\exists$) et non l'inf (respectivement le sup), ceci est dû au fait que par la suite nous considérerons des ensembles finis qui permettent l'utilisation de ces deux opérateurs sans que $T_x\{v(A(x))\} = 0$ (respectivement $S_x\{v(A(x))\} = 1$).

3.3.3 Coalitions floues concordantes et discordantes

Nous présentons dans cette section la construction d'une relation d'indifférence floue. Nous proposons une méthode d'agrégation non-totalement compensatoire permettant de comparer deux points quelconques x et y de l'espace des critères. Elle prend appui sur un principe de concordance et de non-discordance traduisant le double soucis :

- de respecter le point de vue majoritaire qui se dégage d'un ensemble de critères,
- de prendre en compte les minorités de forte opposition dans l'agrégation (droit de veto donné aux critères).

Ces principes issus des processus de vote ont été introduits par Roy dans [93] pour contrôler le caractère compensatoire des procédures d'agrégation multicritères. Nous utilisons ici une généralisation de ces principes proposés par [82] pour l'agrégation des préférences floues.

Il s'agit de construire une relation d'indifférence floue $\sim$ définie sur l'espace des critères, telle que, pour tout couple de points (x, y) , $\sim(x, y)$ évalue le degré avec lequel le décideur ou l'expert est indifférent entre les actions représentées par les points x et y . Cette indifférence peut traduire une absence de préférence en faveur de l'un des deux objets mais aussi selon le problème considéré une similarité des deux objets ou une indiscernabilité entre ceux-ci. L'idée de base proposée par [82] pour évaluer une proposition $\sim(x, y)$, où $\sim$ traduit cette notion d'indifférence, consiste à distinguer dans l'ensemble des critères deux sous ensembles flous disjoints la coalitions concordante et la coalition discordante.

La coalition concordante est le sous ensemble flou des critères en accord avec la proposition $\sim(x, y)$. On note $C_j^\sim(x, y)$ le degré d'appartenance du critère j à ce sous ensemble. En général, $C_j^\sim(x, y) = c(|x_j - y_j|)$ où c est une fonction décroissante telle que $c(0) = 1$.

Exemple 19 (degré d'appartenance à la coalition concordante)

Comme nous l'avons vu à la section 3.3.1, certaines différences de performance ne sont pas nécessairement significatives pour confirmer une préférence entre deux actions a et b sur un critère j . C'est pourquoi on introduit un modèle de préférence à deux seuils p_j et q_j (pseudo-critère), où q_j correspond à la différence de performances $|x_j - y_j|$ maximale compatible avec une totale indifférence entre a et b sur le critère j , ce seuil est appelé seuil d'indifférence. p_j correspond quant à lui à la différence minimale de performances totalement incompatible avec l'indifférence sur le critère j , ce seuil est appelé seuil de discrimination. Afin d'éviter les phénomènes de seuils vus à la section 3.3.1 51, la coalition concordante est définie comme un sous-ensemble flou. Cela permet d'obtenir une transition graduelle entre les seuils q_j et p_j et donne au critère la possibilité d'avoir une appartenance graduelle à la coalition. Ainsi une très faible différence de performances ne peut faire changer radicalement l'appartenance d'un critère à la coalition. Nous proposons de définir le degré d'appartenance du critère j à la coalition concordante par (cf. figure 3.4) :

$$C_j^\sim(x, y) = \begin{cases} 0 & \text{si } x_j - y_j \leq -p_j \\ \frac{1}{2} + \frac{1}{2} \sin \frac{\pi}{p_j - q_j} \left(x_j - y_j + \frac{p_j + q_j}{2} \right) & \text{si } -p_j \leq x_j - y_j \leq -q_j \\ 1 & \text{si } |x_j - y_j| \leq q_j \\ \frac{1}{2} + \frac{1}{2} \sin \frac{\pi}{p_j - q_j} \left(x_j - y_j - \frac{p_j + q_j}{2} \right) & \text{si } q_j \leq x_j - y_j \leq p_j \\ 0 & \text{sinon} \end{cases}$$

la coalition discordante est le sous ensemble flou des critères en violent désaccord avec l'indifférence. On note $D_j^\sim(x, y)$ le degré d'appartenance du critère j à ce sous ensemble. En général, $D_j^\sim(x, y) = d(|x_j - y_j|)$ où d est une fonction croissante telle que $d(z) = 1$ dès que $z > v_j(x_j, y_j)$. La valeur limite $v_j(x_j, y_j)$ est un seuil positif appelé *seuil de veto*. Ce seuil est choisi en concertation avec l'expert ou le décideur et représente l'écart minimal à partir duquel un critère oppose strictement son veto à l'indifférence entre les deux actions. La figure (3.4) montre un exemple de construction de ces ensembles flous pour le critère j . Tout comme pour la coalition concordante, l'appartenance d'un critère à la coalition discordante va être évaluée graduellement. Les remarques faites précédemment sur les effets de seuil sont d'autant plus importante dans le cadre de la coalition discordante. En effet le pouvoir donné à un critère lorsqu'il appartient à cette coalition est extrêmement fort : de part l'effet de veto un critère à lui seul va être capable de remettre en cause l'indifférence entre deux actions, il apparaît alors comme impératif de contrôler l'effet d'une faible différence de performance sur un seul critère. Prenons pour exemple deux actions a et b ayant des performances extrêmement proches sur $n - 1$ critères ce qui aboutit à une forte coalition concordante, le $n^{\text{ème}}$ critère possède alors un pouvoir énorme, si $|g_n(a) - g_n(b)| < v_n$ alors les deux actions vont être fortement indifférentes dans le cas contraire comme un critère oppose son veto, elles ne pourront être indifférentes.

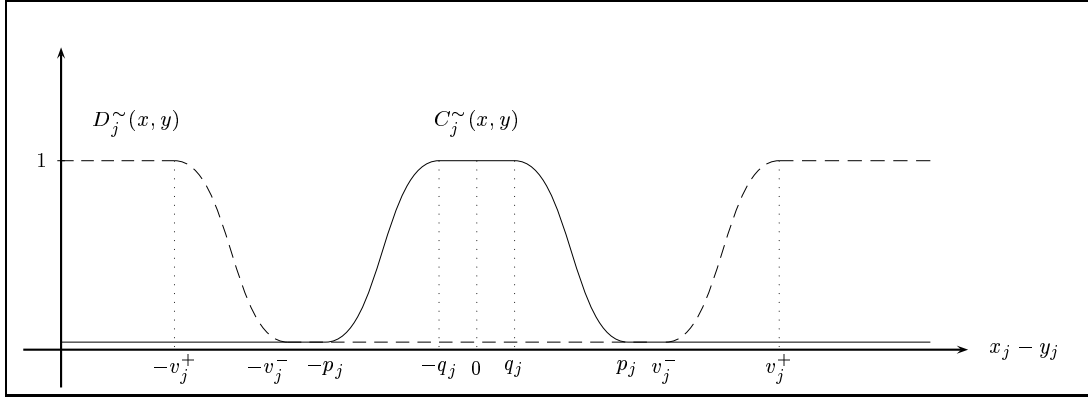


FIG. 3.4 – Coefficient d'appartenance aux coalitions concordante et discordante

Le phénomène de seuil est toujours présent et il apparaît encore plus gênant que dans le cadre de la concordance.

Exemple 20 (degré d'appartenance à la coalition discordante)

De la même manière que pour le seuil d'indifférence, il n'est pas toujours aisé de définir un seuil de veto. Nous allons estimer celui-ci à l'aide d'un intervalle $[v_j^-, v_j^+]$. La valeur v_j^- correspond à la différence de performance maximale que l'on pense être compatible avec l'indifférence globale. La valeur v_j^+ correspond, quant à elle, à la différence de performance minimale que le décideur estime être totalement incompatible avec l'indifférence globale. Cela aboutit à une transition graduelle entre les deux seuils de veto. Nous proposons de définir l'appartenance à la coalition discordante par (cf. figure 3.4) :

$$D_j^~(x, y) = \begin{cases} 1 & \text{si } x_j - y_j \leq -v_j^+ \\ \frac{1}{2} - \frac{1}{2} \sin \frac{\pi}{v_j^+ - v_j^-} \left(x_j - y_j + \frac{v_j^+ + v_j^-}{2} \right) & \text{si } -v_j^+ \leq x_j - y_j \leq -v_j^- \\ 0 & \text{si } |x_j - y_j| \leq v_j^- \\ \frac{1}{2} - \frac{1}{2} \sin \frac{\pi}{v_j^+ - v_j^-} \left(x_j - y_j - \frac{v_j^+ + v_j^-}{2} \right) & \text{si } v_j^- \leq x_j - y_j \leq v_j^+ \\ 1 & \text{sinon} \end{cases}$$

3.3.4 Agrégation multicritère fondée sur un principe de concordance / non-discordance

La définition des coalitions concordante et discordante produit n indices de concordance monocritère $C_j^~(x, y)$ et autant d'indices de discordance $D_j^~(x, y)$, $j = 1, \dots, n$. L'étape d'agrégation consiste alors à évaluer les deux indices suivants :

concordance globale : $C^~(x, y)$ qui évalue le poids de la coalition concordante. Cet indice est calculé à l'aide d'une fonction de $[0, 1]^n$ dans $[0, 1]$ croissante des indices de concordance monocritère. Il est évalué par :

$$C^~(x, y) = \psi(C_1^~(x, y), \dots, C_n^~(x, y), \omega_1, \dots, \omega_n) \quad (3.9)$$

où ψ est un opérateur de type compromis pondéré et $\omega_j \in [0, 1]$ pour $j = 1, \dots, n$ sont des coefficients positifs sommant à 1 et traduisant l'importance relative des critères.

Exemple 21 (moyenne quasi arithmétique)

$$C^\sim(x, y) = \phi^{-1} \left(\sum_{i=1}^n \omega_i \phi(C_i^\sim(x, y)) \right)$$

où ϕ est une fonction continue strictement monotone sur l'intervalle $[0, 1]$

discordance globale : $D^\sim(x, y)$ évalue l'existence (et éventuellement l'importance) des critères en violent désaccord avec la relation globale d'indifférence. Cet indice est calculé à l'aide d'une fonction croissante des indices de discordance monocritère, définie sur $[0, 1]^n$ et à valeurs dans $[0, 1]$. Cette fonction est définie telle que dès lors qu'un critère oppose son veto la discordance globale doit être maximale. $D^\sim(x, y)$ est évalué par :

$$D^\sim(x, y) = \chi(D_1^\sim(x, y), \dots, D_n^\sim(x, y), \omega_1, \dots, \omega_n) \quad (3.10)$$

où χ est un opérateur disjonctif (e.g. une co-norme pondérée) ou un opérateur de compromis vérifiant la condition (pour traduire l'effet de veto) :

$$((\exists i \in \{1, \dots, n\}) / x_i = 1) \Rightarrow \chi(x_1, \dots, x_n) = 1$$

Exemple 22 (les k plus discordants critères)

$$D^\sim(x, y) = 1 - \prod_{j \in V_k} (1 - D_j^\sim(x, y))$$

où V_k représente l'ensemble des k plus discordants critères.

Remarque 4 Cette agrégation présente l'avantage de pouvoir contrôler la synergie entre les critères discordants tous en conservant le droit de veto. L'utilisation d'opérateur d'agrégation comme des co-normes non-idempotentes présente un effet de synergie. Cela signifie concrètement que plusieurs critères faiblement discordants vont avoir la possibilité de faire croître fortement la discordance globale, éventuellement autant qu'un seul critère assez fortement discordant et donc de remettre en cause l'indifférence. Or on peut vouloir éviter qu'une discordance ne soit forte si aucun critère n'est fortement discordant. L'utilisation des k -plus discordants critères permet ce type d'agrégation en n'autorisant une synergie qu'entre un nombre réduit de critères afin d'éviter l'effet de masse évoqué précédemment. Voyant cela sur un petit exemple : supposons deux actions a et b définies par une vingtaine de critères ayant les vecteurs performances suivants :

$$g(a) = (10, 12, 18, 8, 16, 7, 11, 18, 17, 16, 2, 13, 14, 10, 12, 4, 8, 18.8, 15.8, 8)$$

$$g(b) = (11, 12, 16, 9, 14, 6, 13, 17, 18, 14, 4, 15, 14, 18.8, 3.2, 12.8, 16.8, 10, 7, 16.8)$$

Considérons que les poids sont définis tels que $\sum_{j=1}^{13} \omega_j = 0.8$ et que les seuils sont les mêmes sur tous les critères et définis par : $q_j = 2$, $p_j = 4$, $v_j^- = 8$ et $v_j^+ = 12$, $\forall j \in \{1, \dots, 20\}$. Si l'appartenance à la coalition concordante (respectivement discordante) est définie par une transition linéaire entre q_j et p_j (respectivement v_j^- et v_j^+) on a :

$$C^\sim(a, b) = \sum_{j=1}^{13} \omega_j C_j^\sim(a, b) + \sum_{j=14}^{20} \omega_j C_j^\sim(a, b) = 0, 8 + 0 = 0, 8$$

$$\forall j \in \{14, \dots, 20\}, D_j \sim(a, b) = 0, 2$$

$$\forall j \in \{1, \dots, 13\}, D_j \sim(a, b) = 0$$

k	1	2	3	4	5	6	7
$D^\sim(a, b)$	0,2	0,36	0,49	0,59	0,67	0,74	0,79
$\sim(a, b)$	0,64	0,51	0,41	0,33	0,26	0,21	0,17

TAB. 3.2 – *Discordance et indifférence en fonction de k*

En faisant varier k de 1 à 7 on obtient les valeurs du tableau 3.2. Une valeur de k égale à 1 correspond à une agrégation des indices de discordance monocritères par la t -norme idempotente et une valeur de 7 à l'agrégation par la duale-géométrique. Les valeurs intermédiaires sont obtenues en contrôlant la synergie à l'aide de l'opérateur d'agrégation des k plus discordants critères. On voit sur cet exemple que cet opérateur offre une grande souplesse pour contrôler cette synergie. Les résultats obtenus pour les valeurs 1 et 7 sont extrêmement différents. En effet pour $k = 1$ l'indifférence entre les deux actions est relativement forte alors qu'elle est très faible voire peu significative pour $k = 7$.

Le principe de concordance et de non-discordance consiste à valider la proposition $\sim(x, y)$ si et seulement si la coalition concordante est suffisamment forte et la coalition discordante suffisamment faible. Formellement on peut traduire ce principe par l'équivalence logique suivante:

$$\sim(x, y) \equiv (C^\sim(x, y) \wedge \neg D^\sim(x, y)) \quad (3.11)$$

Lorsque $\sim$ est une relation floue, on pose donc :

$$\sim(x, y) = T(C^\sim(x, y) \cdot N(D^\sim(x, y))) \quad (3.12)$$

où T est une t -norme et N une négation stricte.

Cette technique permet de réaliser une agrégation de type compromis vérifiant les propriétés importantes d'*indépendance binaire*, de *respect de l'unanimité* et de *respect du droit de veto* (pour plus de détails, voir [80]). Remarquons que l'existence d'au moins un critère totalement discordant avec la proposition $\sim(x, y)$ suffit à invalider cette proposition quelle que soit l'importance de la coalition concordante. Ainsi deux objets qui se ressemblent sur de nombreuses dimensions seront malgré tout considérés comme totalement non indifférents dès lors qu'il existe un critère selon lequel ils diffèrent radicalement. Cette caractéristique vaut souvent à ce type de méthode le qualificatif de non-totalement compensatoire.

3.4 Affectation multicritère

3.4.1 Définition formelle

Un problème d'affectation consiste à affecter des actions à des catégories prédéfinies. Nous proposons de représenter les catégories à l'aide de points de référence ou profils. Il s'agit d'actions fictives ou non, permettant de modéliser les catégories. Nous distinguons trois types de points de référence :

points de référence limites (ou profils limites) Il s'agit de points de référence qui vont représenter les limites d'une catégorie. Une catégorie sera donc représentée par un ensemble de profils inférieurs et un ensemble de profils supérieurs. Les profils inférieurs et supérieurs correspondent respectivement aux plus mauvaises et meilleures actions méritant d'entrer dans une catégorie. Nous notons $Y_t = Y_t^{sup} \cup Y_t^{inf}$ l'ensemble des points de référence limites de

la catégories C_t où Y_t^{sup} est l'ensemble des profils supérieurs et Y_t^{inf} l'ensemble des profils inférieurs. Ce type de profil sera généralement utilisé si les catégories peuvent être ordonnées même partiellement.

points de référence centraux (ou profils centraux ou encore prototypes) Il s'agit de points de référence qui correspondent à des actions typiques de la catégorie. Les catégories sont alors représentées par un ensemble de prototypes. Nous notons Y_t l'ensemble des prototypes de C_t . Ce type de profil sera généralement utilisé si les catégories ne peuvent être ordonnées.

points de référence excluants (ou anti-prototypes) Il s'agit de profils qui correspondent à des actions dont on est sûr de la non appartenance à la catégorie. Imaginons que l'on veuille définir une catégorie comme l'ensemble des actions qui n'ont pas certaines propriétés. Il est alors nécessaire de considérer des points de référence qui ne sont ni des bornes ni des prototypes des catégories. Ces points de référence servent à typer une partie de l'espace des critères qui ne correspond pas à la catégorie : il faut *ne pas leur être indifférent* pour appartenir à une catégorie. Ils correspondent à des raisons négatives pour limiter l'appartenance à une catégorie. Ce type de points de référence est généralement utilisé conjointement avec les deux autres types de point de référence.

Le décideur peut vouloir utiliser de tels points de référence pour définir une catégorie pour les deux raisons suivantes :

- il n'est pas capable de spécifier complètement une catégorie uniquement par des profils centraux ou limites,
- le nombre de profils limites ou centraux qui seraient nécessaires pour spécifier une catégorie est beaucoup trop important.

Nous notons $Y_t = Y_t^+ \cup Y_t^-$ l'ensemble des points de référence de C_t où Y_t^+ représente l'ensemble de points de référence limites ou centraux de C_t et Y_t^- l'ensemble des points de référence excluants de C_t .

Ainsi, nous pouvons maintenant définir formellement une Méthode Multicritère d'Affectation par Comparaisons Binaires (MMACB) par la définition 15.

Définition 15 (MMACB) *Soient :*

- A un ensemble d'actions,
- $F = \{1, \dots, n\}$ une famille cohérente de critères,
- $C_1, \dots, C_q$ un ensemble fini de catégories respectivement définies par $Y = Y_1 \cup \dots \cup Y_q$ les ensembles de points de référence,
- un Système Relationnel de Préférences utilisé pour comparer les actions à affecter aux points de référence des catégories (le plus souvent le s.r.p. utilisé se réduira à une seule relation binaire),
- une fonction f_t à valeur dans $[0, 1]$, permettant d'obtenir le degré d'affectation d'une action a à une catégorie C_t , prenant comme arguments les résultats des comparaisons entre a et les points de référence de toutes les catégories. Si le s.r.p. se réduit à une seule relation H non nécessairement symétrique, le degré d'affectation $\mu_t(a)$ est défini tel que :

$$\mu_t(a) = f_t(H(a, y_1), \dots, H(a, y_m), H(y_1, a), \dots, H(y_m, a)) \quad (3.13)$$

avec :

$$Y = \{y_1, \dots, y_m\}$$

$$f_t(0, \dots, 0) = 0$$

Une Méthode Multicritère d'Affectation par Comparaisons Binaires est une application h définie telle que :

$$\begin{aligned} h : A &\longrightarrow [0, 1]^q \\ a &\longmapsto (\mu_1(a), \dots, \mu_q(a)) \end{aligned}$$

Par ailleurs, notre définition est plus large que celle proposée par Larichev et Moshkovich [63] qui s'intéresse exclusivement à des catégories ordonnées et à des critères qualitatifs et utilise uniquement des relations de dominance (forcément non réflexive et transitive). Ces restrictions ne permettent alors pas l'utilisation de méthodes fondées sur cette définition (ex. ORCLASS) aux problèmes de diagnostic (profils centraux, catégories non exclusives). Il n'est pas non plus possible d'utiliser des profils centraux et encore moins d'envisager une affectation graduelle du fait du caractère net de la relation de dominance.

3.4.2 Propriétés Générales d'une Méthode Multicritère d'Affectation par Comparaisons Binaires

Nous présentons dans cette section des propriétés générales qu'une MMACB peut présenter, certaines de ces propriétés se doivent d'être respectées. Nous présenterons aux trois sections suivantes des propriétés plus spécifiques en fonction du type de méthode (filtrage par indifférence ou par préférence).

Propriété 1 (Universalité) *Chaque action doit être affectée à au moins une catégorie*

Propriété 2 (Unicité) *Chaque action doit affectée à au plus une catégorie*

Ces deux propriétés (voir [113] et [100]) imposent que le résultat de l'affectation aux catégories forme une partition de l'ensemble des actions. Cependant, Yu et Roy les appliquent dans le cadre plus restreint du filtrage par préférence. Dans un cadre général nous envisageons à la fois la possibilité d'affecter à plusieurs catégories et d'affecter à aucune. Cela est indispensable afin de traiter le problème très général du diagnostic. Si nous considérons un problème de détection de panne où les pannes correspondent aux catégories, il est tout à fait envisageable de détecter plusieurs pannes ou, dans d'autres cas, de ne détecter aucune des pannes connues. Cependant, il est toujours possible d'imposer la propriété d'universalité si l'on introduit une catégorie "*action non affectée*".

Propriété 3 (Indépendance des actions) *L'affectation d'une action ne dépend pas de l'affectation des autres actions.*

Cela signifie que l'évaluation résultant de la méthode d'affectation est uniquement fonction des performances de l'action. Cette propriété met en évidence une différence fondamentale avec les méthodes d'aide à la décision de type rangement ou choix qui est ce que l'on nomme souvent l'évaluation intrinsèque des actions. Le résultat obtenu avec une méthode de rangement ou de choix provient d'une évaluation qui intègre la totalité de l'ensemble des actions, ce qui implique que l'adjonction ou le retrait d'un élément de cet ensemble peut changer le résultat des comparaisons entre les autres actions. Par exemple dans certaines méthodes de rangement, l'adjonction d'une nouvelle action peut inverser l'ordre des autres actions.

Propriété 4 (Affectation par Comparaisons) *L'affectation d'une action est fonction uniquement de la manière dont elle se compare aux points de référence.*

Cette propriété implique que deux actions notoirement différentes qui se comparent de la même manière aux points de référence doivent être affectées à la même catégorie. Seul le résultat des comparaisons est utilisé pour évaluer l'affectation.

Propriété 5 (Neutralité des points de profils) *Le résultat de l'affectation ne dépend pas de la manière dont les profils sont pris en compte.*

Cette propriété signifie que l'ordre dans lequel les profils sont considérés ne doit pas influencer sur le résultat de l'affectation. Ainsi si l'on utilise des profils limites, dès lors que tous les profils sont pris en compte dans la procédure d'affectation, le fait de considérer les profils du meilleur au moins bon ou du plus mauvais au meilleur n'influe pas sur l'affectation des actions si la procédure demeure la même.

Propriété 6 (Stabilité par rapport au regroupement des catégories) *Le regroupement de catégories ou la division de catégories n'influe pas sur l'affectation aux catégories non concernées.*

Cette propriété correspond à la propriété de stabilité de Yu [113]. Elle signifie que l'affectation d'une action à une catégorie C_t ne dépend pas de la manière dont les autres catégories sont considérées. Ainsi le fait de regrouper ou subdiviser les catégories autres que C_t n'influe pas sur l'affectation à C_t .

Les actions vont être comparées aux points de référence des catégories à l'aide de relations binaires. Le choix de celles-ci dépend du type des points de référence des catégories. Si l'on se place dans l'hypothèse des points de référence centraux, la relation binaire H doit être une relation de similitude (satisfaisabilité, inclusion ou ressemblance [14] et [109]), de dissimilarité (distance, ...) ou encore d'indifférence [80]. En revanche si les points de référence sont des profils limites une relation définissant au moins un pré ordre partiel sera nécessaire (relations de préférence, [80]).

Utiliser un ensemble de relations de comparaisons non réduit à une seule relation est une nécessité dans le cas où à la fois des profils limites et centraux définissent les catégories. En revanche l'utilisation d'un s.r.p. avec plus d'une relation n'implique pas forcément que les catégories soient définies par plusieurs types de profil. Voyons cela sur l'exemple suivant.

Exemple 23 *Soit un ensemble de catégories définies par un seul type de points de référence, des profils limites, il est possible d'imposer pour l'affectation d'une action que celle-ci domine la borne inférieure de la catégorie et que la borne supérieure lui soit seulement préférée. Deux relations sont alors utilisées, dominance et préférence. Autrement on a la règle d'affectation suivante :*

$$a \in C_t \iff (a \Delta b^-) \wedge (b^+ P a)$$

avec b^+ et b^- les bornes supérieures et inférieures de C_t et Δ la relation de dominance.

Enfin, dans le cas où les catégories sont ordonnées, la propriété suivante doit impérativement être respectée :

Propriété 7 (Respect de la dominance avec des catégories ordonnées) *Dans le cas où les catégories sont ordonnées, soient deux actions a et b définies telles que a domine b , alors a doit être nécessairement affectée à une catégorie au moins aussi bonne que celle à laquelle b est affectée.*

Le respect de cette propriété va dépendre principalement de la manière dont les catégories sont représentées et donc de la définition des points de référence.

À partir de la définition que nous venons de donner nous pouvons distinguer plusieurs types de méthodes de classification multicritères qui vont différer à la fois par le type de relation de comparaison utilisée et par le type de points de référence. Pour désigner ces méthodes nous utilisons le terme *filtrage* dans le sens où la méthode de classification agit comme un filtre, un crible sur l'ensemble des actions. Nous distinguons donc le *filtrage par préférence* qui utilise une relation de préférence $\succ$ et généralement des profils limites et le *filtrage par indifférence* $\sim$ qui utilise une relation d'indifférence et habituellement des profils centraux. Enfin, on peut aussi envisager du *filtrage combiné* considérant à la fois certaines catégories représentées par des profils centraux et d'autres représentées par des profils limites, et utilisant conjointement des relations de préférence et d'indifférence. Nous présentons ces trois types de filtrage aux trois sections suivantes.

3.4.3 Filtrage par préférence

Ce type de filtrage considère des catégories représentées par des familles de points de référence limites. Formellement une catégorie C_t est représentée par deux sous ensembles des profils Y_t^{sup} et Y_t^{inf} correspondant respectivement aux profils supérieurs et inférieurs de C_t . Tout d'abord la propriété d'*universalité* se doit d'être respectée. Elle impose qu'il est toujours possible d'affecter une action à au moins une catégorie.

Cette propriété signifie que le résultat de l'affectation des actions forme un regroupement de celles-ci. Cependant nous n'imposons pas nécessairement que les catégories forment une partition (ajout de la propriété d'unicité). En effet, il est tout à fait envisageable qu'une action soit affectée à plusieurs catégories successives, cela pouvant par exemple traduire une incertitude dans l'affectation.

Propriété 8 (Catégories ordonnées) *Les catégories sont ordonnées que ce soit par une notion de préférence ou autre.*

Ces deux propriétés (*Universalité* et *Catégories ordonnées*) impliquent que les catégories n'ont une signification que les unes par rapport aux autres. Ce qui implique que si l'on modifie la définition d'une catégorie alors cela modifie nécessairement la définition d'au moins une autre catégorie (sauf si la modification porte sur la frontière supérieure de la meilleure ou sur la frontière inférieure de la plus mauvaise catégorie).

Nous distinguerons deux types de configurations avec catégories ordonnées. Le plus souvent l'ordre sur les catégories peut s'exprimer en terme de préférence. Cela signifie que la catégorie C_{t-1} (par convention) est meilleure que la catégorie C_t . C'est le cas envisagé par Moscarola [70] dans la technique de classification trichotomique de dossiers de crédits. Trois catégories sont envisagées : celle des dossiers acceptés (C_1), celle des dossiers incertains (C_2) et celle des dossiers refusés (C_3).

Cependant l'ordre sur les catégories ne peut pas toujours être envisagée en terme de préférence, on ne peut pas dire si les objets d'une catégorie sont meilleurs ou moins bons que ceux d'une autre. C'est le cas par exemple si l'on cherche à évaluer la cuisson d'un biscuit sortant d'un four en l'affectant à l'une des catégories suivantes : très peu cuit (C_1), peu cuit (C_2), bien cuit (C_3), un peu trop cuit (C_4) et beaucoup trop cuit (C_5) ([85] et [33]). Ces catégories sont effectivement ordonnées mais il n'est pourtant pas raisonnable d'affirmer que les objets de C_1 sont meilleurs que ceux de C_3 . Il semblerait même plus naturel de donner le préordre de préférence figure 3.5. Ainsi, nous voyons que même si l'ordre des catégories ne fait pas intervenir de notions de préférence, il est tout de même possible d'envisager des méthodes de filtrage par préférence. Dans cet exemple la notion de préférence considérée par la méthode correspond à la cuisson : $\succ = \text{"plus cuit que"}$.

L'affectation à C_t est effectuée en utilisant une relation de préférence $\succ$ et repose sur le principe consistant à trouver un encadrement de l'action à affecter par les profils inférieurs et supérieurs des

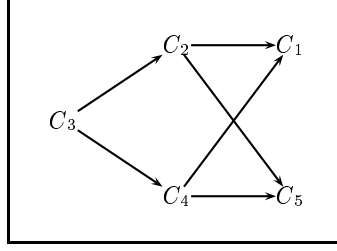


FIG. 3.5 – Préordre de préférence sur les catégories de cuisson

catégories. Formellement on distingue deux principes d'affectation :

affectation ascendante Ce mode d'affectation consiste à considérer les catégories de la plus mauvaise à la meilleure et à asseoir l'affectation en cherchant un profil supérieur préféré à l'action sans que les profils inférieurs ne lui soient préférés, ce principe se traduit par :

$$(a \in C_t) \equiv (\exists y \in Y_t^{sup}, y \succ a) \wedge (\forall y' \in Y_t^{inf}, \neg(y' \succ a))$$

affectation descendante Ce mode d'affectation est le pendant du mode précédent, les catégories sont considérées de la meilleure à la plus mauvaise et l'affectation consiste à chercher un profil inférieur auquel a est préféré sans pour autant que a ne soit préféré aux profils supérieurs, ce qui se traduit par :

$$(a \in C_t) \equiv (\exists y \in Y_t^{inf}, a \succ y) \wedge (\forall y' \in Y_t^{sup}, \neg(a \succ y'))$$

Les méthodes des filtrage par préférence se doivent de respecter la propriétés suivantes définie par Perny [81] :

Propriété 9 (affectation ordonnée) Soit un entier $t^* \in [1, q]$ défini tel que $\mu_{t^*}(a) = \max_{t=1}^q \mu_t(a)$, les degrés d'affectation sont ordonnés de la manière suivante :

- si $j \leq t^*$ alors $\mu_{j-1}(a) \leq \mu_j(a)$
- si $j \geq t^*$ alors $\mu_{j+1}(a) \leq \mu_j(a)$

Cette propriété permet d'éviter qu'une action ne soit affectée à deux catégories non successives sans être affectée à une catégorie intermédiaire. En effet, si l'on considère trois catégories C_1 , C_2 et C_3 ayant pour signification bon, moyen et mauvais cela n'aurait aucun sens d'affecter à C_1 et à C_3 sans affecter à C_2

Propriété 10 (Sens de variation de f_t) La fonction f_t est :

- croissante de $\succ(a, y^-)$ et de $\succ(y^+, a)$ avec $y^- \in Y^{inf}$ et $y^+ \in Y^{sup}$,
- décroissante de $\succ(a, y^+)$ et de $\succ(y^-, a)$ si y^- et y^+ sont des profils limites respectivement inférieure et supérieures de C_t .

3.4.4 Filtrage par indifférence

Modèle classique

Dans le cadre du filtrage par indifférence les catégories sont représentées par des profils centraux. Formellement une catégorie C_t sera représentée par un ensemble de prototypes Y_t . Les catégories ne possèdent pas de structure particulière (aucune condition d'ordre contrairement au filtrage par préférence).

L'affectation aux catégories est effectuée en utilisant une relation d'indifférence $\sim$. Le principe d'affectation consiste à affecter une action à une catégorie dès lors que celle-ci sera indifférente à un prototype de la catégorie. Formellement ce principe peut être représenté par l'équation logique suivante :

$$(a \in C_t) \equiv (\exists y \in Y_t, a \sim y) \quad (3.14)$$

Dans le cadre du filtrage par indifférence, l'équation 3.13 s'exprime plus simplement :

$$\mu_t(a) = f_t(\sim(a, y_1), \dots, \sim(a, y_m)) \quad (3.15)$$

Il n'est pas nécessaire de supposer la symétrie de la relation $\sim$. En effet ce qui est recherché est l'indifférence entre l'action et le point de référence et non pas l'indifférence entre un point de référence et l'action. La relation est orientée. Pour s'en convaincre, il suffit de regarder la classification des mesures de similitude donnée par [92]. Celle-ci présente des mesures de similitude comme les mesure de satisfiabilité ou d'inclusion non symétriques. Dans le cas d'une relation d'inclusion, on va rechercher si une action possède les propriétés suffisantes lui permettant d'être affectée à une classe définie par le point de référence.

Nous pouvons ré-écrire la fonction d'affectation de la manière suivante en divisant Y en deux sous-ensemble tels que $Y = Y_t \cup Y_{\neg t}$ avec $Y_{\neg t} = \bigcup_{i \in \{1, \dots, q\} \setminus \{t\}} Y_i$:

$$\mu_t(a) = f_t(\sim_{Y_t}(a), \sim_{Y_{\neg t}}(a)) \quad (3.16)$$

avec :

$$\sim_{Y_t}(a) = (\sim(a, y_1), \dots, \sim(a, y_{|Y_t|})) \text{ avec } y_i \in Y_t$$

$$\sim_{Y_{\neg t}}(a) = (\sim(a, y'_1), \dots, \sim(a, y'_{|Y_{\neg t}|})) \text{ avec } y'_i \in Y_{\neg t}$$

L'indifférence avec les points de référence de la catégorie C_t pouvant être interprétée comme des raisons positives à l'appartenance et l'indifférence avec les points de référence des autres catégories comme des raisons négatives à l'appartenance.

Ce qui nous permet de formaliser la propriété de *neutralité des points de référence* dans le cas du filtrage par indifférence. La fonction d'affectation f_t va être définie telle que :

$$\begin{aligned} f_t(\sim_{Y_t}(a), \sim_{Y_{\neg t}}(a)) &= (x_1, \dots, x_m, x'_1, \dots, x'_{m'}) \\ &= (x_{\sigma(1)}, \dots, x_{\sigma(m)}, x'_{\sigma'(1)}, \dots, x'_{\sigma'(m')}) \end{aligned}$$

avec σ une permutation quelconque de $\{1, \dots, m\}$ et σ' une permutation quelconque de $\{1, \dots, m'\}$.

Propriété 11 (Sens de variation de f_t) *La fonction f_t est croissante de $\sim(a, y)$ si $y \in Y_t$ et décroissante de $\sim(a, y')$ si $y' \notin Y_t$*

Lors du calcul de μ_t , il peut être utile de s'intéresser aux points de référence limites des catégories autres que C_t . Sous l'hypothèse de **catégories exclusives**, la similitude avec ceux-ci peut être une raison négative pour l'appartenance à C_t . Dès lors que l'indifférence avec un point de référence d'une autre catégorie est déjà établie, l'appartenance à C_t devra être remise en cause quelle que soit le résultat des comparaisons avec les points de Y_t .

Propriété 12 (Catégories Exclusives) *Soit $\sim$ une relation d'indifférence, alors si C_r et C_t sont des catégories exclusives :*

$$f_t(\sim(a, y_1), \dots, \sim(a, y_m)) = 0$$

dès lors que : $\exists y_i \in Y_{\neg t}, (\sim(a, y_i) = 1)$

Il s'agit là d'un comportement pessimiste. Dans le cas où une action a serait indifférente à un prototype de C_t , l'appartenance de a à C_t pourra être remise en cause si a est aussi indifférente à un point de référence d'une autre catégorie exclusive de C_t . Cette attitude permet de ne pas prendre le risque d'une affectation pouvant être remise en cause par l'appartenance à une autre catégorie exclusive de C_t .

Exploitation des points de référence excluants

Nous proposons dans cette section un moyen d'exploiter des catégories décrites par des points de référence centraux et des points de référence excluants. Il s'agit d'une extension du modèle d'affectation de filtrage par indifférence. Nous supposons donc ici qu'une catégorie C_t est décrite à l'aide d'un ensemble de points de référence $Y_t = Y_t^- \cup Y_t^+$. Le principe consiste à affecter une action a à une catégorie C_t dès lors que celle-ci est indifférente à un point de référence de Y_t^+ et qu'elle n'est indifférente à aucun anti-prototype de Y_t^- . Nous avons donc la règle suivante :

$$(a \in C_t) \equiv (\exists y \in Y_t^+, a \sim y) \wedge \neg(\exists y' \in Y_t^-, a \sim y') \quad (3.17)$$

Cette estimation peut être vue comme la combinaison des raisons de l'affectation et des raisons de la non-affectation. On considère pour cela deux évaluations : μ_t^+ l'évaluation de l'indifférence avec les points de Y_t^+ et μ_t^- l'évaluation de l'indifférence avec les points de Y_t^- . Afin de considérer ces deux évaluations, il est possible dans un premier temps de les agréger directement afin d'obtenir un degré d'affectation global. Ainsi, nous pouvons calculer l'affectation global par :

$$\mu_t(a) = T(\mu_t^+(a), 1 - \mu_t^-(a)) \quad (3.18)$$

avec T une t-norme qui traduit le **et** $\wedge$ de l'équation 3.17 et où $\mu_t^+(a)$ et $\mu_t^-(a)$ sont définis par :

$$\mu_t^+(a) = S_{y \in Y_t^+}(\sim(a, y)) \quad (3.19)$$

$$\mu_t^-(a) = S_{y \in Y_t^-}(\sim(a, y)) \quad (3.20)$$

avec S une co-norme qui traduit dans une logique multivalente le quantificateur $\exists$ de l'équation 3.17.

Cependant il est possible d'avoir une autre attitude. L'utilisation de tels points de référence peut engendrer différentes situations dont certaines sont complexes et dont l'interprétation n'est pas immédiate. L'indifférence avec les points de référence centraux peut être interprétée comme des raisons positives à l'appartenance à la catégorie, ce sont des raisons qui confirment l'appartenance. Au contraire, l'indifférence avec les points de référence excluants peut être interprétée comme des raisons négatives à l'appartenance. Ainsi, il est possible d'avoir uniquement des raisons positives, auquel cas il va être possible de conclure à l'appartenance. Il est aussi possible de n'avoir que des raisons négatives dans ce cas, la non-appartenance à la catégorie va être validée. Enfin il est encore possible d'avoir à la fois des raisons positives et négatives ce qui aboutit à une situation contradictoire ou de n'avoir aucune raison, ni positive ni négative, auquel cas il s'agit d'un manque d'information.

Une manière simple de gérer ces différentes situations serait de considérer que seule la première situation, la situation d'affectation, donne lieu à une affectation et que les trois autres donnent lieu à une non-affectation. Cependant cela se traduit par une information beaucoup plus pauvre. Une situation d'information contradictoire est notoirement différente d'une situation de manque d'information. Dans un cas, aucun élément de réponse n'est disponible alors que dans l'autre il existe des raisons qui concourent à l'affectation même si leur prise en compte doit être modérée par l'existence de raisons négatives. Il semble alors réducteur de vouloir traduire de la même manière, uniquement par une non affectation, ces différentes situations. Il nous semble donc judicieux de pouvoir tenir compte de la contradiction et du manque d'information. Pour cela, nous envisageons une modélisation fondée sur une extension continue d'une logique à quatre valeurs, voir [83]. Considérant les quatre valeurs **Vrai**, **Faux**, **Contradictoire** et **Inconnu**, les deux premières valeurs correspondent aux situations d'affectation et de non-affectation où l'information apparaît comme étant de bonne qualité. Les deux valeurs suivantes permettent de gérer une information de mauvaise qualité. La prise en compte de ces imperfections de l'information peut être effectuée en mesurant d'une part la force de l'affectation et de la non affectation et d'autre part la qualité de l'information (contradiction ou inconnu). Pour cela, nous proposons de fonder l'analyse de l'affectation d'une action par l'évaluation de quatre degrés, un degré d'affectation, un degré de non affectation, un degré de contradiction et un degré d'inconnu.

Introduisons maintenant les notations suivantes afin de formaliser ces situations. Soit p une proposition, nous posons un prédicat $\mathcal{B}$, où $\mathcal{B}(p)$ qui signifie *j'ai des raisons de penser que p est vrai*. Considérant les deux propositions " a est affectée à C_t " et " a n'est pas affectée à C_t ", on pose maintenant :

$$\mathcal{B}(a \in C_t) \equiv \exists y \in Y_t^+ / a \sim y \quad (3.21)$$

$$\mathcal{B}(a \notin C_t) \equiv \exists y \in Y_t^- / a \sim y \quad (3.22)$$

On évalue $\mathcal{B}(a \in C_t)$ et $\mathcal{B}(a \notin C_t)$ respectivement par $\mu_t^+(a)$ et $\mu_t^-(a)$.

Notre objectif va être d'exprimer les 4 situations évoquées précédemment à l'aide de 4 degrés, $\mu_t^T(a)$, $\mu_t^F(a)$, $\mu_t^K(a)$ et $\mu_t^U(a)$ traduisant respectivement l'affectation, la non-affectation, la part d'information contradictoire et la part d'inconnu de l'évaluation de l'appartenance de a à la catégorie C_t . On suppose 4 t-normes T_1 , T_2 , T_3 et T_4 non nécessairement identiques.

Degré d'affectation Tout d'abord a peut être indifférente à au moins un point de référence central de C_t sans être indifférente aux anti-prototypes. L'indifférence avec les points de référence apparaît alors comme une *raison positive* de l'appartenance de l'action a à C_t . Cette situation peut être modélisée par :

$$\mathcal{B}(a \in C_t) \wedge \neg \mathcal{B}(a \notin C_t) \quad (3.23)$$

On définit maintenant le degré d'évaluation de l'appartenance par :

$$\mu_t^T(a) = T_1(\mu_t^+(a), (1 - \mu_t^-(a))) \quad (3.24)$$

Degré de non-affectation L'action a peut aussi être indifférente à au moins un anti-prototype de C_t sans être indifférente aux points de référence centraux de C_t . L'indifférence avec les anti-prototypes est alors une *raison négative* de l'appartenance à C_t . Cette situation se traduit par :

$$\neg \mathcal{B}(a \in C_t) \wedge \mathcal{B}(a \notin C_t) \quad (3.25)$$

On définit le degré d'évaluation de la non-appartenance par :

$$\mu_t^F(a) = T_2((1 - \mu_t^+(a)), \mu_t^-(a)) \quad (3.26)$$

Degré de contradiction Il est aussi possible d'être confronté à une situation contradictoire, c'est le cas si l'action à classer est simultanément indifférente à un points de référence central et à un anti-prototype. Il y a ici à la fois des *raisons positives* et des *raisons négatives* de l'appartenance de a à C_t . Cette situation de contradiction se traduit par :

$$\mathcal{B}(a \in C_t) \wedge \mathcal{B}(a \notin C_t) \quad (3.27)$$

Le degré de contradiction est alors défini par :

$$\mu_t^K(a) = T_3(\mu_t^+(a), \mu_t^-(a)) \quad (3.28)$$

Degré d'inconnu Enfin, il est aussi possible que l'action a ne soit indifférente à aucun point de référence. Il n'y a ni *raison positive* ni *raison négative* de l'appartenance. Cette situation se traduit par :

$$\neg \mathcal{B}(a \in C_t) \wedge \neg \mathcal{B}(a \notin C_t) \quad (3.29)$$

Le degré d'inconnu est alors défini par :

$$\mu_t^U(a) = T_4((1 - \mu_t^+(a)), (1 - \mu_t^-(a))) \quad (3.30)$$

Il peut sembler naturel de respecter $\mu_t^T(a) + \mu_t^F(a) + \mu_t^K(a) + \mu_t^U(a) = 1$, on peut alors par exemple utiliser la t-norme produit pour T_1 , T_2 , T_3 et T_4 .

Si l'on considère de manière plus restrictive que l'information sur l'affectation ne peut être simultanément contradictoire et manquante ($\forall (a, t) \in A \times \{1, \dots, q\}, \min(\mu_t^U(a), \mu_t^K(a)) = 0$) alors Perny et Tsoukiàs [83] propose l'utilisation des t-normes suivantes :

$$\mu_t^T(a) = \min(\mu_t^+(a), 1 - (\mu_t^-(a))) \quad (3.31)$$

$$\mu_t^F(a) = \min(1 - (\mu_t^+(a)), \mu_t^-(a)) \quad (3.32)$$

$$\mu_t^K(a) = \max(\mu_t^+(a) + \mu_t^-(a) - 1, 0) \quad (3.33)$$

$$\mu_t^U(a) = \max(1 - \mu_t^+(a)) - (\mu_t^-(a)), 0) \quad (3.34)$$

L'utilisation de ces 4 évaluations permet alors d'émettre un jugement nuancé sur l'affectation de l'action a à la catégorie C_t . Cela peut aussi bien permettre de rejeter l'affectation en situation de contradiction ou d'information manquante mais aussi de pouvoir éventuellement réviser la représentation des catégories.

3.4.5 Filtrage combiné

Enfin, nous pouvons envisager des problèmes de classification plus complexes où les catégories sont représentées à la fois par des points de référence limites et par des prototypes. Dans ce cas la méthode de classification utilise à la fois des relations de préférence et d'indifférence ($\succ$ et $\sim$). On peut distinguer plusieurs raisons :

Raisons structurelles Tout d'abord, il se peut que certaines catégories soient ordonnées et pas d'autres. Auquel cas les catégories ordonnées sont représentées par des profils limites et les catégories *incomparables* par des profils centraux.

Exemple 24 *Si l'on reprend l'exemple de la classification des automobiles, on peut envisager plusieurs catégories dont certaines sont ordonnées et d'autres pas. On retrouve des telles catégories dans la presse spécialisée [21], avec une classification en 20 catégories. Certaines peuvent être facilement ordonnées comme les 4 catégories représentant des voitures plus ou moins sportives : voitures des copains, GTI, familiales sportives et ultra-sportives. Ces voitures sont classées suivant leur caractère sportif. En revanche d'autres ne peuvent être ordonnées comme les catégories cabriolets, monospaces et tout terrains.*

Raisons pratiques D'un point de vue opérationnel, il peut exister des raisons pour représenter certaines catégories par des profils limites et d'autres par des profils centraux. Même si les catégories ne sont pas ordonnées, il peut être plus aisé de représenter une catégorie par des profils limites que par des profils centraux. C'est le cas si l'on est en présence de catégories correspondant à une partie ouverte de l'espace des critères (ou même une partie importante de l'espace des critères). Il est plus pratique de représenter une catégorie par quelques profils limites que par une multitude de profils centraux. L'utilisation de profils limites permet une économie en terme de profils (et donc aussi du nombre de comparaisons nécessaires pour asseoir une affectation ce qui améliore la rapidité d'exécution) et une définition plus aisée : il va être plus facile pour un décideur de fixer quelques bornes plutôt qu'une multitude de prototypes. De plus les catégories *ouvertes* ne peuvent être représentées par des profils centraux, elles doivent être représentées soit par des bornes inférieures soit par des bornes supérieures. Les profils limites apparaissent alors plus comme des contraintes que l'on place sur les catégories.

Exemple 25 *Si l'on reprend l'exemple de objectifs photo, représenter une catégorie téléobjectif est beaucoup plus aisé à l'aide de profils limites que de profils centraux du point de vue de la distance focale. On considère généralement que les téléobjectifs ont une distance focale supérieure à 70 mm, l'échelle peut être considérée comme ouverte : il est beaucoup plus simple de considérer une borne inférieure que de décrire l'étendue de l'échelle.*

Considérant des problèmes d'affectation où les catégories sont représentées à la fois par des profils limites et centraux, la règle d'affectation apparaît comme une combinaison des règles exposées en filtrage par préférence et en filtrage par indifférence. Le filtrage par préférence étant appliqué aux catégories décrites par des points de référence limites et le filtrage par indifférence aux catégories décrites par des prototypes.

Exemple 26 (affectation par filtrage combiné) *Voyons un exemple de règle d'affectation faisant intervenir des profils limites et centraux à travers une affectation combinée par préférence et par indifférence. Soit une catégorie C_t définie comme le regroupement de deux sous-catégories $C_t^\sim$ et $C_t^\succ$ décrites pour l'une par un profil central y et pour l'autre par deux profils respectivement inférieur et supérieur y^{inf} et y^{sup} . Soient $\mu_t^\sim$ le degré d'affectation à la première sous-catégorie calculé selon filtrage par indifférence avec le profil y et $\mu_t^\succ$ le degré d'affectation à la seconde sous-catégorie*

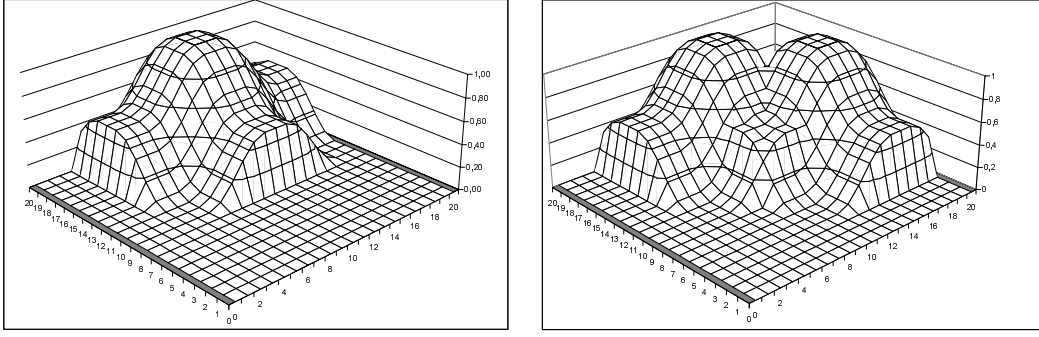


FIG. 3.6 – Filtrage par indifférence avec catégories monopprofil et multiprofil

calculé selon filtrage par indifférence avec les profils limites. Considérons comme règle d'affectation que une action a est affectée à C_t si elle est affectée soit à $C_t^\sim$ soit à $C_t^\succ$ on a donc l'équation logique :

$$(a \in C_t) \equiv (a \in C_t^\sim) \vee (a \in C_t^\succ)$$

ainsi le degré d'affectation à C_t peut être défini par :

$$\mu_t(a) = S(\mu_t^\sim(a), \mu_t^\succ(a))$$

avec S une co-norme, si l'on utilise la co-norme idempotente la catégorie C_t peut être représentée par partie droite de la figure 3.7.

3.4.6 Représentation graphique

Afin de mettre en évidence certaines différence entre les trois types de filtrage que nous venons de présenter, il est possible de visualiser graphiquement comment une catégorie est représentée. Pour cela nous supposons un problème bi-critère, où chaque critère prend ses valeurs dans l'intervalle $[0, 20]$. Nous présentons ensuite sur un troisième axe (dans l'intervalle $[0, 1]$) la valeur de la fonction d'affectation à une catégorie. Nous présentons le cas du filtrage par indifférence à la figure 3.6. La figure de gauche représente une catégorie monopprofil et la figure de droite une catégorie multiprofil (2 prototypes). Dans ce type de filtrage les catégories sont définies comme des zone d'"indifférence" autour des prototypes. La figure 3.7 présente à gauche le cas d'une catégorie définie par profils limites. Dans ce cas le catégories sont définies comme des zones de "préférence" par rapport aux profils. Le cas du filtrage combiné, est présenté dans la partie droite de la figure 3.7. Il s'agit ici d'une règle d'affectation qui combine par une disjonction une affectation par profil central et une affectation par profil limite, voir exemple 26.

3.5 Méthodes de type Plus Indifférents Prototypes

3.5.1 Introduction

On suppose que l'on distingue q catégories $C_1, \dots, C_q$ non-nécessairement disjointes. Chaque catégorie C_t est définie par un ensemble de prototypes Y_t . Les éléments de Y_t peuvent être construits ou appris en collaboration avec le décideur ou l'expert et figurent le ou les profils typiques qu'il faut

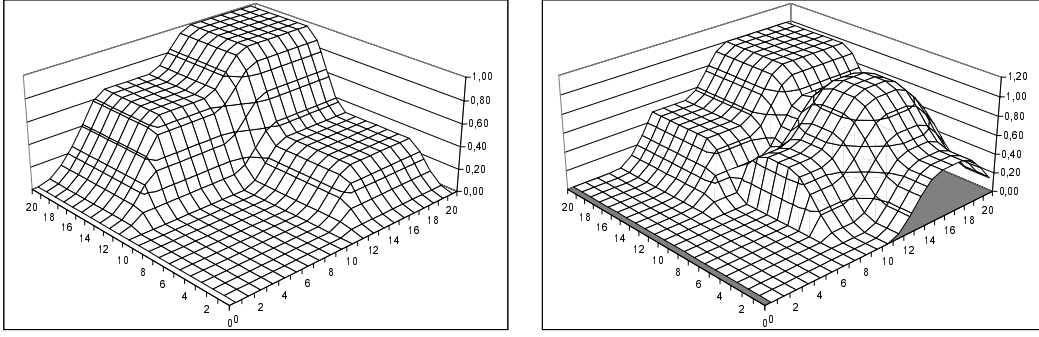


FIG. 3.7 – Filtrage par préférence et combiné

présenter pour entrer dans la catégorie C_t . On note Y l'ensemble de tous les points de références, $Y = Y_1 \cup \dots \cup Y_q$. Nous notons $\mu_t^o(y)$ le degré de représentation du prototype y de la catégorie C_t . Ce degré correspond au degré d'affectation que l'on devrait obtenir si l'on cherchait à affecter le prototype y à C_t . Par ailleurs, on dispose d'un ensemble d'actions $A = \{a_1, a_2, \dots\}$ à classer. Toute action a est caractérisé par un point $x = (x_1, \dots, x_n)$ de E où $x_j = g_j(a)$ représente la performance caractérisant l'action a du point de vue du $j^{\text{ème}}$ critère. On note $X = \{x_1, x_2, \dots\}$ l'image de A dans l'espace des critères.

3.5.2 Filtrage flou par indifférence

Le *filtrage flou par indifférence* consiste à comparer les éléments de X et ceux de Y sur la base de la relation d'indifférence floue construite au §3.3. Le principe d'affectation que nous utilisons est une généralisation de l'algorithme des k plus proches voisins flous [58] exploitant la notion de voisinage à travers la relation $\sim$. Une action a sera affectée à la catégorie C_t si et seulement si le proche voisinage de a dans X contient des éléments de Y_t (cf. figure 3.8). Formellement si $\sim$ est une relation d'indifférence floue définie sur $X \times Y$ l'équation logique 3.14 (p. 65) se traduit par :

$$\mu_t(a) = \begin{cases} S_{y \in V_k(a) \cap Y_t} \{T(\sim(a, y), \mu_t^o(y))\} & \text{si } V_k(a) \cap Y_t \neq \emptyset, \\ 0 & \text{sinon.} \end{cases} \quad (3.35)$$

où S est une co-norme, T une t-norme, $V_k(a)$ désigne l'ensembles des k plus indifférents prototypes de a dans Y , c'est-à-dire les k premiers points $y \in Y$ maximisant $\sim(a, y)$.

Le choix de la co-norme S dépend de l'effet de synergie positive que l'on désire faire jouer entre les voisins du point à classer. L'utilisation de la co-norme idempotente ne fait jouer aucun effet de synergie alors que tout autre co-norme va faire jouer un effet de synergie. Ainsi on peut distinguer :

Le Plus Indifférent Prototype (méthode PIP) ($S(x, y) = \max(x, y)$)

La co-norme idempotente ne fait jouer aucun effet de synergie, on calcule le degré d'appartenance d'une action a à la catégorie C_t comme le plus grand degré d'indifférence entre l'action à affecter et les points de référence de la catégorie C_t figurant parmi ses k plus indifférents voisins. Il suffit donc d'être indifférent à un seul point de référence de la catégorie pour en faire partie, l'indifférence avec d'autres points de référence ne va pas influencer sur le degré d'affectation. On pose alors :

$$\mu_t(a) = \begin{cases} \max_{y \in V_k(x) \cap Y_t} \{T(\sim(a, y), \mu_t^o(y))\} & \text{si } V_k(a) \cap Y_t \neq \emptyset, \\ 0 & \text{sinon.} \end{cases} \quad (3.36)$$

Données :

- a une action à classer,
- q catégories $C_1, \dots, C_q$ représentées par les ensembles de prototypes $Y = Y_1 \cup \dots \cup Y_q$,
- $\sim$ une relation d'indifférence,
- k le nombre de plus indifférents prototypes,
- $V_k(a) = \{v_1, \dots, v_k\}$ la liste **ordonnée** des k plus indifférents prototypes (v_1 est le voisin le plus indifférent à a).

Résultat : un vecteur de degrés d'affectation $\mu(a) = (\mu_1(a), \dots, \mu_q(a))$.

DÉBUT

i=1

RÉPÉTER JUSQU'A (obtention des k plus indifférents prototypes)

calculer $\sim(a, y_i)$

SI ($i \leq k$)

ALORS intégrer y_i dans $V_k(a)$

SINON

SI (si $\sim(a, y_i) > \sim(a, v_k)$)

ALORS intégrer y_i dans $V_k(a)$

FINSI

FINSI

i = i + 1

FIN RÉPÉTER JUSQU'A

POUR t **ALLANT DE** 1 **à** q

SI ($V_k(a) \cap Y_t = \emptyset$)

ALORS $\mu_t(a) = 0$

SINON $\mu_t(a) = S_{y \in V_k(a) \cap Y_t} \sim(a, y)$

FINSI

FINPOUR

FIN

FIG. 3.8 – *Algorithme des k plus indifférents prototypes*

Les k Plus Indifférents Prototypes (méthode k -PIP) ($S(x, y) = x + y - xy$)

L'utilisation d'une co-norme différente de la co-norme idempotente fait opérer une synergie positive entre les points de $V_k(a)$, il va permettre de renforcer l'appartenance à une catégorie, si l'individu à classer possède dans la liste de ses k plus indifférents prototypes, plusieurs points de référence d'une même catégorie. On pose alors :

$$\mu_t(a) = \begin{cases} 1 - \prod_{y \in V_k(a) \cap Y_t} (1 - T(\sim(a, y), \mu_t^o(y))) & \text{si } V_k(a) \cap Y_t \neq \emptyset, \\ 0 & \text{sinon.} \end{cases} \quad (3.37)$$

Remarque 5 Pour un même jeu de données, le résultat de l'affectation par la méthode k -PIP produit toujours un degré d'affectation au moins supérieur au résultat de la méthode PIP. En effet la méthode PIP utilise la co-norme idempotente qui respecte $\forall(x, y) \in [0, 1]^2, \max(x, y) \leq S(x, y)$ avec S une co-norme quelconque. Si l'on note $\mu_t^k(a)$ l'affectation par la méthode k -PIP et $\mu_t^1(a)$ l'affectation par la méthode PIP, on a :

$$\max_{y \in V_k(a) \cap Y_t} \{T(\sim(a, y), \mu_t^o(y))\} \leq S_{y \in V_k(a) \cap Y_t} \{T(\sim(a, y), \mu_t^o(y))\} \quad (3.38)$$

$$\mu_t^1(a) \leq \mu_t^k(a) \quad (3.39)$$

Lorsque nous avons définies les catégories, nous avons supposé que celles-ci étaient représentées par des points de référence centraux. Deux types de points de référence peuvent être distingués :

prototypes Les prototypes sont des objets typiques de la catégorie. Ils représentent les profils typique qu'une action doit présenter pour pouvoir être affectée à une catégorie. Une catégorie sera donc représentée par autant de prototypes qu'elle présente de profils. Généralement même les catégories complexes et multiprofiles ne nécessitent qu'un nombre réduit de prototypes pour leur description. De plus, de part le caractère multiprofil, les prototypes n'ont pas tendance à être proches les uns des autres, au contraire ils ont plutôt tendance à être contrastés.

points exemples Au contraire les points exemples ne correspondent pas nécessairement exactement à des profils présentés par les catégories. En effet, le décideur peut ne pas être à même d'isoler et de spécifier les profils et donc les prototypes des catégories. Il peut donc être amené à caractériser une catégorie non pas par des prototypes mais juste par des points exemples qui vont correspondre à une liste de points qui méritent juste d'entrer dans la catégorie mais qui ne possèdent pas de caractère prototypique. Ces points peuvent provenir d'un historique, d'objets déjà classés par un expert, d'objets dont l'appartenance aux catégories a été connue *a posteriori* ou même éventuellement d'une classification effectuée par une autre méthode. De par leur provenance et contrairement aux prototypes ces points exemples peuvent se retrouver en grand nombre pour décrire une catégorie. De plus, il est tout à fait possible que plusieurs points exemples soient très proches (toujours contrairement aux prototypes).

Cette distinction des types de points de référence peut être mise en parallèle avec le type de méthode (PIP ou k -PIP). Ainsi la méthode PIP (dite méthode sans renforcement) va être plus adaptée à des problèmes où les catégories sont décrites par des prototypes. En effet, lorsque les points de référence correspondent à des profils typiques que l'on doit présenter pour intégrer la catégorie, il va suffire d'être indifférent à un seul prototype pour asseoir l'appartenance. De plus comme les prototypes correspondent à des profils différents que l'on doit présenter, un objet à classer ne sera généralement pas indifférent à plusieurs prototypes. Si l'on reprend l'exemple de la catégorie des voitures sportives représentées par trois prototypes : une GTI (ex. Peugeot 206 S16), une familiale sportive (ex. Audi S4) et une sportive GT (ex. Porsche 911) il semble difficile qu'une

automobile à classer soit indifférente à plusieurs de ces trois prototypes. En revanche, lorsque les catégories sont décrites par des points exemples la méthode k -PIP (dite avec renforcement) va être plus adaptée. En effet, dans ce cas le nombre de points de référence peut être assez important et il existe généralement plusieurs points de référence *proches* les uns des autres auquel cas l'indifférence avec plusieurs points de référence va renforcer l'appartenance à la catégorie.

3.5.3 Amélioration de la rapidité du calcul des plus indifférents prototypes

Afin de réduire les temps de calcul et donc d'améliorer l'efficacité et le rendement de notre méthode nous proposons une approche inspirée de celle proposée par Kittler [61] et par Naranjo *et al.* [72] utilisée pour améliorer la procédure de recherche des k -ppv. Le principe est de réduire le nombre de points de référence candidats à l'accès à l'ensemble V_k par l'utilisation d'une autre mesure de proximité que celle utilisée dans la procédure d'affectation. L'utilisation de cette autre mesure permet la construction d'un sous-ensemble réduit de points de référence candidats à V_k , la recherche de V_k est alors effectuée sur ce sous-ensemble. Ainsi, en moyenne la construction du sous-ensemble réduit et la recherche de V_k dans ce sous-ensemble est elle plus rapide que la recherche de V_k dans l'ensemble de départ.

Nous proposons d'effectuer une sélection de points candidats à partir du calcul de la discordance. Les prototypes qui sont complètement discordants ne pourront jamais être indifférents. La discordance agit comme un majorant du degré d'indifférence car celui-ci est calculée par la t -norme de la concordance et de la négation de la discordance ($\sim(a, b) = T(C^\sim(a, b), 1 - D^\sim(a, b)) \leq 1 - D^\sim(a, b)$). Ainsi tous les points de référence discordants peuvent être éliminés avant le calcul de la concordance. Cette approche est d'autant plus adaptée qu'elle n'induit pas de surcoût contrairement à la méthode de [61]. L'indice de discordance doit nécessairement être calculé alors que la méthode de Kittler nécessitait le calcul d'une distance supplémentaire.

3.6 Explication de l'affectation

3.6.1 Structure des explications

Dans le cadre de l'aide à la décision, il est aussi important si ce n'est de plus de construire un résultat satisfaisant que de pouvoir l'expliquer. Nous avons distingué au chapitre 2 les méthodes possédant un caractère explicatif de celles qui en sont dépourvues. Nous montrons dans cette section le caractère explicatif de notre approche du filtrage flou. Nous proposons donc un moyen de construire automatiquement une explication de l'affectation des actions aux catégories. Cette explication est fondée sur l'interprétation des différents degrés utilisés dans la phase de calcul de la méthode.

Tout d'abord nous pouvons distinguer deux types d'explication :

Explications positives Ce sont celles qui vont permettre de justifier l'affectation à une catégorie.

Ce type d'explication donne les raisons pour lesquelles une action est affectée à une catégorie.

Explications négatives Ce sont les explications qui justifient qu'une action n'appartient pas à une catégorie. Elles donnent les raisons du rejet de l'affectation.

Après avoir distingué les deux types d'explications précédents nous distinguons trois niveaux d'explications. Les niveaux d'explication correspondent aux différentes étapes de l'agrégation de la

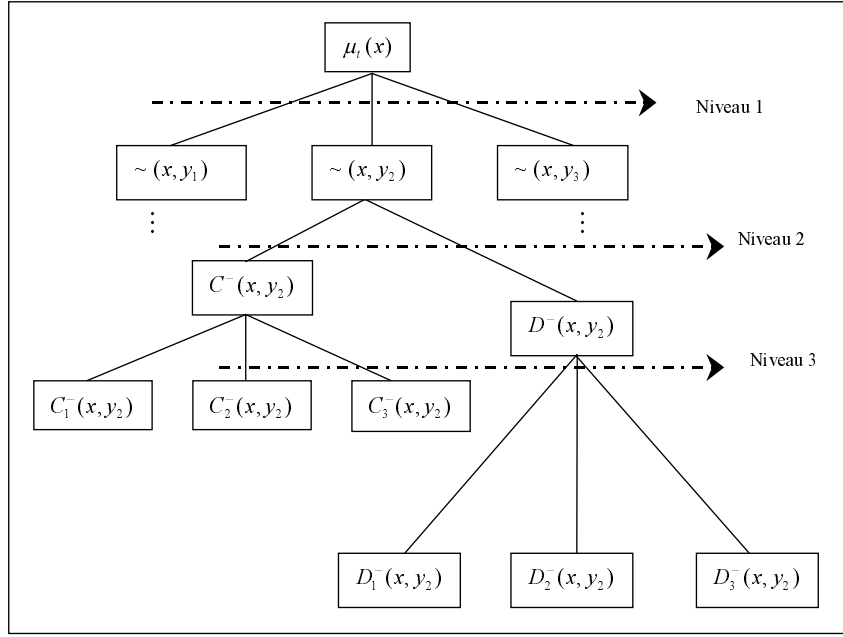


FIG. 3.9 – *Arbre de calcul*

procédure d'affectation. Nous utilisons cinq degrés pour générer des explications (*cf.* figure 3.9), ceux-ci sont regroupés selon les trois niveaux suivants :

- 1^{er} niveau** Ce sont les explications engendrées à partir des degrés d'indifférence ($\sim$),
- 2^{ème} niveau** Ce sont les explications engendrées à partir des degrés de concordance et de discordance globale ($C^\sim$ et $D^\sim$),
- 3^{ème} niveau** Ce sont les explications engendrées à partir des degrés de concordance et discordance monocritères ($C_j^\sim$ et $D_j^\sim, \forall i \in \{1, \dots, n\}$).

Ces explications doivent être vue comme des explications emboîtées. Les explications de premier niveau expliquent l'affectation, les explication de deuxième niveau expliquent l'indifférence entre deux objets et les explications de troisième niveau apportent des précisions sur les degrés de concordance et discordance globale.

Voyons maintenant à travers un exemple comment les explications peuvent être engendrées. Les explications sont engendrées à partir de degrés de l'intervalle $[0, 1]$, une discrétisation de cet intervalle va être nécessaire. Nous faisons les hypothèses suivantes en notant $x = g(a)$:

$$\begin{aligned}
 \mu_t(a) &= \max_{y \in Y_t} \sim (x, y) \\
 \sim (x, y) &= \min(C^\sim(x, y), 1 - D^\sim(x, y)) \\
 C^\sim(x, y) &= \sum_j \omega_j C_j^\sim(x, y) \\
 D^\sim(x, y) &= \max_j D_j^\sim(x, y) \\
 0 &< \alpha < \beta < 1
 \end{aligned}$$

Les valeurs α et β servent à la discrétisation de l'intervalle $[0, 1]$ et sont interprétées de la manière

suivante :

- L'intervalle $[0, \alpha]$ est associé au terme *faible*, la situation $0 < \mu_t(a) \leq \alpha$ permet d'engendrer l'explication suivante *a est faiblement affecté à la catégorie C_t* .
- L'intervalle $]\alpha, \beta]$ est associé au terme *moyen*, la situation $\alpha < \mu_t(a) < \beta$ permet d'engendrer l'explication suivante *a est moyennement affecté à la catégorie C_t* .
- L'intervalle $]\beta, 1]$ est associé au terme *fort*, la situation $\beta < \mu_t(a) \leq 1$ permet d'engendrer l'explication suivante *a est fortement affecté à la catégorie C_t* .

La discrétisation de l'intervalle et les valeurs α et β doivent être définies conjointement avec le décideur suivant le niveau de détail que celui-ci recherche pour le problème de classification considéré.

3.6.2 Les différents types d'explication

Nous présentons dans cette section les différentes explications qui peuvent être engendrées.

Explication de premier niveau

L'indifférence avec un prototype L'explication à ce niveau peut être soit positive soit négative.

Nous pouvons engendrer des explications comme :

- *explication positive* : si $\exists y \in Y_t / \sim (a, y)$ alors $(a \in C_t)$ et nous avons l'explication *a est affectée à C_t car a est fortement indifférent à y_j* .
- *explication négative* : si $\forall y \in Y_t / \neg \sim (a, y)$ alors $(a \notin C_t)$ et nous avons l'explication *a ne peut être affecté à C_t car a n'est indifférent à aucun prototype de C_t* .

Nous pouvons distinguer trois principales situations qui engendrent les explications suivantes :

- La situation $(\mu_t(a) > \beta)$: $\exists y \in Y_t / \sim (x, y) > \beta$ permet de produire l'explication : *a est fortement affecté à C_t car a est fortement indifférent à y*.
- La situation $(\alpha < \mu_t(a) < \beta)$: $\forall y \in Y_t / \sim (x, y) < \beta$ et $\exists z \in Y_t / \sim_X (x, z) > \alpha$ permet de produire l'explication : *a est moyennement affecté à C_t car a est moyennement indifférent à y sans être fortement indifférent à aucun autre prototype de C_t* .
- La situation $(\mu_t(a) < \alpha)$: $\forall y \in Y_t / \sim (x, y) < \alpha$ permet de produire l'explication *a est faiblement affecté à C_t car a est au mieux faiblement indifférent à y*.

Exemple 27 (Explication de premier niveau) Si la Renault Clio 16 V est un prototype de la catégorie sportive peut générer l'explication : La Peugeot 206 S16 est fortement affectée aux voitures sportives car vous êtes fortement indifférent entre elle et la Renault Clio 16 V.

Explication de deuxième niveau

Les explications de deuxième niveau font intervenir la concordance et la discordance globale. Elles doivent donc être utilisées conjointement pour expliquer l'indifférence entre un objet et un prototype.

La concordance globale Tous comme au niveau précédent, nous pouvons produire des explications positives et négatives. Ce type d'explication permet d'expliquer la force de l'indifférence

entre un objet et un prototype.

- Explication positive: si $\exists y \in Y_t / (C^\sim(x, y) > \beta) \wedge (D^\sim(x, y) < \alpha)$ nous pouvons produire l'explication *a peut être affecté à la catégorie C_t car a et y sont indifférents sur la majorité des critères.*
- Explication négative: si $\forall y \in Y_t / C^\sim(x, y) < \alpha$ nous pouvons produire l'explication *a ne peut être affecté à C_t car a n'est indifférent sur la majorité des critères à aucun prototype.*

Nous pouvons distinguer les situations suivantes :

- La situation $(\sim(x, y) > \beta) : C^\sim(x, y) > \beta$ permet de produire l'explication: *a est fortement indifférent à y car ils sont indifférents sur la majeure partie des critères.*
- La situation $(\alpha < \sim(x, y) < \beta) : C^\sim(x, y) > \alpha$ permet de produire l'explication: *a est moyennement indifférent à y car le poids des critères indifférents est moyen.*
- La situation $(\sim(x, y) < \alpha) : C^\sim(x, y) < \alpha$ permet de produire l'explication *a est faiblement indifférent faiblement indifférent à y car les poids des critères indifférents est faible.*

Exemple 28 (Explication fondée sur la concordance globale) *Si la Renault Laguna est un prototype de la catégorie familiale, on peut générer l'explication : La Peugeot 206 S16 et la Renault Laguna ne peuvent être considérées comme indifférente car elles ne sont pas indifférentes sur la majeure partie des critères.*

La discordance globale Les explications produites au niveau de la discordance globale sont des explications négatives. Elles expriment la force des divergences qui interdisent l'indifférence entre un objet et un prototype. Nous pouvons produire des explications négatives telles que si $D^\sim(x, y) > \beta$ alors nous avons *a ne peut être indifférent à y car le poids des critères en violent désaccord avec l'indifférence est trop grand.* Nous distinguons à cet effet les situations suivantes :

- Si $\sim(x, y) > \beta$ alors $D^\sim(x, y) < 1 - \beta$ et nous pouvons produire l'explication: *a est fortement indifférent à y car le poids des critères en violent désaccord est faible.*
- Si $(\alpha < \sim(x, y) < \beta)$ alors nous distinguons les deux cas suivants:
 - Soit $\alpha < C^\sim(x, y) < \beta$ et $D^\sim(x, y) < 1 - \alpha$ nous pouvons produire l'explication suivante: *a est moyennement indifférent à y car le poids des critères en violent désaccord est moyen.*
 - Soit $C^\sim(x, y) > \beta$ et $1 - \beta < D^\sim(x, y) < 1 - \alpha$ nous pouvons produire l'explication suivante: *a est moyennement indifférent à y car le poids des critères en violent désaccord est moyen.*
- Si $(\sim(x, y) < \alpha)$ nous distinguons les deux cas:
 - Soit $C^\sim(x, y) < \alpha$ et la discordance n'intervient pas dans l'explication.
 - Soit $C^\sim(x, y) > \alpha$ et $D^\sim(x, y) > 1 - \alpha$ et nous pouvons produire l'explication suivante: *a est faiblement indifférent à y car le poids des critères en violent désaccord est trop fort.*

Exemple 29 (Explication fondée sur la discordance globale) *Soit la catégorie des petites voitures économique, on peut générer l'explication : La BMW 730d ne peut être affectée à la catégorie des petites voitures économiques car il existe des divergences extrêmes sur certains critères entre la BMW 730d et tous les prototypes de la catégorie*

Explication de troisième niveau

Les explications de troisième niveau servent à affiner les explications précédentes.

La concordance monocritère Les explications de ce type permettent au décideur de voir quels sont les principaux critères qui entrent dans confirmer l'indifférence. De ce fait les explications produites seront principalement des explications positives. Pour chaque critère j prenant part à la concordance nous pouvons produire des explications telles que si $C_j^\sim(x, y) > \beta$ nous avons *a est fortement indifférent à y sur le critère j*.

Exemple 30 (Explication fondée sur la concordance monocritère) *On peut expliquer l'indifférence entre la 206 S16 et la Clio 16 V par des explications pour chaque critère du type : La 206 S16 est indifférente à la Clio 16 V car elle ont des performances proches en vitesse de pointe.*

La discordance monocritère Les explications de ce type sont principalement des explications négatives. Elles vont permettre d'exprimer à cause de quel critère un objet ne va pas pouvoir être indifférent à un prototype. Nous pouvons produire des explications telles que si $D_j^\sim(x, y) > \beta$ nous avons *a est au mieux faiblement indifférent à y car ils sont extrêmement divergent sur le critère j*.

Exemple 31 (Explication fondée sur la discordance monocritère) *On peut expliquer la non affectation de la BMW 730d à la catégorie de voitures économiques par une explication du type : La BMW 730d ne peut être affectée à la catégorie des voitures économiques car son prix est beaucoup trop éloigné de celui des prototypes de la catégorie.*

Chapitre 4

Construction des poids des critères

Résumé

Les méthodes d'affectation que nous avons décrites supposent la connaissance de différents paramètres, comme les poids et les seuils des critères, pour évaluer l'indifférence entre les actions à affecter et les prototypes des catégories. Cependant dans bien des cas, la détermination de tels paramètres ne va pas de soi. Nous proposons donc de construire ces coefficients à partir d'un ensemble de points déjà affectés aux catégories, il s'agit d'un problème d'apprentissage supervisé. Nous proposons dans cette partie un modèle pour la construction des poids sous forme de programme mathématique. Cette construction des coefficients d'importance peut être effectuée de manière automatique cependant nous proposons de l'intégrer dans un processus interactif de construction des poids en tant que phase de calcul. Ce processus fait alterner des phases de calcul et des phases de dialogue pendant lesquelles les contraintes du programme mathématique peuvent être remises en cause par le décideur.

Mots clés : apprentissage supervisé, poids, coefficient d'importance, interactif, programmation mathématique.

4.1 Un modèle pour la construction des poids des critères

4.1.1 Pourquoi construire les poids des critères?

Les modèles d'évaluation multicritères, que ce soit ceux dans le cadre de la problématique du rangement, de choix ou du tri, supposent connus plusieurs paramètres préférentiels (seuils et poids). Cependant le décideur désirant utiliser ces modèles n'est pas nécessairement à même de spécifier ces paramètres. Nous nous intéressons dans ce chapitre aux paramètres que sont les coefficients d'importance des critères aussi appelés poids des critères. Ces paramètres représentent l'importance relative avec laquelle les critères interviennent dans l'évaluation d'une comparaison entre deux actions. Nous considérons dans ce chapitre que le décideur n'est pas capable de spécifier de telles valeurs. Nous supposons en revanche que le décideur est à même de spécifier une information beaucoup moins technique que ces paramètres et beaucoup plus en amont, plus générale : un ensemble de points déjà affectés.

Nous proposons un modèle pouvant être intégré dans une approche interactive permettant de construire les poids des critères à partir d'un ensemble d'apprentissage constitué de points représentant des actions dont l'affectation est connue, Z . Le problème se pose de la manière suivante : on recherche un ensemble de coefficients $\omega = \{\omega_1, \dots, \omega_n\}$ reflétant l'importance relative des critères connaissant :

- un modèle multicritère d'affectation constitué de catégories (méthode PIP) fourni par l'homme d'étude,
- un ensemble de catégories $C_1, \dots, C_q$ fournies par le décideur, représentées chacune par leur ensemble de points de référence, $Y_1, \dots, Y_q$,
- un ensemble de points d'apprentissage, $Z = \{z_1, \dots, z_m\}$ fourni par le décideur, formé de l'image d'actions dont l'affectation est connue (cette affectation pouvant éventuellement être graduelle) dans l'espace des critères. On note $\mu_t^o(z)$ l'appartenance du points d'apprentissage z à la c

L'information nécessaire au fonctionnement de notre approche concerne l'appartenance des points d'apprentissage aux catégories et non l'évaluation de l'indifférence entre les points d'apprentissage et les profils de catégories. Dans le cas de catégories multiprofiles, on ne demande pas au décideur à quel profil les points d'apprentissage sont indifférents. Ainsi, même si le décideur ne raisonne pas nécessairement avec des profils explicites (et n'est donc pas à même d'évaluer l'indifférence entre les points de Z et les profils), il peut exprimer son avis sur l'appartenance des points d'apprentissage aux catégories.

4.1.2 Objectif

L'objectif exposé est de minimiser la divergence d'appréciation entre l'affectation effectuée provenant de l'ensemble d'apprentissage et celle effectuée par la méthode. Nous considérons par ailleurs un jeu de poids par catégorie, l'apprentissage est effectué au niveau de chaque catégorie. L'objectif général s'exprime alors par :

$$\min_{\omega} F = \psi_{z \in Z_t} ((\mu_t(z) - \mu_t^o(z)))^2 \quad (4.1)$$

avec ψ opérateur monotone croissant tel que $\psi(0, \dots, 0) = 0$ et $\psi(1, \dots, 1) = 1$, $\mu_t^o(z)$ le degré avec lequel le point z mérite d'entrée dans C_t au vue de l'expert et $\mu_t(z)$ le degré d'affectation calculé

par la méthode en considérant le jeu de poids $\omega_1, \dots, \omega_n$.

Remarque 6 (Degré d'appartenance graduelle pour les points d'apprentissage)

Habituellement, les points d'apprentissage ont une appartenance nette aux catégories, cependant on peut être amené à disposer des points d'apprentissage dont l'appartenance n'est pas nette et ce pour les raisons suivantes :

- *L'expert peut vouloir exprimer des informations plus nuancées que la simple appartenance de type binaire. Sans pour autant vouloir faire varier le degré d'appartenance en continu dans l'intervalle $[0, 1]$, il peut vouloir par exemple exprimer une hésitation (catégorie incertaine) ou un manque de précision (catégories imprécises) et utiliser les valeurs de l'ensemble $\{0, \frac{1}{2}, 1\}$ pour le degré d'affectation où cette la valeur $\frac{1}{2}$ exprime cette hésitation ou imprécision.*
- *Si les données d'apprentissage proviennent d'un historique, il se peut très bien que l'appartenance aux catégories soit floue. Le caractère flou de l'appartenance pouvant provenir soit d'une autre méthode soit d'une affectation floue connue a posteriori.*
- *Enfin, il est possible de construire des degrés d'appartenance à partir d'histogrammes de fréquences provenant d'un historique, soit en construisant directement une estimation d'une probabilité, soit en construisant un degré d'appartenance sur la base d'une interprétation possibiliste des histogramme à la manière de Dubois et Prade [32].*

Considérer le min pour l'opérateur ψ correspond à une attitude totalement non compensatoire où une solution sera considérée comme mauvaise dès lors qu'elle contiendra un seul point d'apprentissage dont l'affectation effectuée par la méthode ne coïncide pas avec l'affectation connue μ_t^o . La somme en revanche permet, en tant qu'opérateur totalement compensatoire, une attitude de compromis équipondéré. On peut enfin vouloir adopter une attitude de compromis plus sophistiquée, en donnant aux points une importance d'autant plus grande que ceux-ci sont mal ré-affectés. Nous pouvons alors pour cela considérer un opérateur de type OWA, voir [116].

Dans le cas particulier où $\sum_{t=1}^q \mu_t(z) = 1, \forall z \in Z$ (la repartition des points dans les catégories forment une partition de l'ensemble des points), la fonction objectif peut être ré-écrite autrement (cf. equt. 4.2). L'objectif va être de bien affecter les points de la catégorie et de ne pas affecter ceux qui n'y appartiennent pas. La fonction objectif va donc être définie sur l'ensembles des points d'apprentissage de toutes les catégories de la manière suivante :

$$\min_{\omega} \quad F = \psi_{z \in Z} \left((\mu_t(z) - \mu_t^o(z))^2 \right) \quad (4.2)$$

avec ψ une fonction croissante de $(\mu_t(z) - \mu_t^o(z))^2$. Le problème de classification devient alors un problème de discrimination.

4.1.3 Contraintes

Nous allons distinguer deux types de contraintes : les Contraintes Structurelles (**CS**) et les Contraintes Évolutives (**CE**).

Contraintes Structurelles : Ce sont les contraintes inhérentes au modèle d'affectation considéré. Elles traduisent la définition d'un jeu de poids c'est-à-dire le fait que ceux-ci sont positifs et qu'ils somment à 1. On note **CS** cet ensemble de contraintes :

$$\mathbf{CS} = \left\{ \begin{array}{l} \sum_{j=1}^n \omega_j = 1 \\ \omega_j \geq 0 \quad \forall j \in \{1, \dots, n\} \end{array} \right. \quad (4.3)$$

Contraintes Évolutives : Ce sont les contraintes qui vont traduire les informations que le décideur est capable de donner, explicitement ou implicitement, sur l'importance des poids. Ces contraintes peuvent prendre la forme de valeurs possibles pour chacun des poids ou d'un ordre partiel sur les coalitions de critères et donc sur les sommes de poids :

$$\mathbf{CE} = \left\{ \begin{array}{l} \delta_i^{inf} \leq \omega_i \leq \delta_i^{sup} \quad \forall i \in \{1, \dots, n\} \\ \sum_{i \in I} \omega_i \geq \sum_{j \in J} \omega_j \quad \text{si } I \text{ est plus important que } J \end{array} \right. \quad (4.4)$$

avec $I \subset \{1, \dots, n\}$, $J \subset \{1, \dots, n\}$, $I \cap J = \emptyset$, et où δ_i^{sup} et δ_i^{inf} , sont les valeurs maximale et minimale que ω_i peut prendre. Par défaut $\delta_i^{sup} = 1$ et $\delta_i^{inf} = 0$. Les deux conditions suivantes doivent être respectées pour que les puissent sommer à 1 :

$$\begin{aligned} - \sum_{i=1}^n \delta_i^{sup} &\geq 1, \\ - \sum_{i=1}^n \delta_i^{inf} &\leq 1. \end{aligned}$$

On a donc le programme mathématique suivant :

$$\begin{aligned} \min_{\omega} \quad & F = \psi_{z \in Z_t} \left((\mu_t^o(z) - \mu_t(z))^2 \right) \\ \text{s.c.} \quad & \left\{ \begin{array}{l} \mathbf{CS} \\ \mathbf{CE} \end{array} \right. \end{aligned} \quad (4.5)$$

L'ensemble de contraintes **CE** permet, si besoin est, d'éviter qu'un critère en particulier se voit attribuer un poids trop important en majorant la valeur de chaque poids par δ_i^{sup} . On évite en particulier d'aboutir à un modèle d'affectation monocritère (si $\exists i \in \{1, \dots, n\}$, $\omega_i = 1$). Cela permet d'orienter la construction vers une distribution équitable du pouvoir décisionnel entre les critères.

Il est de même possible de prendre en compte une information purement ordinale que le décideur serait susceptible de donner sur les poids. Si celui-ci considère qu'un ensemble de critères I est plus important qu'un autre ensemble de critères J , il suffit alors de rajouter la contrainte $\sum_{i \in I} \omega_i \geq \sum_{j \in J} \omega_j$ (contraintes **CE**).

4.1.4 Extraction d'information préférentielle et perspective interactive

Extraction d'information sur les poids de critères par affectation segmentée

Nous proposons dans cette section un moyen de questionner le décideur afin d'obtenir une information ordinale sur les poids des critères (affectation segmentée¹). Cette information peut alors être utilisée pour construire les contraintes de type **CE**. Ce questionnaire est fondé sur la comparaison de l'appartenance d'actions fictives aux catégories. Ainsi, on ne questionne pas le décideur directement sur les poids et leur valeur, aspect qu'il n'appréhende pas nécessairement, mais à un niveau plus élevé qu'il est supposé appréhender : l'appartenance aux catégories.

Définissons tout d'abord la notion d'action multiprofilée :

Définition 16 (action multiprofilée) Soient x et y deux points de l'espace des critères et, I et J deux sous-ensembles formant une partition de F , l'ensemble des critères. On appelle action multiprofilée, l'action représentée par le point construit avec les performances de x sur les critères de I et les performance de y sur les critères de J . On note cette action $x_I y_J$.

1. en référence à la méthode des descriptions segmentées, Roy [94] et aussi [100] p. 301-309

Le principe de ce questionnaire consiste à construire deux actions fictives et à demander au décideur laquelle des deux mérite le plus d'entrer dans la catégorie. Voyons dans le cas d'une catégorie monoprotail l'information obtenue.

Soient une catégorie C_t représentée par un prototype y , I et J deux parties de l'ensemble des critères formant une partition de F et un point x de l'espace des critères défini tel que :

$$C_i^\sim(x, y) = 0 \text{ et } D_i^\sim(x, y) = 0, \forall i \in \{1, \dots, n\}$$

c'est à dire $p_i \geq |x_i - y_i| \geq v_i^-, \forall i \in \{1, \dots, n\}$.

Nous construisons ensuite deux actions multiprotailées a et b représentées par les vecteurs $x_I y_J$ et $x_J y_I$, puis on demande au décideur laquelle de ces deux actions mérite le plus d'entrer dans la catégorie C_t . Si a mérite plus que b d'intégrer la catégorie C_t , la réponse du décideur peut alors être modélisée par :

$$\mu_t(a) \geq \mu_t(b) \quad (4.6)$$

Comme les catégories sont monoprotails, le degré d'affectation à la catégorie est égal au degré d'indifférence avec le protail de C_t . Ainsi l'équation 4.6 est équivalente à :

$$\sim(x_I y_J, y) \geq \sim(x_J y_I, y) \quad (4.7)$$

Par définition de x il n'existe pas de critères discordants avec l'indifférence entre chacune des actions a et b , et le protail de C_t ainsi $D^\sim(x_I y_J, y) = 0$ et $D^\sim(x_J y_I, y) = 0$. Nous avons donc :

$$\sum_{i=1}^n \omega_i C_i^\sim(x_I y_J, y) \geq \sum_{i=1}^n \omega_i C_i^\sim(x_J y_I, y) \quad (4.8)$$

Comme I et J forment une partition de F , on a :

$$\underbrace{\sum_{i \in I} \omega_i C_i^\sim(x, y) + \sum_{i \in J} \omega_i C_i^\sim(y, y)}_0 \geq \sum_{i \in I} \omega_i C_i^\sim(y, y) + \underbrace{\sum_{i \in J} \omega_i C_i^\sim(x, y)}_0 \quad (4.9)$$

car $C_i^\sim(x, y) = 0, \forall i \in \{1, \dots, n\}$, on a donc :

$$\sum_{i \in J} \omega_i C_i^\sim(x_I y_J, y) \geq \sum_{i \in I} \omega_i C_i^\sim(x_J y_I, y) \quad (4.10)$$

$$\sum_{i \in J} \omega_i \geq \sum_{i \in I} \omega_i \quad (4.11)$$

Nous obtenons donc une information ordinale sur l'importance relative des coalitions I et J . Cette information peut être intégrée dans les contraintes évolutives **CE**.

Perspective interactive

On peut envisager de résoudre le programme (4.5) afin d'obtenir un jeu de poids optimal permettant de résumer au mieux l'ensemble d'apprentissage Z_t . Il s'agit là d'une approche orientée apprentissage automatique. Le décideur fournit l'information dont il dispose, et celle-ci pouvant être

extrêmement pauvre. Cependant, une autre approche est aussi possible; elle consiste à avoir une attitude interactive (*cf.* Vanderpooten [110], Roy et Bouyssou chapitre 7 de [100]), avec le décideur et à réaliser une véritable aide à la construction des poids plutôt qu'une construction complètement automatique de ceux-ci.

La procédure débute par une phase de *data mining* dans laquelle le décideur spécifie tout d'abord l'information qu'il est capable de fournir. Ainsi en utilisant les affectations segmentées et les informations éventuelles du décideur l'ensemble de contraintes $\mathbf{CE}_0$ peut être construit. Ensuite le programme mathématique est résolu en considérant $\mathbf{CS}$ et $\mathbf{CE}_0$. L'ensemble de poids ω_0 obtenu sert alors de base pour faire réagir le décideur. Celui-ci a la possibilité de réagir, soit en acceptant le jeu de poids, soit en le rejetant tout en soulignant éventuellement les éléments constitutifs de son insatisfaction. Ces éléments doivent être pris en compte pour la définition du nouvel ensemble de contraintes $\mathbf{CE}_i$ à travers un protocole de dialogue. La procédure est itérée jusqu'à obtention d'un jeu de poids satisfaisant (voir figure (4.1)).

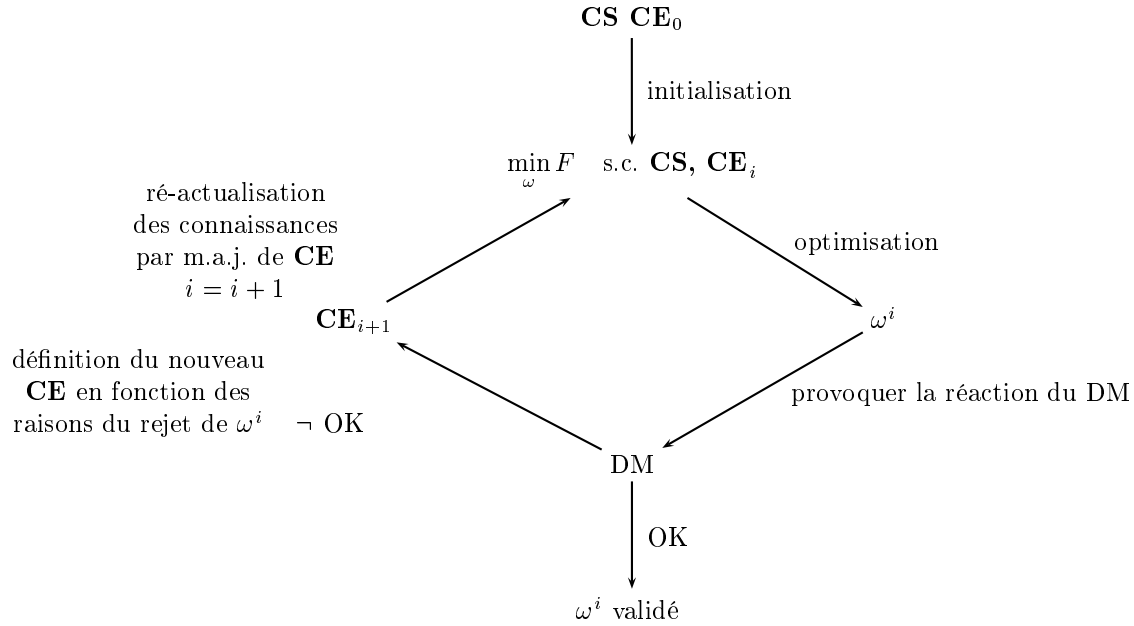


FIG. 4.1 – Construction interactive des poids

4.2 Exemples de programmes sous forme linéaire

4.2.1 Hypothèses sur les opérateurs d'agrégation

Pour pouvoir exprimer aisément le programme (4.5) sous forme linéaire certaines hypothèses doivent être mises sur les opérateurs ψ du programme mathématique, S et T du modèle d'affectation (méthode PIP et k -PIP).

Pour la t -norme T agrégeant les degrés de concordance et de discordance, nous utiliserons soit le produit soit le min selon le pouvoir que l'on veut donner à la discordance. La t -norme idempotente fait agir la discordance uniquement comme un minorant de l'indifférence alors que le produit permet de donner un pouvoir plus important à la discordance en lui permettant d'agir quelque soit le niveau de concordance dès lors qu'un critère se trouve discordant. Ces deux opérateurs seront aisément linéarisables si l'on pose comme condition que la discordance est calculée sans faire appel aux

coefficients d'importance des critères. Dans la pratique, nous utiliserons une des deux formules suivantes pour agréger les coefficients d'appartenance à la coalition discordante :

$$\begin{aligned} - D^\sim(a, b) &= \max_{j=1}^n (D_j^\sim(a, b)) \\ - D^\sim(a, b) &= \prod_{j=1}^n (1 - D_j^\sim(a, b)) \end{aligned}$$

Les deux instances de l'opérateur (ψ) d'agrégation des coefficients d'appartenance des points d'apprentissage que nous avons envisagés précédemment, à savoir le min et la somme, sont aussi aisément linéarisables.

Dans cette partie, nous faisons l'hypothèse que les catégories sont représentées par un seul prototype, que $\mu_t^o = 1$ et que l'appartenance des points d'apprentissage est nette ($\forall z \in Z, \mu_t^o(z) \in \{0, 1\}$). L'opérateur d'agrégation des degrés d'indifférence, S , n'est donc pas nécessaire. L'objectif prend alors la forme suivante :

$$\max_{\omega} F = \psi_{z \in Z_t} (T(C^\sim(z, y), N(D^\sim(z, y)))) \quad (4.12)$$

avec y l'unique profil de C_t , T une t-norme et N un opérateur de négation.

4.2.2 Indifférence définie sur un principe de concordance

L'objectif est de maximiser la bonne affectation des points d'apprentissage de Z_t dans la catégorie C_t . Chaque catégorie étant représentée par un prototype, nous considérerons l'indifférence entre les points d'apprentissage et y , le profil de C_t . On a donc l'objectif suivant, qui forme avec les contraintes **CE** et **CS** un programme linéaire aisé à résoudre :

$$\begin{aligned} \max_{\omega} F &= \sum_{z_i \in Z} C^\sim(z_i, y) = \sum_{z_i \in Z} \left(\sum_{j=1}^n \omega_j \times C_j^\sim(z_i, y) \right) \\ \text{s.c.} \quad &\begin{cases} \mathbf{CE} \\ \mathbf{CS} \end{cases} \end{aligned}$$

4.2.3 Indifférence définie sur le principe de concordance et non-discordance

Nous supposons ici l'indifférence définie par :

$$\begin{aligned} \sim(a, y) &= \min(C^\sim(a, y), 1 - D^\sim(a, y)) \\ &= \min \left(\sum_{j=1}^n \omega_j \times C_j^\sim(a, y), K \right) \end{aligned}$$

avec $K = 1 - D^\sim(a, y)$

Il faut alors maximiser la fonction suivante :

$$F = \sum_{z_i \in Z} \left(\min \left(\sum_{j=1}^n \omega_j \times C_j^\sim(z_i, y), K \right) \right)$$

profils	logiciel	développement	mobilité	vente
ingénieurs	15	15	10	10
commercial	10	10	15	15
polyvalent	15	15	15	15

TAB. 4.1 – *Profils des catégories*

Cette expression peut être linéarisée. On a alors le programme linéaire suivant constitué de $2|Z| + n$ contraintes en plus des contraintes de **CE** et de $n + |Z|$ variables :

$$\begin{aligned} \max_{\omega, m_i} \quad & \sum_{i=1}^{|Z|} m_i \\ \text{s.c.} \quad & \left\{ \begin{array}{l} \forall z_i \in Z \left\{ \begin{array}{l} m_i \leq \sum_{j=1}^n \omega_j \times C_j^\sim(z_i, y) \\ m_i \leq K \end{array} \right. \\ \mathbf{CE} \\ \mathbf{CS} \end{array} \right. \end{aligned} \quad (4.13)$$

4.3 Exemple

4.3.1 Définition du problème

Nous considérons un problème d'embauche de candidats issu de [80] à des postes d'ingénieurs, de commerciaux et ingénieurs-commerciaux (candidats polyvalents). On a donc un problème de classification en trois catégories C_1 , C_2 et C_3 correspondant respectivement aux postes d'ingénieurs, de commerciaux et ingénieurs-commerciaux.

Chaque candidat est évalué sur une famille de quatre critères par un vecteur de notes allant de 0 à 20. Les critères sont les connaissances logiciel, l'expérience en développement, la mobilité et l'expérience de vente. On cherche à retrouver les poids correspondants aux critères à partir d'un ensemble de candidats dont le degré d'affectation est déjà connu. Nous considérons 40 candidats dans chaque catégorie dont les performances et les degrés d'affectation sont donnés respectivement dans les tableaux B.1, B.2 et B.3. Les catégories ne forment pas une partition cependant nous considérons qu'elles sont mutuellement exclusives en effet, du fait de la discordance, il n'existe que quelques rares candidats qui peuvent être affectés à plusieurs catégories (avec des degrés extrêmement faibles).

4.3.2 Formalisation

Chaque catégorie est modélisée à l'aide d'un profil central correspondant à l'idée que le décideur se fait de l'individu typique de la catégorie, voir tableau 4.1.

Chaque critère est construit selon un modèle de pseudo-critère à seuils constants et identiques pour chacun d'entre eux. Les seuils d'indifférence, de préférence et de veto sont respectivement égaux à : 1, 2 et 5. Nous considérons une relation d'indifférence utilisant un modèle d'agrégation non totalement compensatoire, définie par :

$$\sim(x, y) = \sum_{j=1}^n C_j^\sim(x, y) \times K$$

catégorie	g_1	g_2	g_3	g_4
ingénieurs	0, 16	0, 23	0, 36	0, 25
commerciaux	0, 28	0, 23	0, 22	0, 27
ing.-com.	0, 20	0, 30	0, 23	0, 27

TAB. 4.2 – *Poids construits avec les coefficients dégradés de l'ensemble d'apprentissage*

avec $K = 1 - D^\sim(x, y)$ une constante

Nous pouvons considérer deux cas, tout d'abord le cas où les catégories ne forment pas nécessairement une partition puis le cas où elles en forment une. Dans le cas d'une non partition le décideur considère les catégories indépendamment les unes des autres, il ne cherche pas à discriminer les candidats mais plutôt à savoir ce à quoi ils sont aptes (autrement dit un candidat peut être apte à plusieurs postes). En revanche, dans le cas d'une partition on recherche réellement à discriminer les candidats. Cela se traduit par une différence dans la fonction objectif du programme mathématique, dans le cas d'une partition nous allons considérer les points de toutes les catégories alors que dans le cas contraire seuls les points de la catégorie seront considérés.

Si l'affectation des actions aux catégories ne forme pas une partition de A , le programme mathématique à résoudre est donc pour chaque catégorie $t \in \{1, 2, 3\}$:

$$\begin{aligned} \min_{\omega} \quad & F = \sum_{z \in Z_t} (\mu_t^o(z) - \mu_t(z))^2 \\ \text{s.c.} \quad & \begin{cases} \mathbf{CS} \\ \mathbf{CE} \end{cases} \end{aligned} \quad (4.14)$$

avec $\mu_t(z) = \sim(z, y_t)$ où y_t est le prototype de C_t .

Dans le cas d'une partition on a la fonction objectif :

$$F = \sum_{z \in Z} (\mu_t^o(z) - \mu_t(z))^2 \quad (4.15)$$

L'objectif dans le programme mathématique 4.14 diffère de celui exprimé dans l'équation 4.15 par l'utilisation des points d'apprentissage de toutes les catégories (Z) plutôt qu'uniquement des points d'apprentissage de Z_t .

4.3.3 Résultats

Les affectations des points d'apprentissage ont été effectuées à l'aide du même vecteur poids pour toutes les catégories, $\omega = (0.25, 0.25, 0.25, 0.25)$. Voyons s'il est possible de retrouver ce vecteur à partir d'une information dégradée (modification des degrés d'affectation) puis à partir d'une encore plus frustrée, l'appartenance stricte.

Nous avons dans un premier temps effectué un apprentissage en utilisant les degrés d'appartenance initiaux. On retrouve alors exactement les jeux de poids initial. Cependant il s'agit là d'un cas tout à fait idéal, voyons maintenant ce qui se passe si nous introduisons une perturbation de 10 % en moyenne dans les coefficients initiaux, voir tableau B.4. Nous verrons enfin les résultats obtenus dans le cas où l'on utilise des degrés d'affectation nets.

Chaque ensemble de candidats déjà affectés est partagé en ensemble d'apprentissage et ensemble de test de 20 points chacun. Nous présentons les poids construits à partir de l'ensemble d'apprentissage dans le tableau 4.2.

ensembles	ingénieurs		commerciaux		ing.-com.		tous	
	max.	moy.	max.	moy.	max.	moy.	max.	moy.
apprent.+ test	0,1118	0,0415	0,0493	0,0147	0,0673	0,0229	0,1118	0,0263
apprent.	0,1118	0,0335	0,0493	0,0167	0,0673	0,0227	0,1118	0,0243
test	0,1118	0,0495	0,0493	0,0126	0,0673	0,0231	0,1118	0,0284

TAB. 4.3 – Différence entre affectations initiales et calculées avec les poids construits

catégorie	g_1	g_2	g_3	g_4
ingénieurs	0,11	0,15	0,29	0,46
commerciaux	0,13	0,27	0,14	0,46
ing.-com.	0,25	0	0,53	0,22

TAB. 4.4 – Poids construits avec les coefficients nets de l'ensemble d'apprentissage

On constate que les poids construits sont assez proches de ceux utilisés pour obtenir l'affectation initiale puisque la différence maximale avec les poids initiaux est de 0,11. Cependant ce n'est pas le seul résultat à considérer, il convient aussi de s'intéresser à la qualité de l'affectation de l'ensemble de test effectuée avec les jeux de poids construits sur l'ensemble d'apprentissage. Nous considérons pour cela la différence d'affectation avec les coefficients initiaux. Deux indicateurs, la moyenne des différences d'affectation et la plus grande différence d'affectation, sont considérés. Nous présentons dans le tableau 4.3 ces indicateurs pour les ensembles d'apprentissage et de test et pour la réunion des deux.

On constate une différence d'affectation moyenne qui varie entre 1,2% et 5%. Ces valeurs sont largement acceptables dans un cadre d'aide à la décision. En effet présenter une affectation à un décideur avec un degré de 0,80 plutôt qu'un degré de 0,85 ne peut être considéré comme une erreur. D'autre part, la différence d'affectation maximale qui ne dépasse pas 12% constitue elle aussi un excellent résultat.

Considérons maintenant que l'affectation connue des points d'apprentissage est une affectation nette. On obtient les vecteurs de poids du tableau 4.4.

Ce jeu de poids semble moins satisfaisant que celui obtenu en utilisant les coefficients dégradés. On retrouve notamment pour la catégorie des ingénieurs commerciaux le critère g_2 avec un poids nul et le critère g_3 avec un poids très élevé de 0,53. Cependant, comme nous l'avons exprimé précédemment, il ne suffit pas de regarder le jeu de poids obtenu, il convient surtout de s'intéresser aux résultats de l'affectation effectuée sur l'échantillon de test. Les résultats de l'affectation qui sont résumés dans le tableau 4.5 donnent des résultats qui peuvent être considérés comme satisfaisants. En effet, la différence d'affectation, en moyenne n'excède pas 8,2%, ce qui apparaît encore comme une différence très acceptable. De plus la différence maximale, si elle paraît importante 0,28, doit être modérée par le fait qu'elle ne correspond qu'à un seul point mal classé.

ensembles	ingénieurs		commerciaux		ing.-com.		tous	
	max.	moy.	max.	moy.	max.	moy.	max.	moy.
apprent.+ test	0,2442	0,0775	0,2286	0,0817	0,281	0,0665	0,281	0,0752
apprent.	0,2442	0,0796	0,2286	0,1014	0,281	0,0638	0,281	0,0816
test	0,2442	0,0753	0,2286	0,0620	0,281	0,0691	0,281	0,0688

TAB. 4.5 – Différence entre affectations initiales et calculées avec les poids construits à partir des coefficients nets

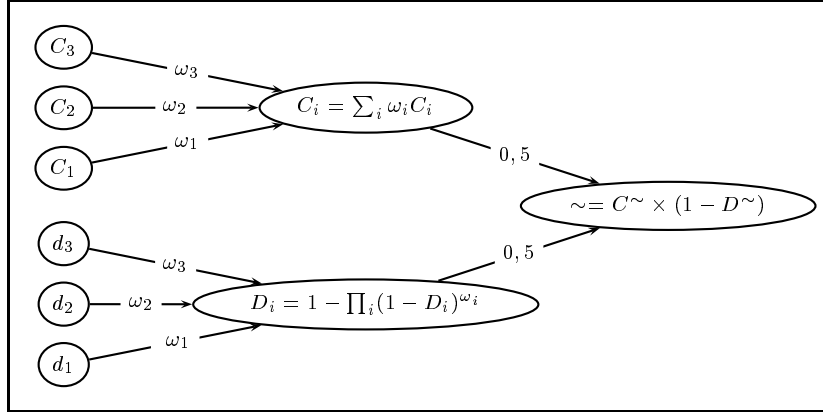


FIG. 4.2 – Structure d'un RN pour l'apprentissage de coefficients d'importance

4.4 Perspective utilisant une approche connexioniste

Les réseaux de neurones, en particulier les Perceptrons MultiCouches (PMC), sont des méthodes permettant de reproduire une fonction d'affectation par l'intermédiaire d'un apprentissage de poids. Cependant ces poids ne sont pas interprétables directement par un décideur. Les réseaux de neurones sont souvent qualifiés de méthodes de type boîte noire, c'est-à-dire de méthodes ne possédant aucun caractère explicatif. Nous proposons d'utiliser une approche neuronale se fondant sur un PMC à une couche cachée. En donnant aux neurones une signification particulière, celle d'opérateurs d'agrégation multicritère, les poids du réseau peuvent être interprétés comme les coefficients d'importance des critères.

Celui ci est composé d'une couche d'entrée ayant $2n$ neurones, n neurones prenant comme valeurs les indices de discordance, notés $D_1, \dots, D_n$, et n neurones prenant les indices de concordance, notés $C_1, \dots, C_n$. La couche cachée est composée de deux neurones, un pour agréger les indices de concordance, noté C , et un pour agréger les indices de discordance, noté D . Les fonctions de transferts à utiliser correspondent aux fonctions d'agrégation multicritères. La couche de sortie est elle constituée d'un seul neurone, noté S , dont la fonction de transfert est la fonction d'agrégation de la discordance globale et de la concordance globale.

Tous les neurones de la couche d'entrée ne vont pas être reliés à tous les neurones de la couche cachée. Les neurones correspondants aux indices de concordance monocritères sont uniquement reliés au neurone caché de concordance et les neurones d'entrée de discordance monocritère sont reliés au neurone caché de discordance. Il y a donc $2n$ poids à considérer. Une question se pose: "est il significatif de considérer que les opérateurs d'agrégation de concordance et de discordance utilisent des jeux de poids différents?". On peut dans un premier temps considérer que non et imposer que les poids soient les mêmes pour un même pour un critère, $\omega_1, \dots, \omega_n$. L'algorithme de rétropropagation de gradient est alors utilisé pour l'apprentissage des poids. On voit sur la figure 4.2 p. 90 la structure d'un RN pour l'apprentissage de poids d'un problème à trois critères.

Chapitre 5

Construction des prototypes et problématique de la typologie

Résumé

Nous nous intéressons, dans ce chapitre, à la construction de la représentation des catégories. En effet un décideur n'est pas toujours à même de spécifier les points de référence des catégories. Nous proposons à cet effet une approche permettant la construction de points de référence (prototypes) à partir d'un ensemble de points exemples. Nous nous plaçons dans le cadre de la problématique de la typologie. Pour cela nous proposons deux approches : une fondée sur une méthode de clustering la méthode des nuées dynamiques et l'autre fondée sur une approche meta-heuristique : un algorithme génétique. La première présente l'avantage d'une convergence très rapide compatible avec une utilisation en temps réel en vue d'une approche interactive mais nécessite de fixer le nombre de points de référence par catégorie. La seconde présente une convergence plus lente mais n'impose pas de fixer le nombre de profils. Enfin nous précisons que nous ne nous présentons pas comme spécialistes de ces techniques mais que nous les considérons comme des outils pour construire les catégories.

Mots clés : algorithme génétique, apprentissage supervisé et non supervisé, construction de prototypes, meta-heuristique, nuées dynamique, typologie.

5.1 Pourquoi construire des points de référence?

Les méthodes de classification multicritères que nous avons présentées au chapitre 3 supposent la connaissance de la représentation des catégories, c'est-à-dire des points de référence de celles-ci. Dans la pratique ces points de référence sont généralement spécifiés soit par le décideur soit par un acteur faisant fonction d'expert. Cependant, le décideur ou l'expert peut ne pas être à même de spécifier de réels prototypes de chaque catégorie. Il peut ne disposer que d'un ensemble d'actions dont l'affectation aux catégories est connue mais ne bénéficiant d'aucun caractère prototypique. Ce type d'action pouvant provenir d'un historique, d'une classification connue *a posteriori* ou encore d'une classification effectuée par une autre méthode.

Nous avons proposé deux méthodes, les méthodes PIP et k -PIP, qui vont permettre de gérer des points de référence de nature différente (*cf.* p. 73). Ainsi la méthode PIP va être adaptée dans le cas où les points de référence sont des prototypes des catégories et la méthode k -PIP dans le cas où les points de référence ne sont que des points exemples. Nous disposons donc de méthodes permettant de gérer les deux situations. Cependant le décideur peut ne disposer que de points exemples et vouloir effectuer une affectation fondée sur l'utilisation de prototypes et ce pour les raisons suivantes :

- Pour des raisons d'utilisation plus rapide (temps réel par exemple) : une procédure d'affectation considérant des catégories décrites par un nombre faible de prototypes est beaucoup plus efficace en terme de temps d'exécution qu'une procédure considérant un nombre de points exemples important. En effet le nombre de calcul de degrés d'indifférence est d'autant plus faible que l'ensemble des points de référence est réduit. On peut citer par exemple le site *Information and Computer Science, University of California, Irvine* [53] qui présente une liste de bases de données téléchargeables utilisables pour la classification, certaines d'entre elles sont très volumineuses comme la *Forest Covertype Database* constituée de 581012 points exemples répartis en 8 catégories.
- Pour des raisons explicatives : l'utilisation de prototypes permet de générer une explication de l'affectation contrairement à l'utilisation de points exemples. En effet, il est possible de générer des explications du type "l'action a est affectée à la catégorie C_t car elle est fortement indifférente à y ". Si le point de référence y est un prototypes de C_t , ce genre d'explication a effectivement avoir un sens pour le décideur car y reflète les aspects typiques de la catégorie. En revanche si la catégorie sont représentées par des points exemples (méthode k -PIP) le degré d'affectation fait intervenir plusieurs points de référence et donc prend en compte l'indifférence avec k points de référence, il est donc plus difficile de fournir une explication concise compréhensible pour un décideur. De plus les explications que l'on pourrait générer ne font intervenir que des points exemples qui ne reflètent pas nécessairement les aspects typiques de la catégorie.

Nous supposons donc que le décideur n'est pas à même de caractériser les catégories par des prototypes mais dispose seulement d'ensemble d'exemples. Nous distinguons à cet effet, l'ensemble des points exemples $Z = \cup_{t=1}^q Z_t$ qui va servir de base à la construction des prototypes et l'ensemble des points de référence qui sont ici des prototypes des catégories, $Y = \cup_{t=1}^q Y_t$. L'objectif va être de construire des prototypes des catégories à partir de l'ensemble de points exemples, nous nous plaçons donc dans un cadre proche de celui de la problématique de la typologie.

Définition 17 (Problématique de la Typologie) *Cette problématique considère un ensemble d'actions non triées. Elle vise à regrouper, soit de manière automatique soit en concertation avec un décideur ou un expert, ces actions en catégories, où les actions entrant dans un même catégorie méritent de recevoir la même décision. Les catégories ainsi construites ne correspondent pas à des classes où les actions sont proches mais à des groupes d'actions méritant les mêmes décisions.*

Cette problématique vise à regrouper les actions nécessitant des traitements similaires. Notre objectif de construction de prototypes est proche de cette problématique dans le sens où l'on cherche à regrouper les points exemple afin de construire pour chaque groupe un prototype supposé le représenter. Dans cette optique la problématique de la typologie peut nous amener à résoudre les deux problèmes suivants :

- Cette problématique vise à faire ressortir les éléments principaux qui vont intervenir dans une décision en mettant en avant les points communs entre les points exemples des ensembles $Z_1, \dots, Z_q$ à travers les prototypes construits. L'objectif recherché est alors de type cognitif et apparaît proche de la problématique de la description. La procédure résultant peut alors être qualifiée de *procédure de description de décision*.
- Cette problématique peut aussi apparaître comme un complément à la problématique du tri. Elle vise alors à construire les catégories auxquelles doivent être par la suite affectées les actions à classer. La procédure résultant est alors qualifiée de *procédure de construction de catégories*.

Nous nous plaçons dans le cadre d'une procédure de construction de catégories. Cette optique de construction de description de catégories a fait l'objet d'une littérature importante surtout par la communauté *machine learning* que ce soit pour les méthodes de type k -voisins ou pour les méthodes de construction d'arbres de décision (voir [50] pour une bibliographie). On peut citer par exemple, De Jong et Spears [27] avec un système en charge d'apprendre des règles de classification. Skalak [105] a quant à lui développé un algorithme génétique pour l'apprentissage de prototypes dans le cas du plus proche prototypes (méthodes de Hart [48] *condensed nearest neighbor rule* et Gates [40] *reduced nearest neighbor rule*).

Par ailleurs, il nous paraît important d'ajouter que l'effet de renforcement induit par la méthode k -PIP (p. 73) n'est pas nécessaire lorsque les points de référence sont de réels prototypes des catégories. Bezdek et Castela [9], à propos des méthodes de type k -voisins, montrent à cet effet que l'utilisation d'ensembles d'exemples conjointement à la méthode des k plus proches voisins (avec $k \geq 2$) est équivalente en terme de mauvaise classification à l'utilisation d'ensembles réduits de prototypes conjointement à la méthode du plus proche prototype.

Nous envisageons donc un ensemble de points d'apprentissage, $Z_1, \dots, Z_q$, qu'il s'agit de résumer au mieux, catégorie par catégorie, pour obtenir un ensemble représentatif de prototypes $Y_1, \dots, Y_q$. Pour cela nous proposons d'utiliser une méthode de *clustering* qui exploite le concept d'indifférence floue dans une méthode de type nuées dynamiques [28], [29]. Nous envisageons ensuite une approche de type meta-heuristique fondée sur un algorithme génétique [43]. Les approches que nous proposons se distinguent principalement de celles citées au paragraphe précédent par leur capacité à considérer des catégories multiprofiles. Cet aspect est évoqué par Sen et Knight [103], certaines catégories (comme les catégories non linéairement séparable) ne peuvent être caractérisées par une représentation monoprofil.

5.2 Construction de prototypes par nuées dynamiques

5.2.1 Principe

La méthode des nuées dynamiques [28] est une méthode de *clustering*, elle permet d'obtenir une partition d'un ensemble de points. L'ensemble de points est partagé en sous-ensembles homogènes appelés *clusters*, chaque *cluster* est représenté par un *noyau* supposé le résumer (le noyau est une représentation du *cluster*), le nombre de clusters est supposé connu.

Soient Z l'ensemble d'apprentissage, q le nombre de classes recherchées, $\mathcal{P} = \{\mathcal{P}_1, \dots, \mathcal{P}_q\}$ une partition de Z , $Y = \{y_1, \dots, y_q\}$ l'ensemble des noyaux représentant $\mathcal{P}$.

tirer q noyaux initiaux $\{y_1, \dots, y_q\}$ (tirage aléatoire ou non)

RÉPÉTER

- construire $\mathcal{P}$ en associant chaque $z \in Z$ à la partie $\mathcal{P}_i$ représentée par le y_i le plus proche

- calculer les nouveaux noyaux $\{y_1, \dots, y_q\}$

JUSQU'À stabilité (Y n'évolue plus)

FIG. 5.1 – *Algorithme des nuées dynamiques*

Les noyaux représentant les *clusters* peuvent être de différente nature : centroïde, droite ou sous ensemble de points. Dans notre cas, nous cherchons à caractériser chaque catégorie par une famille de prototypes. Cependant celles-ci ne forment pas nécessairement une partition et ne peuvent par conséquent par être assimilées aux *clusters*. Nous avons donc pour objectif d'appliquer la méthode catégorie par catégorie, en recherchant au sein de chacune d'entre elles les noyaux qui représentent le mieux leurs *clusters*. Chaque noyau correspond donc à un point de E .

La méthode consiste à construire des partitions successives de l'ensemble des points. Elle peut être divisée en phases : une phase d'initialisation, une phase d'affectation et une phase de construction.

- La phase d'initialisation consiste choisir les noyaux initiaux.
- La phase d'affectation consiste à affecter les points aux noyaux qui sont les plus proches afin de former les *clusters*.
- La phase de construction consiste à calculer les nouveaux noyaux correspondant *clusters* construits.

Ainsi nous pouvons donner l'algorithme 5.1, le déroulement de cet algorithme peut être illustré par la figure 5.2.

Nous donnons l'exposé de la méthode des nuées dynamiques en annexe C. Nous montrons entre autres les hypothèses sous-jacente à la méthode et la convergence.

5.2.2 Mise en œuvre

Afin d'exposer la mise en œuvre de notre méthodes nous précisons les notations suivantes :

- Z ensemble des points d'apprentissage,
- $\mathcal{P} = \{\mathcal{P}_1, \dots, \mathcal{P}_r\}$ l'ensemble des parties (*clusters*) de Z ,
- $Y = \{y_1, \dots, y_r\}$, l'ensemble des prototypes (*noyaux*) que nous cherchons à construire.
- $dissim(x, y) = 1 - \sim(x, y) + \alpha \|x - y\|$ la mesure de dissimilarité utilisée pour mesurer l'éloignement entre deux points.

Tirage des noyaux initiaux

Le résultat produit par la méthode des nuées dynamiques est fonction des noyaux initiaux. Ceux-ci peuvent être fixés de plusieurs manières. La plus simple consiste à effectuer un tirage

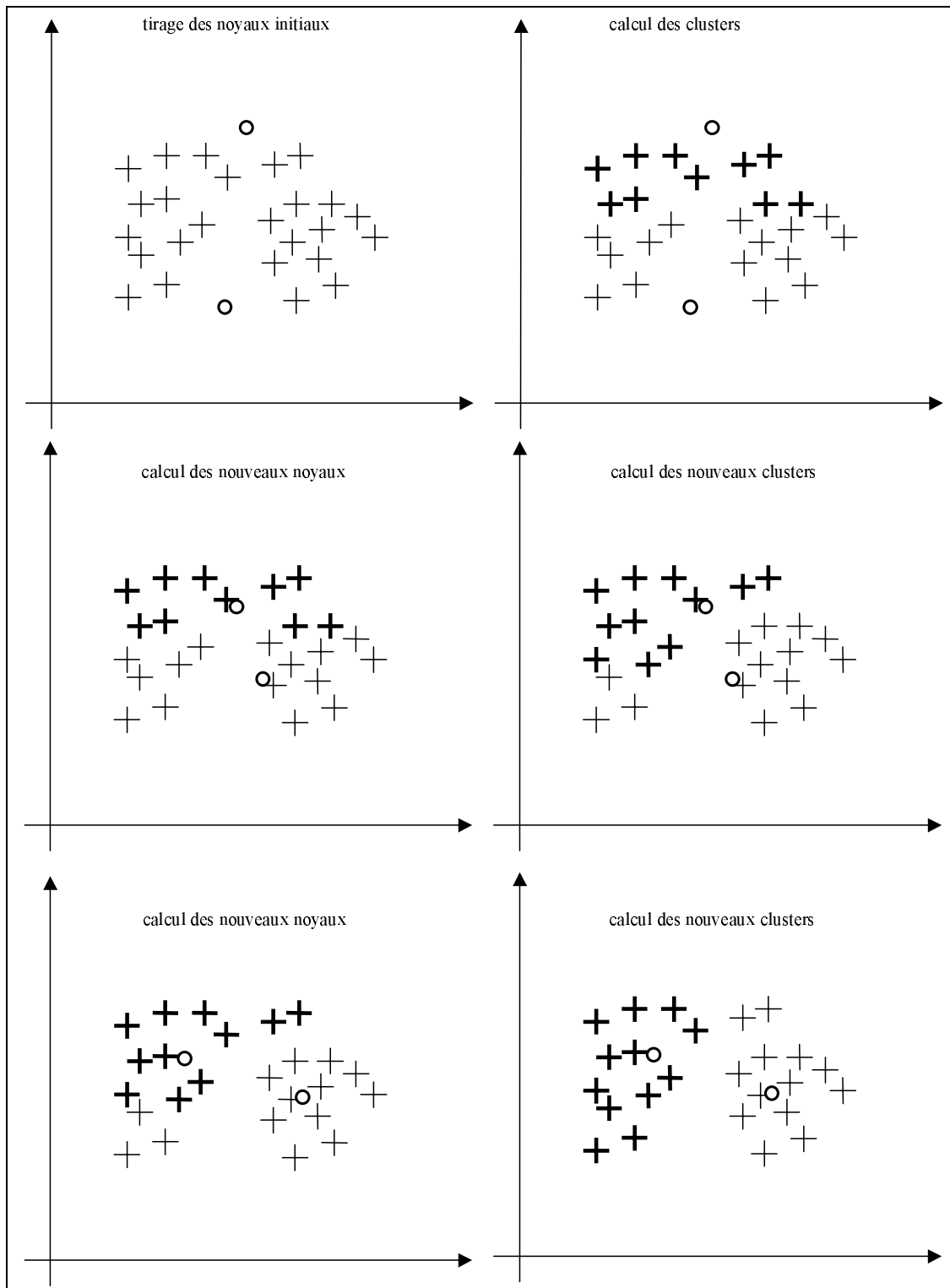


FIG. 5.2 – *Calcul de clusters par nuées dynamiques*

aléatoire parmi les points d'apprentissage. Cependant cette technique simple peut être améliorée aisément. Il faut alors effectuer une sélection de noyaux initiaux suffisamment différents au sens de la relation d'indifférence $\sim$. Cela permet alors d'améliorer à la fois la rapidité de convergence (nombre d'itérations nécessaires) et la qualité de la représentation (critère de convergence W , cf. annexe C).

Construction de la partition P

La partition $\mathcal{P}$ est définie à l'aide de la fonction d'affectation par la règle suivante :

$$\mathcal{P}_i = \{z \in Z \mid dissim(y_i, z) \leq dissim(y_j, z) \ \forall j \in \{1, \dots, r\} / \{i\}\}$$

Cette règle correspond à la méthode PIP exposée précédemment. En effet chaque point de Z est affecté à la partie correspondant au noyau dont il est le plus proche, il s'agit là de la règle du plus indifférent prototype.

Calcul des noyaux

Les nouveaux noyaux sont calculés à chaque itération, ainsi les performances des prototypes sont définies telles que :

$$g_i(y_l) = \phi_{z \in P_i}(g_i(z)) \quad \forall i \in [1, n]$$

La fonction ϕ est un opérateur de compromis, la plupart du temps les noyaux correspondent aux barycentres des *clusters*, la fonction ϕ est alors une moyenne arithmétique. Cependant pour les critères qualitatifs des opérateurs plus ordinaux comme le mode ou la médiane doivent être utilisés.

Choix du nombre de prototype par catégorie

Si l'on note Z_t l'ensemble des points exemples de la catégorie C_t , l'algorithme doit être appliqué à chacun des ensemble Z_t . Ne sachant pas a priori combien de prototypes sont nécessaires à la description d'une catégorie, le nombre r de *clusters* à distinguer au sein de chacune d'entre elle n'est pas connu. L'algorithme doit donc théoriquement être appliqué pour rechercher $1, 2, \dots, |Z_t|$ *clusters* dans Z_t . Néanmoins, la pratique nous montre que souvent quelques prototypes bien choisis sont suffisants pour résumer l'information contenue dans Z_t .

Pour évaluer l'aptitude d'un ensemble de prototypes, Y_t , à résumer une ensemble de points d'apprentissage Z_t nous utilisons deux critères qui vont évaluer d'un côté la qualité d'une représentation et d'un autre la non redondance des prototypes.

Qualité de la Représentation Ce critère va évaluer l'aptitude d'un ensemble de prototypes à résumer les points d'apprentissage de la catégorie, c'est-à-dire la qualité de la représentation de Z_t par Y_t . Autrement dit, les points de Z_t doivent être correctement ré-affectés à l'aide de prototypes de Y_t . Cette quantité est évaluée par :

$$QR(Z_t, Y_t) = 1 - \Phi(\delta_1, \dots, \delta_{|Z_t|})$$

où $\delta_i = |\mu_t(z_i) - \mu_t^o(z_i)|, \forall i \in \{1, \dots, |Z_t|\}$, $\mu_t(z)$ est le degré d'affectation de z_i à C_t en considérant les prototypes de Y_t , Φ une fonction croissante de $[0, 1]^{|Z_t|}$ dans $[0, 1]$ et $\mu_t^o(z)$ le degré avec lequel z entre dans C_t .

La quantité $|\mu_t(z) - \mu_t^o(z)|$ évalue la divergence d'affectation entre la méthode PIP utilisant les prototypes calculés et l'affectation connue des points de Z_t .

exemple : $QR(Z_t, Y_t) = 1 - \frac{1}{|Z_t|} \sum_{z \in Z_t} (|\mu_t(z) - \mu_t^o(z)|)$

Non Redondance des Prototypes Dans un ensemble de prototype il n'est pas nécessaire de retrouver des prototypes indifférents (tels que $\sim(y, y') = 1$), il y a alors une redondance d'information. Cette quantité peut être évaluée par :

$$NR(Y_t) = 1 - \max_{(y, y') \in Y_y \times Y_t} \sim(y, y')$$

En pratique, à mesure que le nombre r de noyaux augmente, on observe une amélioration de la qualité du résumé mais la représentation devient potentiellement plus coûteuse : plus de prototypes implique plus de comparaisons dans l'affectation (*cf.* figure 5.3 dans l'exemple de la section suivante).

5.2.3 Exemple

Nous avons choisi d'illustrer notre méthode sur un ensemble d'automobiles qu'il s'agit d'affecter dans l'une des 4 classes suivantes :

- les petites voitures économiques, C_1
- les voiture familiales, C_2
- les berlines de luxe, C_3
- les voitures sportives, C_4

La construction des points de référence (prototypes) des catégories requiert une liste de points d'apprentissage déjà affectés à l'une des quatre catégories. Pour construire un modèle explicatif de l'affectation, nous considérons les critères suivants :

- Prix du véhicule (F)
- Vitesse maximale (km.h⁻¹)
- Consommation à 120 km.h⁻¹ (l)
- Taille du coffre (dm³)
- Accélération (temps (s) pour passer de 0 à 100 km.h⁻¹)
- Freinage (distance (m) de freinage de 130 km.h⁻¹ à 0 km.h⁻¹)
- Confort sonore (bruit en dB)

Il suffit alors de considérer l'image dans l'espace des 7 critères de l'ensemble des véhicules déjà affectés à la catégorie C_t pour constituer l'ensembles des Z_t . Dans ce cas, les performances des véhicules sont extraites des notices techniques que l'on trouve dans la presse spécialisée [21]. En appliquant notre algorithme de construction de prototypes à chacune des ensembles $Z_t (t = 1, \dots, 4)$, nous obtenons les ensembles des points de référence $Y_t (t = 1, \dots, 4)$. En examinant les indices $QR(Z_4/Y_4)$ et surtout $NR(Y_4)$, il apparaît que la meilleure représentation est obtenue par une configuration à trois prototypes qui constituent l'ensemble Y_4 (*cf.* figure 5.3 qui montre l'évolution

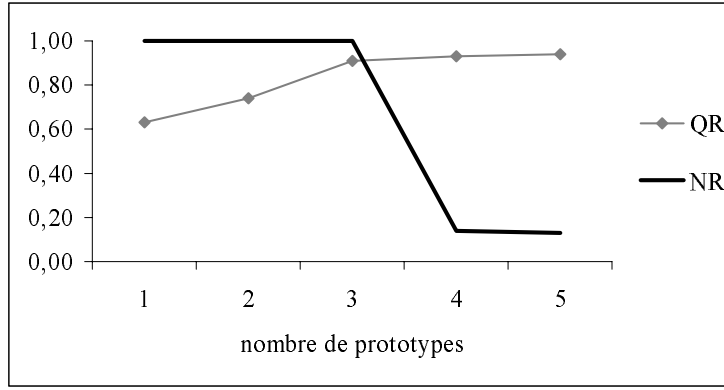


FIG. 5.3 – *Qualité de Représentation et non-redondance*

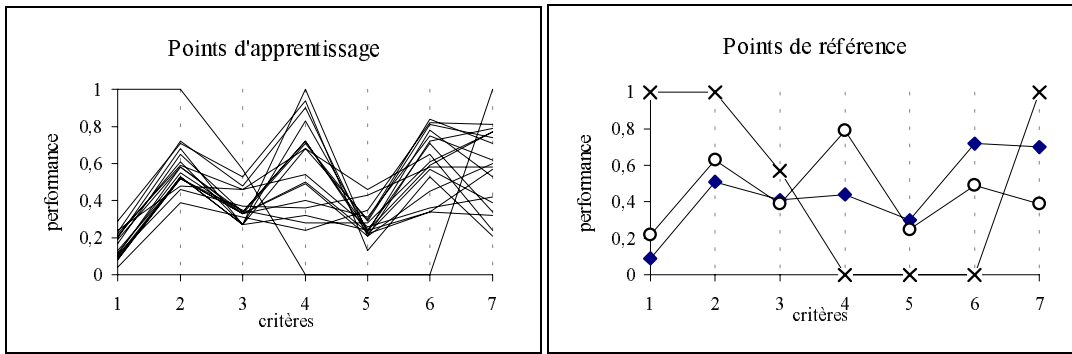


FIG. 5.4 – *Points Z_4 et de Y_4*

du nombre des quantités QR et NR en fonction du nombre de prototypes). Les figures 2 et 3 ci-dessous permettent de comparer les ensembles Z_4 et Y_4 et la table 1 donne le détail des points de référence contenus dans Y_4 .

En procédant de même pour chacune des catégories, on dispose d'un ensemble complet de points de référence. On est donc en mesure de procéder à l'affectation de nouveaux points. Imaginons par exemple qu'un nouveau véhicule se présente avec les performances suivantes : (120000; 200; 9; 350; 9; 80; 73), le résultat de l'affectation utilisant la méthode k -PIP (avec $k = 3$) et les ensembles $Z_i, i = 1, \dots, 4$, et la méthode PIP et les ensembles $Y_i, i = 1, \dots, 4$ est représenté dans le tableau 5.2.

Cette affectation montre la faible différence d'affectation obtenue entre l'utilisation des données brutes comme points de référence, Z , conjointement avec la méthode k -PIP et l'utilisation des prototypes calculés par nuées dynamiques, Y , conjointement avec la méthode PIP. Ce qui montre que malgré la perte d'information induite par la réduction du nombre de points de référence (passage d'une centaine de points d'apprentissage à 12 prototypes), l'affectation n'est pas significativement dégradée.

5.3 Construction de prototypes par algorithme génétique

Comme nous l'avons vu à la section précédente, la méthode des nuées dynamiques permet de trouver rapidement les prototypes d'une catégorie cependant elle impose à l'utilisateur de fixer

Y_4	prix	v.max.	conso.	coffre	accel.	frein.	bruit
proto. 1	106811.11	203.58	8.18	247.56	8.87	81.79	72.72
proto. 2	530200.00	257.20	9.60	82.00	6.00	63.10	75.80
proto. 3	165137.50	216.25	8.59	377.88	8.45	75.85	69.49

TAB. 5.1 – Prototypes de C_4 construits par nuées dynamiques

affectation	économique	familiale	luxe	sport
PIP avec Y_t	0.00	0.94	0.00	0.82
k - PIP avec Z_t	0.00	0.99	0.00	0.82

TAB. 5.2 – Affectation avec Z_t et Y_t

le nombre de prototypes par catégorie. Nous proposons maintenant une approche qui présente le double avantage de s'affranchir de cette contrainte et de pouvoir proposer plusieurs solutions au décideur.

Les algorithmes génétiques sont des techniques d'optimisation de type meta-heuristique fondées sur les principes de l'évolution biologique. Elles ont été introduite par Holland [51] et sont utilisées pour résoudre des problèmes complexes d'optimisation. Elles se distinguent principalement de autres techniques d'optimisation par les 4 points suivants (voir Goldberg [43]) :

- l'utilisation directe des variables du problème à optimiser,
- la capacité à considérer parallèlement plusieurs solutions réalisables (exploration à l'aide d'une population plutôt que d'un point unique),
- une exploration des solutions réalisables à l'aide uniquement de l'évaluation d'une fonction d'évaluation,
- l'utilisation d'opérateurs probabilistes et de règles non-déterministes.

Nous présentons dans cette section une procédure de construction de prototypes fondée sur un algorithme génétique. Cette procédure diffère principalement de la procédure fondée sur les nuées dynamiques par sa capacité à mettre en compétition des représentations d'une catégorie composées d'un nombre différent de prototypes.

5.3.1 Principe des algorithmes génétiques

Solutions et population

Un individu au sens des algorithmes génétiques correspond à une solution réalisable, pour notre problème il s'agit d'un ensemble de prototypes décrivant une catégorie. La procédure doit être utilisée autant de fois qu'il y a de catégories.

Définition 18 (individu) *Un individu s (on parle aussi de solution ou de chromosome) est une liste de prototypes représentée par un vecteur (une matrice de performances car chaque prototype est lui même représenté par un vecteur de performances) :*

$$s = \begin{pmatrix} y_1 & = & (g_1(y_1), g_2(y_1), \dots, g_n(y_1)) \\ y_2 & = & (g_1(y_2), g_2(y_2), \dots, g_n(y_2)) \\ \vdots & & \\ y_m & = & (g_1(y_m), g_2(y_m), \dots, g_n(y_m)) \end{pmatrix}$$

avec : $y_1, \dots, y_m$ les m prototypes de la solution.

Remarque 7 les solutions sont constituées de points de l'espace des critères ce qui signifie que nous utilisons un codage réel des paramètres (real coded voir Goldberg [42]) contrairement à une grande partie des algorithmes qui utilisent un codage binaire.

Définition 19 (gènes et allèles) Les gènes sont les composantes des chromosomes, pour nous si l'on considère des solutions constituées de 3 prototypes, il y aura alors 3 gènes. Les allèles sont les valeurs possibles que les gènes peuvent prendre. Pour notre cas, il s'agit de l'ensemble des profils qui peuvent être construits (exemple : $y \in E_1 \times \dots \times E_n$).

Le nombre de prototypes par catégorie n'est pas fixé a priori, nous pouvons considérer et mettre en compétition des solutions de tailles différentes.

Définition 20 (population) Une population P est un ensemble d'individus, elle est représentée par un vecteur.

Initialisation de la procédure

La phase d'initialisation de l'algorithme consiste en un tirage d'un nombre de solutions de départ formant la population de départ, [43]. Il est possible d'engendrer les solutions de deux manières différentes. Tout d'abord il est possible d'engendrer des actions fictives de manière aléatoire à partir des domaines de définition des fonctions critères. Cependant cela présente l'inconvénient de pouvoir engendrer des actions non pertinentes, c'est-à-dire des actions trop peu réalistes (on parlera d'actions non significatives, à ne pas confondre avec actions fictives, irréalistes ou potentielles qui sont des actions qui ont un sens et qui peuvent présenter un intérêt, [100]). Pour éviter ce problème nous proposons de constituer les solutions de départ à l'aide des actions de l'échantillon d'apprentissage en effectuant un tirage aléatoire. Il est aussi possible de chercher à sélectionner des actions suffisamment différentes entre elles pour considérer des solutions de départ contrastées. La pratique nous a montré que la convergence est accélérée lorsque la population de départ est constituée de solutions contrastées dont les prototypes au sein de chacune sont proches.

5.3.2 Fonction d'évaluation

Principe

La définition de la fonction d'évaluation pose le problème de l'évaluation de la qualité d'une solution. Pour cela nous utilisons les deux mesures : qualité de la représentations ($QR(Z_t/Y_t)$) et non redondance des prototypes ($NR(Y_t)$), présentées à la section précédente. Nous présentons ici le cas où les points d'apprentissage ont une appartenance nette aux catégories ($\forall z \in Z_t, \mu_t^o(z) \in \{0, 1\}$). Nous cherchons à savoir si chaque point de l'ensemble d'apprentissage peut être expliqué par un prototype de la solution à évaluer sans pour autant que les prototypes de la solution ne soient redondants.

Formulation générale

Qualité de la Représentation Nous proposons une formulation de la qualité de la représentation dans le cas où les points d'apprentissage appartiennent de manière non graduelle aux catégories. On a donc si s est une solution supposée représenter la catégorie C_t , z_i un point de l'ensemble d'apprentissage Z_t :

$$QR(Z_t/s) = \varphi_{z \in Z_t}(\mu_t(z))$$

avec $\mu_t(z)$ le degré d'affectation calculé en tenant compte des prototypes de s , φ un opérateur d'agrégation conjonctif ou de compromis,

Les coefficients d'affectation $\mu_t(z)$, peuvent être calculés à l'aide d'une méthode d'affectation multicritère de type filtrage flou. Généralement on utilise la méthode PIP, il suffit alors qu'un point ressemble à un seul prototype d'une catégorie pour qu'il soit affecté à celle ci.

Exemple 32 (opérateur d'agrégation φ) *Les opérateurs de type t-norme, et de manière plus générale ceux ayant le 0 comme élément absorbant doivent être utilisés avec précaution. En effet ce type, d'opérateur implique que dès lors qu'un point de Z est mal ré-affecté la solution est considérée comme mauvaise. Les opérateurs suivants sont donc envisageable mais leur usage ne doit pas être généralisés :*

$$QR(Z_t/s) = \min_{z \in Z_t}(\mu_t(z)), QR(Z_t/s) = \prod_{z \in Z_t} (\mu_t(z))^{\frac{1}{|Z_t|}}$$

Leur intérêt réside dans leur caractère restrictif, ils doivent être utilisés lorsque une solution est jugé inacceptable au premier point mal classé.

En revanche des opérateurs de compromis comme la moyenne arithmétique ne présentent pas ce caractère si restrictif et permettent de discriminer plus finement des solutions ayant quelques points mal ré-affectés :

$$QR(Z_t/s) = \frac{1}{|Z_t|} \sum_{z \in Z_t} \mu_t(z)$$

Pour plus de détails sur les opérateurs utilisables on peut consulter [79], [44] pour les intégrales floues et [116] et [117] pour OWA.

Concision d'une représentation

Cette mesure va nous permettre d'évaluer la redondance des prototypes au sein d'une solution. Avoir plusieurs prototypes très proche (au sens de la relation d'indifférence $\sim$) engendre une redondance d'information. Celle ci n'améliore pas la description de la catégorie au contraire, elle engendre un coût supplémentaire en terme de taille des données. On ne doit donc pas retrouver dans un ensemble de prototypes Y_t deux points x et y parfaitement similaires (tels que $\sim(y, y') = 1$). La non redondance des prototypes $NR(Y_t)$ peut donc être évaluée graduellement par la quantité :

$$NR(Y_t) = 1 - \max_{(y, y') \in Y_t \times Y_t} \{\sim(y, y')\}$$

Agrégation de la non redondance et de la qualité d'une représentation

En pratique, à mesure que le nombre de prototypes augmente, la qualité s'améliore et QR augmente mais la représentation devient plus coûteuse et NR diminue. Ces deux mesures doivent alors être agrégées par une fonction f définie sur $[0, 1]^2$ à valeurs réelles dans $[0, 1]$. Où f est une fonction croissante de ses deux arguments, définie telle que $f(1, 1) = 1$

Une solution s supposée résumer un ensemble d'apprentissage Z_t peut donc être évaluée par la fonction h à maximiser définie telle que :

$$h(s) = f(QR(Z_t/s), NR(s))$$

Exemple 33

$$h(s) = \max(QR(Z_t/s), NR(s))$$

$$h(s) = QR(Z_t/s) \times NR(s)$$

5.3.3 Opérateurs d'évolution utilisés

Deux niveaux d'évolution

Une solution se présente sous la forme d'une matrice, les prototypes étant en ligne et les performances en colonnes par convention. Il est possible d'opérer à la fois sur les prototypes et sur les performances. Nous allons utiliser les deux opérateurs que sont le croisement et la mutation, [43]. Il y aura donc deux types de croisement, un sur les prototypes et un sur les performances, et deux types de mutation, une sur les prototypes et une sur les performances.

Ces deux opérateurs présentent une propriété intéressante, ce sont des opérateurs dont le résultat fait partie de l'ensemble de définition des arguments :

Propriété 13 $\forall (x_1, \dots, x_n) \in X^n \ f(x_1, \dots, x_n) \in X$

Ces opérateurs sont donc parfaitement compatibles avec l'utilisation de critères définis sur des échelles qualitatives. En effet, ces opérateurs n'effectuent que des re-combinaisons des performances.

L'opérateur de croisement (*cross-over*) possède quant à lui une propriété encore plus restrictive, le résultat de l'opération de croisement fait partie de l'ensemble des arguments de cette fonction :

Propriété 14 $f(x_1, \dots, x_n) \in \{x_1, \dots, x_n\}$

Une interprétation peut être donnée à ces opérateurs, l'opérateur de croisement sert à explorer les différentes combinaisons possibles entre des solutions ou des profils, c'est un *opérateur d'intensification* : il effectue une recherche locale entre plusieurs solutions. La mutation quant à elle sert à introduire de manière modérée une certaine nouveauté parmi les solutions et d'explorer des parties de l'espace non encore éventuellement décrites, c'est un *opérateur de diversification* qui permet de sortir des situations d'optima locaux.

Opérateurs agissant sur les prototypes

Dans cette partie, nous considérons que les gènes sont les prototypes.

Opérateur de croisement Le croisement consiste en l'échange de deux ou plusieurs prototypes entre deux solutions.

Définition 21 (Croisement en un point) Soient $s_i = (y_1^i, \dots, y_m^i)$ la solution i formée de n prototypes où y_j^i est le $j^{\text{ième}}$ prototype de la solution i alors on définit χ_k la fonction de croisement associée à $k \in \{1, \dots, m\}$ (k est appelé le point de croisement) si l'on a :

$$\chi_k(s_1, s_2) = (s_{12}, s_{21})$$

avec

$$s_{12} = (y_1^1, \dots, y_k^1, y_{k+1}^2, \dots, y_m^2) \text{ et } s_{21} = (y_1^2, \dots, y_k^2, y_{k+1}^1, \dots, y_m^1)$$

Ce résultat est constitué de deux solutions filles. Si l'on considère 2 points de croisement nous obtenons alors un résultat formé de 8 solutions filles.

Exemple 34 (croisement au niveau des prototypes) Soient deux solutions constituées de 3 prototypes évalués sur 5 critères représentées par les matrices suivantes :

$$s_1 = \begin{pmatrix} 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

$$s_2 = \begin{pmatrix} 5, 5, 5, 5, 5 \\ 5, 5, 5, 5, 5 \\ 5, 5, 5, 5, 5 \end{pmatrix}$$

Le résultat du croisement de ces deux solution avec le point de croisement $k = 2$ fournie les deux solutions :

$$s_{12} = \begin{pmatrix} 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \\ 5, 5, 5, 5, 5 \end{pmatrix}$$

$$s_{21} = \begin{pmatrix} 5, 5, 5, 5, 5 \\ 5, 5, 5, 5, 5 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

Opérateur de mutation L'opérateur de mutation consiste à faire modifier un gène d'une solution avec une probabilité fixée a priori. Cet opérateur nous permet de parcourir une partie de l'espace qui ne serait éventuellement pas couverte par les prototypes des solutions de départ.

Définition 22 Soit y^a un prototype aléatoire et s une solution alors ϖ_i est la fonction de mutation du prototype $i \in \{1, \dots, m\}$ si :

$$\varpi_i(s) = (y_1^1, \dots, y_{i-1}^1, y^a, y_{i+1}^1, \dots, y_n^1)$$

Exemple 35 (mutation au niveau des prototypes) Soient une solution constituées de 3 prototypes évalués sur 5 critères représentées par les matrices suivantes :

$$s = \begin{pmatrix} 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

Un résultat de la mutation du prototype 1 de s est :

$$\varpi_1(s) = \begin{pmatrix} 5, 9, 7, 6, 8 \\ 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

Opérateurs agissant sur les performances

Dans cette partie nous considérons que les performances sont les gènes.

Opérateur de croisement Nous allons effectuer une opération de croisement au niveau de chaque prototype. Une alternative existe. Tout d'abord il est possible d'utiliser le même point de croisement pour tous les prototypes des solutions à croiser. Si l'on désire raffiner la procédure, il est alors possible de choisir un point de croisement par type de prototype.

Définition 23 χ^k est la fonction de croisement associée au point de croisement $k \in \{1, \dots, m\}$ si on a $\chi^k(s_1, s_2) = (s_{12}, s_{21})$ avec :

$$s_{12} = \begin{pmatrix} (g_1(y_1^1), \dots, g_k(y_1^1), g_{k+1}(y_1^2), \dots, g_n(y_1^2)) \\ (g_1(y_2^1), \dots, g_k(y_2^1), g_{k+1}(y_2^2), \dots, g_n(y_2^2)) \\ \vdots \\ (g_1(y_m^1), \dots, g_k(y_m^1), g_{k+1}(y_m^2), \dots, g_n(y_m^2)) \end{pmatrix}$$

$$s_{21} = \begin{pmatrix} (g_1(y_1^2), \dots, g_k(y_1^2), g_{k+1}(y_1^1), \dots, g_n(y_1^1)) \\ (g_1(y_2^2), \dots, g_k(y_2^2), g_{k+1}(y_2^1), \dots, g_n(y_2^1)) \\ \vdots \\ (g_1(y_m^2), \dots, g_k(y_m^2), g_{k+1}(y_m^1), \dots, g_n(y_m^1)) \end{pmatrix}$$

Cette fonction effectue un croisement entre des prototypes en considérant le même point de croisement pour chaque prototype. Il est bien sûr possible d'utiliser un point de croisement par prototype, il y aurait alors m points de croisement. Des prototypes de lignes différentes peuvent aussi être croisés entre eux. Il faudra obligatoirement utiliser le même point de croisement pour ces deux prototypes. Il y aura alors globalement $m/2$ points de croisement pour la fonction χ .

Exemple 36 (croisement) Soient deux solutions constituées de 3 prototypes évalués sur 5 critères représentées par les matrices suivantes :

$$s_1 = \begin{pmatrix} 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

$$s_2 = \begin{pmatrix} 5, 5, 5, 5, 5 \\ 5, 5, 5, 5, 5 \\ 5, 5, 5, 5, 5 \end{pmatrix}$$

Le résultat du croisement de ces deux solutions avec le point de croisement $k = 3$ fournit les deux solutions :

$$s_{12} = \begin{pmatrix} 1, 1, 1, 5, 5 \\ 1, 1, 1, 5, 5 \\ 1, 1, 1, 5, 5 \end{pmatrix}$$

$$s_{21} = \begin{pmatrix} 5, 5, 1, 1, 1 \\ 5, 5, 1, 1, 1 \\ 5, 5, 1, 1, 1 \end{pmatrix}$$

Opérateur de mutation Les mutations sont effectuées au niveau des performances des prototypes. Une performance d'un prototype d'une solution est sélectionnée puis remplacée par une performance aléatoire du domaine de définition du critère considéré. Pour cet opérateur encore il demeure plusieurs éventualités. On peut tout d'abord modifier à chaque mutation, une seule performance d'un seul prototype d'une solution. Il est aussi possible de modifier les performances sur différents critères de différents prototypes d'une solution.

Définition 24 Soient s une solution et y^a un prototype aléatoire alors $\varpi_{i,j}$ est la fonction de mutation de la performance du critère j du profil i si :

$$\varpi_{i,j}(s_1) = \begin{pmatrix} (g_1(y_1^1), \dots, g_j(y_1^1), \dots, g_n(y_1^1)) \\ \vdots \\ (g_1(y_i^1), \dots, g_j(y_i^a), \dots, g_n(y_i^1)) \\ \vdots \\ (g_1(y_m^1), \dots, g_j(y_m^1), \dots, g_n(y_m^1)) \end{pmatrix}$$

Exemple 37 (mutation au niveau des performance) Soient une solution constituées de 3 prototypes évalués sur 5 critères représentées par les matrices suivantes :

$$s = \begin{pmatrix} 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

Un résultat de la mutation du prototype 1 de s au niveau du critère 3 est :

$$\varpi_{1,3}(s) = \begin{pmatrix} 1, 1, 7, 1, 1 \\ 1, 1, 1, 1, 1 \\ 1, 1, 1, 1, 1 \end{pmatrix}$$

Opérateurs de variation de la taille des solutions

Nous allons utiliser l'opérateur de mutation pour faire varier la taille des solutions. Pour cela nous allons définir deux opérateurs, la mutation par ajout et la mutation par suppression. Ces opérateurs vont agir au niveau des prototypes.

Mutation par ajout : La mutation par ajout consiste à ajouter un prototype à une solution dès lors que le nombre maximum de prototypes n'est pas atteint.

Définition 25 Soient y^a un prototype aléatoire et s une solution alors ϖ_i^+ est la fonction de mutation par ajout définie par :

$$\varpi_i^+(s) = (y_1^1, \dots, y_i^1, y^a, y_{i+1}^1, \dots, y_n^1)$$

Mutation par suppression : La mutation par suppression consiste à retirer un prototype d'une solution dès lors que sa taille est supérieure ou égale à 2.

Définition 26 Soit s une solution alors ϖ_i^- est la fonction de mutation par suppression du prototype i définie par :

$$\varpi_i^-(s) = (y_1^1, \dots, y_{i-1}^1, y_{i+1}^1, \dots, y_n^1)$$

Soit *nbIteration* le nombre d'itération de la procédure

DÉBUT

 tirage d'une population initiale

POUR *i* **ALLANT DE** 1 **À** *nbIteration*

 croisement prototypes

 croisement performances

 mutation prototypes

 mutation performances

 mutation par ajout

 mutation par suppression

 sélection

FINPOUR

FIN

FIG. 5.5 – Procédure générale de l'algorithme génétique

5.3.4 Algorithmes

Nous présentons dans cette section les algorithmes que nous avons utilisés. Tout d'abord la procédure générale qui montre le principe des algorithmes génétiques figure 5.5. Nous utilisons dans cette procédure les deux types d'opérateur de croisement que nous avons définis auparavant (croisement des prototypes et croisement des performances), la procédure de croisement générale est présentée figure 5.6. La variable *nbGenes* prend une signification différente suivant le type de croisement considéré. Dans la procédure de croisement de prototypes, *nbGenes* correspond au nombre de prototypes de solutions à croiser. Dans ce type de croisement les solutions *pop[j]* et *pop[i]* peuvent être de taille différente, c'est pourquoi le point de croisement considéré prend ses valeur dans l'intervalle $[1, \min(pop[i], pop[j])]$. Dans la procédure de croisement des performances, *nbGenes* correspond au nombre de critères. Par ailleurs nous considérons quatre types de mutations, la mutation des prototypes, la mutation des performances, la mutation par ajout et la mutation par suppression. De même que pour le croisement, la variable *nbGenes* prend comme valeur le nombre de prototypes pour la mutation des prototypes et le nombre de critères pour la mutation des performances.

5.3.5 Exemple

Pour illustrer notre procédure génétique nous avons choisi un exemple d'affectation de véhicules à des catégories provenant de la presse spécialisées [21]. L'objectif est de construire les prototypes des types de voitures considérées. Nous considérons un ensemble d'apprentissage *Z* constitué de 400 points d'apprentissage et 4 catégories : les petites voitures économiques C_1 , les voitures familiales C_2 , les voitures de luxe C_3 et les voitures sportives C_4 . Les véhicules sont affectés en considérant les critères suivant :

- prix (FF),
- vitesse maximale (km./h.),
- consommation (l./100 km.),
- taille du coffre (dm³),
- 1000 m. départ arrêté (sec.),
- reprises de 40 km./h. à 120 km./h. sur le dernier rapport de boîte (sec.),

Soient

- pop un tableau de solutions de taille $length$
- $alea_{\mathbb{R}}(x)$ une fonction retournant une valeur aléatoire réel de l'intervalle $[0, x]$
- $alea_{\mathbb{N}}(n)$ une fonction retournant une valeur aléatoire entière de l'intervalle $[1, n]$
- Pr_{xo} la probabilité de croisement

DÉBUT

$flag = \text{true}$

POUR i **ALLANT DE** 1 **À** $pop.length$

$Pr = \text{probabilité aléatoire}$

SI $Pr < Pr_{xo}$

ALORS

$flag = \neg flag$

SI $flag$

ALORS

$k = alea_{\mathbb{N}}(\min(pop[j].nbGenes, pop[i].nbGenes))$

$(s_{ji}, s_{ij}) = \chi^k(pop[j], pop[i])$

intégrer s_{ji} et s_{ij} dans pop

SINON $i = j$

FINSI

FINSI

FINPOUR

FIN

FIG. 5.6 – *Procédure de croisement*

Soient :

- pop un tableau de solutions de taille $length$
- $alea_{\mathbb{R}}(x)$ une fonction retournant une valeur aléatoire réel de l'intervalle $[0, x]$
- $alea_{\mathbb{N}}(n)$ une fonction retournant une valeur aléatoire entière de l'intervalle $[1, n]$
- Pr_{mut} la probabilité de mutation

DÉBUT

POUR j **ALLANT DE** 1 **À** $pop.length$

$Pr = alea_{\mathbb{R}}(1)$

SI $Pr < Pr_{mut}$

ALORS

$k = alea_{\mathbb{N}}(nbGenes)$

$s = \varpi_k(pop[j])$

intégrer s dans pop

FINSI

FINPOUR

FIN

FIG. 5.7 – *Procédure de mutation*

– autonomie (km).

Ces critères sont construits sur un modèle de pseudo-critère à seuils variables. Chaque catégorie possède son propre jeu de paramètres (poids et seuils d'indifférence, préférence et veto).

Regardons maintenant les résultats obtenus pour la catégorie C_1 en considérant $Pr_{mut} = 0.1$ et $Pr_{xo} = 0.3$. Le nombre maximum de prototypes a été fixé à 9 et la fonction d'évaluation utilise la moyenne arithmétique pour évaluer la qualité de la représentation QR . Nous obtenons 4 prototypes (voir tableau 5.3). Les temps de calculs sont fonctions des probabilités de croisement et de mutation, ils sont ici de l'ordre de quelques minutes et n'excèdent pas 15 minutes sur un Pentium® 166 Mhz muni de 32 Mo de mémoire vive tournant sous MS® Windows 95.

Les 4 prototypes obtenus représentent 4 types de voitures économiques bien différentes. y_1 est typiquement une petite voiture économique (citadine), peu chère et peu performante. Au contraire, y_2 représente un type de petite voiture beaucoup plus dynamique mais beaucoup moins économique (plus performante mais plus chère). y_3 est un prototype bien différent de part la taille de son coffre, il représente typiquement un petit break pas chère (grand coffre, piètres performances et peu coûteux pour un break). Enfin le prototype y_4 , il représente la petite voiture diesel. Il est plus chère que les autres mais consomme beaucoup moins et ne possède pas de bonnes performances (vitesse maxi. et accélération) sauf pour ce qui est des reprises ce qui est typique des moteurs diesel (faible puissance mais couple important). Au total notre procédure nous a permis d'isoler 4 types réellement différents de voitures économiques

proto.	prix	v. max.	conso.	coffre	1000 m.	reprises	autonomie
y_1	57781	148	6.30	170	36.8	39.9	479
y_2	73800	172	6.80	215	33.4	38.8	490
y_3	70979	138	6.46	1271	39.6	46.3	620
y_4	83894	141	5.41	444	39.2	35.7	595

TAB. 5.3 – Prototypes de C_1

Intéressons nous maintenant au rôle que jouent les probabilités de mutation et de croisement. Pour cela nous avons fait tourner notre algorithme plusieurs fois en faisant varier les probabilités. Nous présentons les résultats obtenus sur la catégorie C_1 constituée d'une cinquantaine de points d'apprentissage. Le nombre d'itérations nécessaire à une convergence est de l'ordre de 1000. Nous avons fait varier les probabilités de 0 à 1 par pas de 0.01 et chaque jeu de probabilités a fait l'objet de 10 exécutions. Nous présentons un résumé de ces résultats dans le tableau 5.4, où les valeurs correspondent à la moyenne pour les dix exécutions de la fonction d'évaluation du meilleur individu obtenu à chaque fois.

Tout d'abord remarquons que les plus mauvais résultats sont obtenus dès lors qu'un des deux opérateurs n'est pas actif (probabilité nulle) cf. ligne et colonne en italique dans le tableau 5.4. En revanche, dès que les deux opérateurs peuvent être utilisés conjointement, même avec une faible

		Probabilité de mutation										
		0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	1
Probabilité de croisement	0	<i>0.5417</i>	<i>0.4913</i>	<i>0.5517</i>	<i>0.5509</i>	<i>0.6072</i>	<i>0.5491</i>	<i>0.5528</i>	<i>0.5390</i>	<i>0.5436</i>	<i>0.5483</i>	<i>0.5451</i>
	0.1	<i>0.5462</i>	0.7544	0.7615	0.6750	0.6681	0.6325	0.6153	0.6069	0.5901	0.5902	0.5897
	0.2	<i>0.5857</i>	0.7265	0.7736	0.7506	0.7391	0.7203	0.6656	0.6814	0.6656	0.6685	0.6343
	0.3	<i>0.5536</i>	0.7732	0.7724	0.7665	0.7753	0.7373	0.7188	0.7120	0.6993	0.6775	0.6556
	0.4	<i>0.5299</i>	0.7274	0.7880	0.7022	0.7733	0.7531	0.7522	0.7100	0.7169	0.7039	0.6516
	0.5	<i>0.5410</i>	0.7442	0.7430	0.7105	0.7435	0.7310	0.7563	0.7377	0.7220	0.7242	0.6550
	0.6	<i>0.5329</i>	0.7487	0.7415	0.7577	0.7870	0.8290	0.7710	0.7719	0.7386	0.7134	0.6030
	0.7	<i>0.5535</i>	0.722	0.7386	0.7283	0.7675	0.7466	0.7356	0.7189	0.7459	0.7426	0.6621
	0.8	<i>0.5422</i>	0.7412	0.7280	0.7450	0.7493	0.7360	0.7586	0.7538	0.7445	0.6949	0.6118
	0.9	<i>0.5339</i>	0.7505	0.7445	0.7755	0.7580	0.7664	0.7310	0.7599	0.7125	0.7147	0.6421
	1	<i>0.5184</i>	0.6948	0.7345	0.7390	0.7588	0.7640	0.7347	0.7285	0.7423	0.7261	0.6807

TAB. 5.4 – Fonction d'évaluation

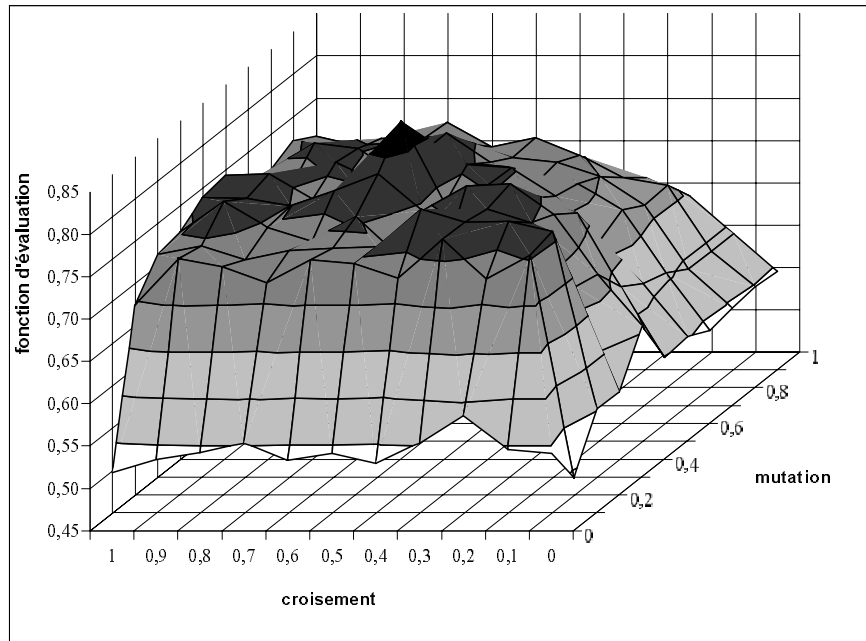


FIG. 5.8 – *Fonction d'évaluation*

probabilité, la valeur de la fonction d'évaluation augmente significativement (de 0,55 à 0,75, cf. ligne 2 et colonne 3). Il y a un effet de synergie entre les deux types d'opérateur, l'utilisation des deux opérateurs est nécessaire pour obtenir de bons résultats. Les meilleurs résultats sont alors obtenus avec une probabilité de mutation comprise entre 0,1 et 0,5 et une probabilité de croisement entre 0,2 et 0,6.

Cependant de trop fortes pour valeurs pour les probabilités de croisement et de mutation ne sont pas nécessairement bénéfiques. En effet, dès que la probabilité de mutation dépasse 0,7 la fonction d'évaluation diminue sensiblement. Lorsque les probabilités sont trop fortes, la population évolue trop vite et la plupart des individus, y compris le meilleur, sont éliminés à chaque génération. On voit sur la figure 5.8, qui représente la valeur de la fonction d'évaluation en fonction des probabilités de croisement et de mutation, que la fonction d'évaluation admet un plafond lorsque les probabilités ne prennent pas des valeurs extrêmes.

Lorsque l'on choisit les probabilités de croisement et de mutation, nous devons aussi prendre en compte le coût des opérateurs d'évolution. Le croisement est largement plus coûteux que la mutation au point que celle-ci peut être considérée comme négligeable. De ce fait on conseille généralement l'utilisation de valeurs faibles pour la probabilité de croisement lorsque le temps de calcul est une variable importante comme pour une utilisation interactive par exemple.

TROISIÈME PARTIE

Mise en œuvre

Chapitre 6

Logiciel FFI

Résumé

Nous présentons dans ce dernier chapitre la mise en œuvre des méthodes que nous avons proposées au chapitre 3. Les méthodes PIP et k-PIP ainsi que la méthode des k plus proches voisins flous de [58] ont été implantées. Ce logiciel permet donc entre autre, de gérer des catégories multiprofiles possédant chacune leur propre jeu de paramètres préférentiels (seuils et poids), d'effectuer une affectation floue et d'obtenir une visualisation graphique des résultats. De plus, ce logiciel présente la particularité d'être multi-plateforme.

Mots clés : Filtrage Flou par Indifférence, java, logiciel.

6.1 Introduction

Le logiciel FFI (Filtrage Flou par Indifférence ou Fuzzy Filtering by Indifference) est actuellement en cours de développement conjointement avec un étudiant, Alvin Sushilen Pillay de l'IUP MIAGE de l'Université Paris IX-Dauphine. Un prototype est disponible sous forme compressée à l'adresse <http://www.lamsade.dauphine.fr/~henriet/FFI>.

Il est développé en Java (version 1.2) en utilisant l'interface graphique **swing** sous l'environnement de développement intégré Kawa 3.13. L'utilisation du langage Java permet une utilisation multi-plateforme non dépendante du système d'exploitation.

6.2 Méthodes mises en œuvre

Le logiciel FFI présente une implémentation des la méthode de filtrage flou présentée au chapitre 3 et de la méthode des k plus proches voisins flous de Keller *et al.* [58].

Filtrage Flou par Indifférence

- Les critères sont construits sur un modèle à 4 seuils présenté au chapitre 3. La concordance entre x et y sur le critère j est évaluée par la formule :

$$C_j^\sim(x, y) = \begin{cases} 0 & \text{si } x_j - y_j \leq -p_j \\ \frac{1}{2} + \frac{1}{2} \sin \frac{\pi}{p_j - q_j} \left(x_j - y_j + \frac{p_j + q_j}{2} \right) & \text{si } -p_j \leq x_j - y_j \leq -q_j \\ 1 & \text{si } |x_j - y_j| \leq q_j \\ \frac{1}{2} - \frac{1}{2} \sin \frac{\pi}{p_j - q_j} \left(x_j - y_j - \frac{p_j + q_j}{2} \right) & \text{si } q_j \leq x_j - y_j \leq p_j \\ 0 & \text{sinon} \end{cases} \quad (6.1)$$

La discordance entre x et y sur le critère j est évaluée par la formule :

$$D_j^\sim(x, y) = \begin{cases} 1 & \text{si } x_j - y_j \leq -v_j^+ \\ \frac{1}{2} - \frac{1}{2} \sin \frac{\pi}{v_j^+ - v_j^-} \left(x_j - y_j + \frac{v_j^+ + v_j^-}{2} \right) & \text{si } -v_j^+ \leq x_j - y_j \leq -v_j^- \\ 0 & \text{si } |x_j - y_j| \leq v_j^- \\ \frac{1}{2} + \frac{1}{2} \sin \frac{\pi}{v_j^+ - v_j^-} \left(x_j - y_j - \frac{v_j^+ + v_j^-}{2} \right) & \text{si } v_j^- \leq x_j - y_j \leq v_j^+ \\ 1 & \text{sinon} \end{cases} \quad (6.2)$$

- La relation d'indifférence peut être paramétrée, nous avons implémenté les variantes suivantes :

Concordance : Nous proposons deux opérateurs d'agrégation, la moyenne arithmétique et le min. :

$$C^\sim(x, y) = \sum_{i=1}^n \omega_i \times C_i^\sim(x, y)$$

$$C^\sim(x, y) = \min_{i=1}^n C_i^\sim(x, y)$$

Discordance : Nous proposons aussi deux opérateurs d'agrégation, la duale géométrique et le max. :

$$D^\sim(x, y) = 1 - \prod_{i=1}^n (1 - D_i^\sim(x, y))$$

$$D^\sim(x, y) = \max_{i=1}^n D_i^\sim(x, y)$$

Indifférence : Les indices de concordance et de discordance doivent être agrégés pour donner l'indifférence selon le principe de concordance et non-discordance, là encore nous proposons deux formulations :

$$\sim(x, y) = C^\sim(x, y) \times (1 - D^\sim(x, y))$$

$$\sim(x, y) = \min(C^\sim(x, y), (1 - D^\sim(x, y)))$$

- Nous avons implanté les deux formulations présentées au chapitre 3, à savoir l'affectation avec renforcement et sans renforcement.

– *Affectation sans renforcement*

$$\mu_t(x) = \begin{cases} \max_{y \in V_k(x) \cap Y_t} \{\sim(x, y)\} & \text{si } V_k(x) \cap Y_t \neq \emptyset, \\ 0 & \text{sinon.} \end{cases} \quad (6.3)$$

– *Affectation avec renforcement*

$$\mu_t(x) = \begin{cases} 1 - \prod_{y \in V_k(x) \cap Y_t} (1 - \sim(x, y)) & \text{si } V_k(x) \cap Y_t \neq \emptyset, \\ 0 & \text{sinon.} \end{cases} \quad (6.4)$$

- Les catégories sont représentées par des points de référence centraux, il s'agit de catégories multiprofiles. Chaque catégorie possède son propre jeu de paramètres préférentiels, à savoir les seuils d'indifférence, de préférence, de veto minimum et de veto maximum ainsi que les poids des critères.

***k* plus proches voisins flous**

Par ailleurs nous avons aussi implanté la méthodes de *k* plus proches voisins flous afin de pouvoir comparer aisément le résultat de l'affectation utilisant des informations préférentielles (filtrage par préférence) et de celle, plus classiques, des *k*-ppv. Dans ce cas aussi la méthode est paramétrable.

- Plusieurs types de distance sont disponibles, nous avons implantés les normes suivantes :

- norme 1 : $\|x - y\|_1 = \sum_{i=1}^n |x_i - y_i|$,

- norme 2 : $\|x - y\|_2 = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$,

- norme ∞ : $\|x - y\|_\infty = \max_{i=0}^n |x_i - y_i|$.

- De même que pour le filtrage par indifférence nous avons implanté deux méthodes d'affectation :

- *affectation sans renforcement* : Il s'agit de la méthode du plus proche prototype, l'objet x est affecté à la classe qui possède le prototype dont il est le plus proche,

- *affectation avec renforcement* : Le degré d'affectation $\mu_t(x)$ à la classe C_t est donné par la formule [58] :

$$\mu_t(x) = \frac{\sum_{j=1}^k u_{tj} (1/\|x - y^j\|^2)}{\sum_{j=1}^k (1/\|x - y^j\|^2)} \quad (6.5)$$

avec $u_{tj} = 1$ si y_j est prototype de C_t et $u_{tj} = 0$ sinon.

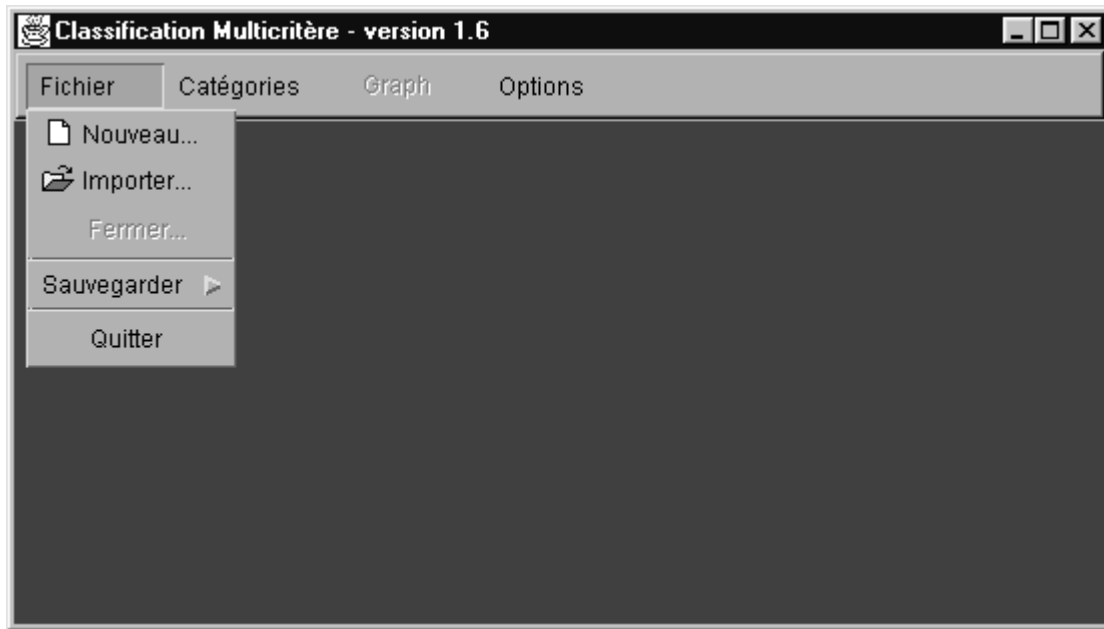


FIG. 6.1 –

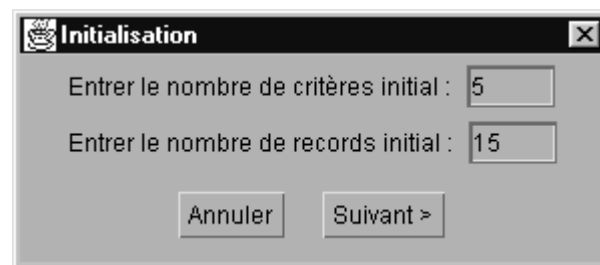


FIG. 6.2 –

6.3 Fonctionnalités

6.3.1 Démarrage

L'application doit être lancée avec le noyau java 1.2 et utilise l'interface graphique `swing`.

6.3.2 Création des fichiers d'actions à classer et de prototypes

Deux fichiers sont utilisés, un pour les actions à classer et un pour la représentation des catégories (prototypes, seuils et poids). L'option **Nouveau** du menu **Fichier** permet la création de nouveaux fichiers (figure 6.1). Il est alors demandé à l'utilisateur d'entrer le nombre d'actions du fichier (**records**) et le nombre de critères (figure 6.2). Les prototypes sont gérés dans le menu **catégories**.

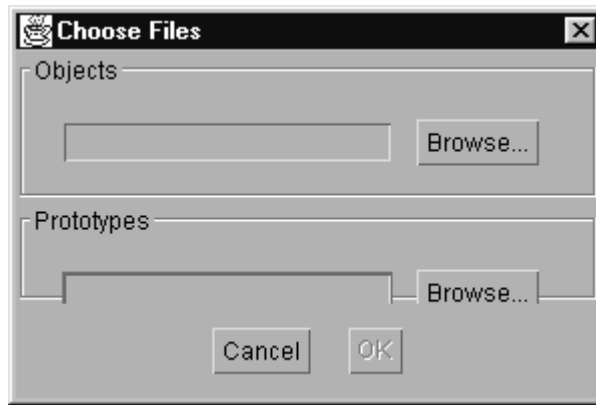


FIG. 6.3 –

6.3.3 Chargement de Fichiers

Les fichiers d'actions à classer et de catégories peuvent être chargés avec l'option **Importer** du menu **Fichier**. Les deux fichiers **Objects** et **Prototypes** doivent être chargés (figure 6.3), les fichiers peuvent être sélectionnés dans un menu déroulant (figure 6.4). Cette opération ouvre deux fenêtres, la fenêtre **Base de données** qui contient les actions à classer et la fenêtre **Catégories** qui contient les prototypes et le paramètres préférentiels (figure 6.5).

6.3.4 Sauvegarde

Les données peuvent être sauvegardées à tout moment soit dans le menu **Fichier** général, soit dans le menu **Fichier** des fenêtres **Base de données** et **Catégories**.

6.3.5 Gestion des données : actions, catégories et critères

Les actions à classer sont représentées dans la fenêtre **Base de données** (figure 6.5). Il est possible de modifier les performances des actions directement dans cette fenêtre, d'ajouter de nouvelles actions et d'en retirer en utilisant les boutons **Ajouter objet** et **Enlever objet(s)**.

La fenêtre **Catégories** permet de gérer les catégories. Celles-ci sont représentées par des onglets (par exemple sur la figure 6.5 trois catégories sont représentées : Familiale, Sport et Luxe). Les prototypes peuvent être modifiés directement dans la fenêtre. Il est possible d'ajouter et de supprimer catégories et prototypes directement dans la fenêtre. Le bouton **Propriétés** permet de modifier les paramètres préférentiels (seuils et poids) (figure 6.6).

Des critères peuvent être ajoutés ou supprimés en utilisant l'option **critères** du menu **Option** (figure 6.7).

6.3.6 Visualisation graphique des actions

Les actions peuvent être visualisées graphiquement par rapport aux prototypes de chaque catégorie en sélectionnant l'option **Objets/Profiles** du menu **Graph** (figure 6.8). On sélectionne alors l'action à visualiser et la catégories de référence (figure 6.9). On obtient le graphique de la figure 6.10. Les critères sont en abscisse et le niveau de performance en ordonnée est normé.

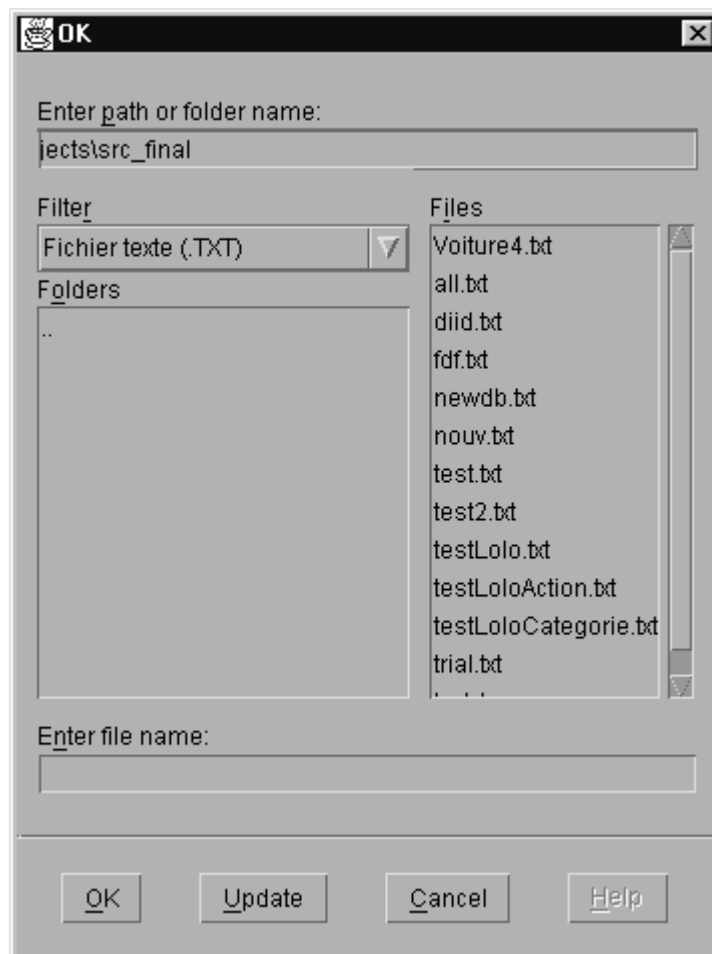


FIG. 6.4 –



FIG. 6.5 –

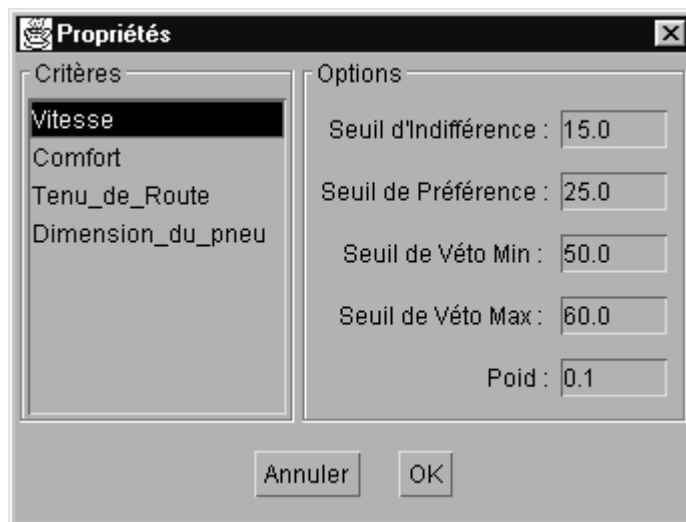


FIG. 6.6 –



FIG. 6.7 –

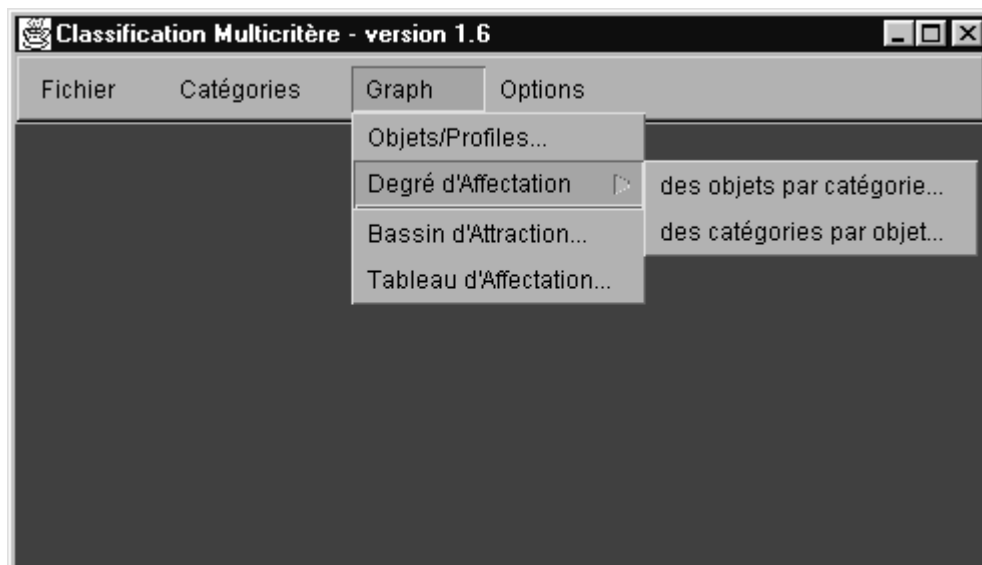


FIG. 6.8 –

Graphe des Profiles par Rapport au Objets [X]

Sélection des paramètres

Category : Familiale [Choisir...]

Record : Laguna [Choisir...]

[Annuler] [OK]

FIG. 6.9 –

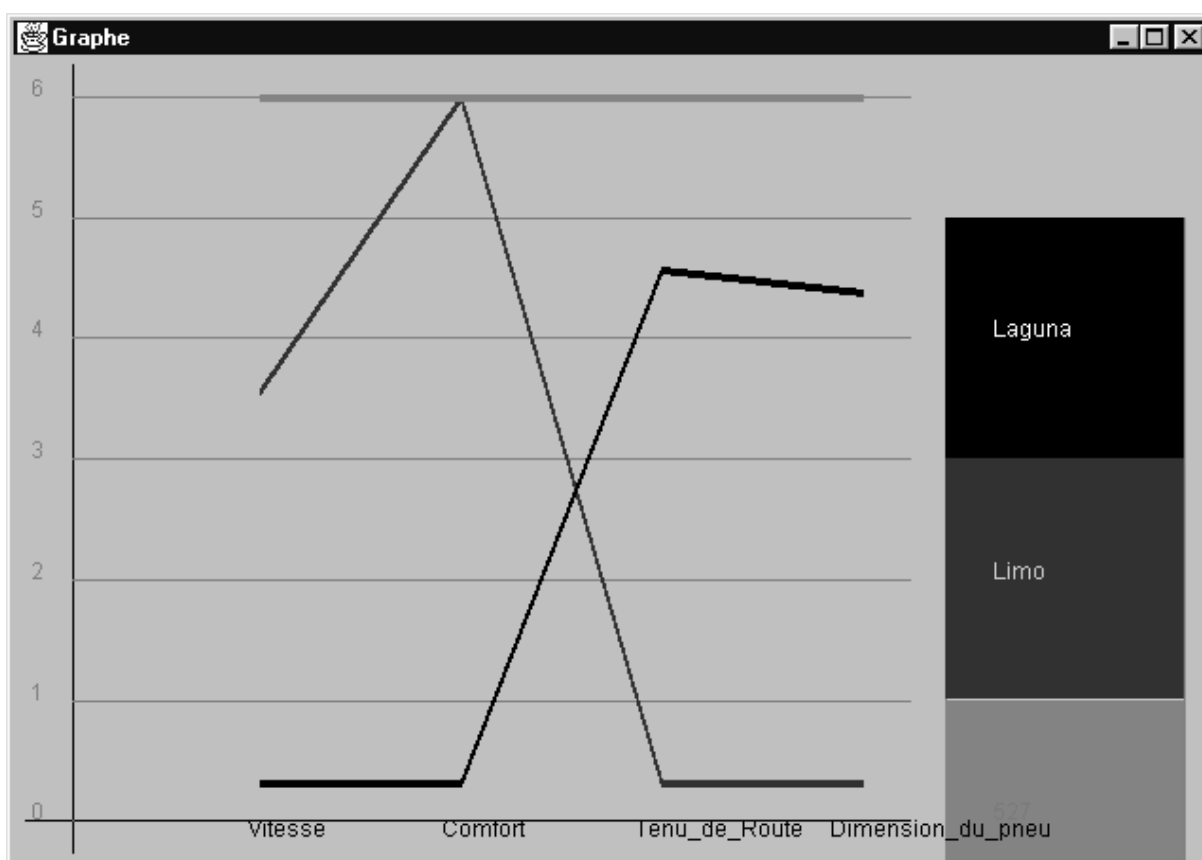


FIG. 6.10 –



FIG. 6.11 –

6.3.7 Affectation des actions

Il est possible de visualiser les résultats de l'affectation de trois façons différentes : histogrammes et tableaux représentant l'affectation d'une action à toutes les catégories (menu `menu` → `Degré d'affectation` → `des catégories par objet`), histogrammes et tableaux représentant l'affectation de toutes les action à une catégories (menu → `Degré d'affectation` → `des objets par catégorie`) et un tableau représentant l'affectation de toutes les actions à toutes les catégories (menu `menu` → `Tableau d'affectation`). Ainsi il est possible d'observer graphiquement à la fois de quelle façon une action va être affectée à toutes les catégories (figure 6.13) et aussi de quelle façon toutes les actions se répartissent dans une catégorie (figure 6.12).

À chaque représentation l'utilisateur a le choix de la méthode (*cf.* §6.2), voir figure 6.11. Il peut ainsi choisir la méthode des k -ppv avec différents types de distance ou une méthode multicritère (PIP ou k -PIP) avec deux types de concordance, discordance et indifférence. Il doit ensuite choisir le nombre de prototypes k qui interviennent dans la méthode avec renforcement. Les résultats de l'affectation sont alors présentés sous forme d'histogrammes et de tableaux (figure 6.12).

6.4 Évolution

Le prototype que nous avons présenté a été développé en Java, il est de ce fait portable sous diverses plates-formes (Windows, Linux, MacOS, ...) du moment qu'un noyau java existe sous le système d'exploitation considéré. Cependant même si java est un langage qui tend à se généraliser, son utilisation n'est pas nécessairement familière à tous les utilisateurs. Nous envisageons de faire évoluer notre application afin de la rendre accessible à travers un *fureteur* via le protocole `http`. Ceci permet une utilisation aisée qui ne nécessite ni installation particulière ni système d'exploitation particulier, la seule exigence étant de posséder un *fureteur* et une connexion.

De plus nous envisageons de porter nos méthodes de construction (pour l'instant développées en C++) en Java afin de les intégrer au sein du logiciel FFI. Dans cette optique les traitements seront effectués sur le serveur et affichés sous forme de page `html`. Les seuls données échangées sont

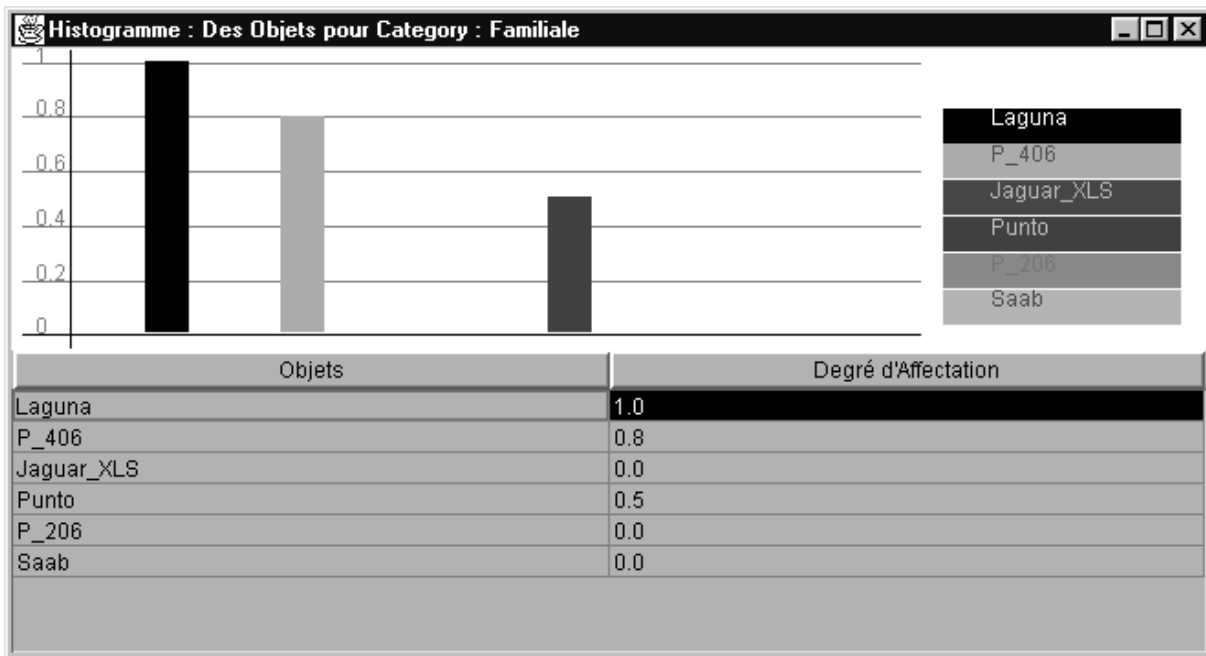


FIG. 6.12 –

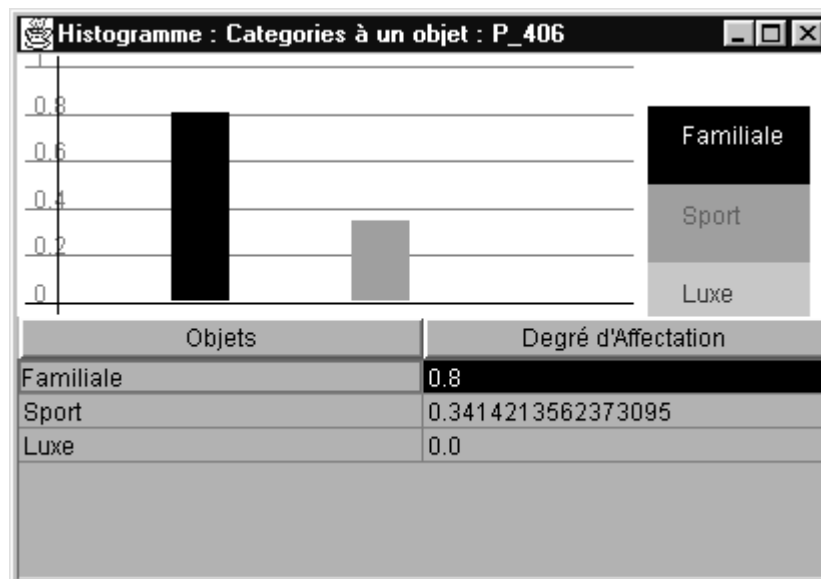


FIG. 6.13 –

ainsi les fichiers de données de l'utilisateur et la page de résultat.

Conclusion

Nous avons défini les méthodes de classification multicritère comme permettant d'affecter des actions potentielles à des catégories prédéfinies. De multiples problèmes font appel à ce type de méthode et de nombreuses approches sont proposées pour les résoudre. Celles-ci font l'objet du chapitre 2. Parmi ces méthodes, rares sont celles qui sont fondées sur la modélisation des préférences (utilisation de relations de préférence et d'indifférence) : nous nous sommes donc intéressés à l'approche de la classification dans un cadre d'aide multicritère à la décision.

Nous avons tout d'abord défini une famille de méthodes baptisées Méthodes Multicritères d'Affectation par Comparaisons Binaires (MMACB) et les propriétés qui les accompagnent ; certaines de ces propriétés sont caractéristiques des MMACB dans le sens où elles s'appliquent à toute méthode de la famille, tandis que d'autres sont spécifiques au type de filtrage utilisé, par préférence ou par indifférence. Nous nous sommes intéressés dans ce cadre aux méthodes fondées sur une relation d'indifférence floue, réflexion qui nous a mené à l'élaboration d'une famille de méthodes qui permettent suivant l'instance choisie de s'adapter à différents types de problème. Les méthodes issues de cette famille diffèrent notamment des autres méthodes multicritères de classification par la nature des points de référence qu'elles permettent de prendre en compte, ainsi suivant que l'on désire prendre en compte des points de référence typiques ou plutôt des points exemples il est possible de choisir une instance avec ou sans renforcement. Leur apport essentiel par rapport aux méthodes multicritères existantes est d'offrir la possibilité d'effectuer une affectation graduelle. Cela permet principalement d'offrir une information plus riche qu'une simple appartenance binaire, cette information graduelle pouvant aussi bien traduire une incertitude sur l'appartenance que l'appartenance à des catégories mal définies. Mais aussi d'offrir une plus grande robustesse des méthodes en évitant qu'une petite différence de performance puisse faire basculer une action à l'intérieur ou hors d'une catégorie (voir effet de seuil section 3.3.1). Par ailleurs elles offrent aussi la possibilité de considérer des catégories multiprofiles représentées par des profils centraux (et éventuellement des points de référence excluants). Notons également la souplesse de nos méthodes qui ne formulent aucune hypothèse restrictive quant à la structure des catégories (ex. structure d'ordre, partition). Enfin, les méthodes que nous avons proposées permettent, et cela est tout particulièrement intéressant et utile dans le cadre de l'AD, d'obtenir une explication de l'affectation en générant automatiquement l'interprétation des calculs effectués. Afin d'illustrer notre famille de méthodes nous avons donné deux instances, les méthodes du Plus Indifférent Prototype (PIP) et des k Plus Indifférents Prototypes (k -PIP).

Nous nous sommes intéressés, dans une seconde partie, à la construction des points de référence des catégories et des coefficients d'importance des critères. En effet, il est souvent plus aisé pour un décideur de donner des exemples de points déjà affectés aux catégories que de produire directement et explicitement les paramètres de méthodes, paramètres dont il n'appréhende pas nécessairement la véritable signification. Aussi proposons-nous deux procédés de construction des points de référence des catégories, l'un fondé sur l'utilisation d'un algorithme de type nuées dynamiques, l'autre sur l'utilisation d'un algorithme génétique. La première solution présente par rapport à la seconde l'avantage d'une convergence nettement plus rapide qui favorise la construction interactive. Elle nécessite en revanche de fixer le nombre de points de référence par catégories. Au contraire, la

procédure de type meta-heuristique permet non seulement de construire les prototypes d'une catégorie mais aussi le nombre de prototypes, elle est toutefois plus coûteuse. En ce qui concerne les coefficients d'importance des critères, nous proposons une modélisation sous forme de programme mathématique pouvant être rapidement résolu s'inscrivant naturellement dans un processus interactif.

Vient ensuite la mise en œuvre. Nous proposons pour cela une implémentation de notre méthode en Java sous la forme du logiciel FFI développé conjointement avec un étudiant de la MIAE Dauphine. Ce logiciel permet la classification suivant les méthodes PIP et k -PIP, ainsi que la méthode de k plus proches voisins flous [58] utilisant différents types de distance : il constitue ainsi un outil de comparaison des résultats d'une affectation avec une méthode classique de classification et nos méthodes fondées sur l'aide multicritère à la décision. D'un point de vue pratique, il est à noter que FFI propose une visualisation graphique des résultats et permet de pouvoir traiter des données de grande taille.

Nous nous sommes intéressés aux méthodes de classification multicritères générales en définissant le concept de MMACB. Cependant notre recherche s'est concentrée sur les méthodes de filtrage par indifférence, nous n'avons à cet effet exploré qu'un pan de ce type de méthodes. Le travail que nous avons présenté initie donc le domaine des méthodes multicritères de classification et, proposant des méthodes, il pose au passage certaines interrogations qui peuvent, entre autres, donner lieu aux axes de recherche suivants :

- Nous avons évoqué la possibilité d'effectuer du *filtrage combiné*, c'est-à-dire d'utiliser conjointement des relations de préférence et d'indifférence afin de pouvoir représenter des catégories à la fois à l'aide de profils limites et de profils centraux. Ce type de filtrage permet de considérer des catégories complexes dont la représentation n'est pas aisée, voire impossible, en utilisant un seul type de profil.
- Nous avons aussi abordé l'existence de points de référence excluants (ou anti-prototypes) et leur exploitation à en utilisant une extension floue d'une logique à 4 valeurs [83]. Il serait intéressant d'approfondir cette voie afin de pouvoir l'intégrer cet aspect dans la méthode FFI. De plus l'utilisation de tels points de référence peut aussi être envisagée dans le cadre du filtrage par préférence.
- Nous avons utilisé la méthode des nuées dynamiques afin de construire les prototypes des catégories. Il semble que cette méthode pourrait également être utilisée dans le cas du filtrage par préférence et de la construction de profils limites en adaptant la fonction et l'espace de représentation.
- Nous avons proposé de construire les poids des critères par la résolution d'un programme mathématique au sein d'un processus interactif. Cette optique pourrait être étendue aux autres paramètres préférentiels que sont les seuils des critères. Pour ce qui est des poids des critères, une procédure fondée sur un algorithme génétique a été proposée conjointement avec un étudiant de DEA [65]. Les algorithmes génétiques présentent le double avantage d'être utilisables sur des type de problème très variés, permettant l'intégration de fonctions complexes (*cf.* fonctions critère et agrégation de concordance et discordance), et de proposer non pas une solution mais un panel de solutions au décideur. En effet, un ensemble de solutions contrastées peut aussi servir de base de discussion dans la phase de fixation des paramètres préférentiels.
- Enfin nous pouvons envisager l'utilisation des techniques que nous avons développées dans un cadre plus large que celui de la classification. Les techniques de filtrage flou par indifférence peuvent en effet être utilisées dans le cadre de la *décision collective* et du *conseil par collaboration* (*cf.* projet Film Conseil <http://www-poleia.lip6.fr/fconseil/> [84]). Le problème de décision collective consiste à affecter des actions à des catégories en considérant que les critères correspondent à l'avis de différents individus. L'activité de conseil apparaît

comme une vision duale du problème d'affectation. Elle consiste à évaluer l'appartenance des individus en prenant comme critères les actions du problème précédent.

Enfin, même si d'un point de vue théorique certains points méritent d'être approfondis, nous pensons avoir donné un cadre d'analyse générale pour ce qui est des problèmes de classification en aide multicritère à la décision. De plus il nous semble important de souligner le caractère opérationnel des méthodes que nous avons proposées. Une démarche d'aide à la décision passe nécessairement par l'étude théorique de modèles (d'agrégation, de préférence et d'affectation dans notre cas) mais aussi par la mise en œuvre de ceux-ci à travers, par exemple, des logiciels évolués. C'est cette double optique, d'étude des modèles et de leur mise en œuvre, qui a guidé nos travaux. De ce point de vue, le prototype de la méthode FFI que nous avons présenté apparaît comme une ébauche d'un logiciel évolué d'aide à la décision. Il s'agit, à notre avis, d'une démarche permettant d'intégrer des modèles avancés de représentation des préférences dans des outils concrets d'aide à la décision.

Annexes

Annexe A

Exemples relatifs au chapitre 1

L'exemple des 6 cavaliers

Voyons l'exemple classique des six cavaliers (exemple 38) [64] qui montre comment deux modélisation différentes d'un même problème peuvent conduire à des résolutions différentes.

Exemple 38 (les six cavaliers) *Considérons un échiquier à neuf cases et six cavaliers formant deux équipes de trois, les ♞ et les ♟ respectivement en a3, b3, c3 et en a1, b1, c1, voir figure (A.1) :*

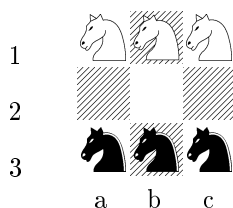


FIG. A.1 – Les six cavaliers

Le but du jeu est de faire passer les cavaliers ♞ des cases a3, b3, c3 aux cases a1, b1, c1 et les cavaliers ♟ des cases a1, b1, c1 aux cases a3, b3, c3 sans que ceux-ci n'occupent la même case et en respectant le déplacement classique des cavaliers. La modélisation de ce problème par la figure (A.1) est la plus naturelle, elle est par ailleurs immédiate. Néanmoins la recherche de la solution, si elle est facile, n'est pas totalement immédiate. En effet, il existe une autre modélisation qui permet d'obtenir immédiatement le résultat. Représentons le problème à l'aide d'un graphe dont les sommets sont les cases de l'échiquier et les arêtes les déplacements possibles des cavaliers. On a la figure (A.2). Les seuls déplacements des pièces autorisés sont donc des déplacements le long des arêtes vers un sommet sans cavalier. La solution est alors obtenue en déplaçant successivement chaque cavalier vers un sommet libre et adjacent jusqu'à ce que chacun soit sur la case adéquate.

Intransitivité de l'indifférence

Nous présentons deux exemples mettant en évidence la non-transitivité de l'indifférence.

Exemple 39 (Armstrong et les sandwiches au fromage) *Armstrong [2] prend l'exemple de sand-*

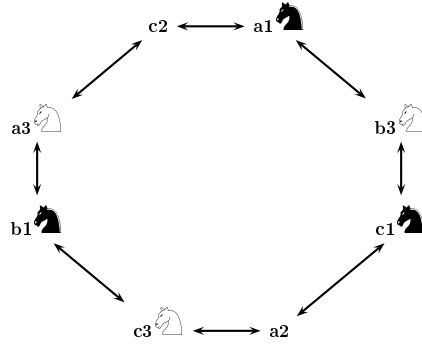


FIG. A.2 – Graphe de déplacement des six cavaliers

wichs au fromage : il considère une série de sandwiches pain/fromage en proportions différentes. Remplaçons progressivement du fromage par du pain de sorte qu'un amateur de sandwich soit indifférent entre deux sandwiches à chaque substitution. Cependant le sandwich initial sera préféré au sandwich final uniquement constitué de pain.

Exemple 40 (Luce et la tasse de café sucrée) Luce [66] décrit ces phénomènes d'intransitivité à l'aide de la structure de quasi ordre à travers l'exemple des sucres dans la tasse de café. Supposons qu'un individu préfère une tasse de café avec 1 sucre plutôt qu'avec 5. On prépare ensuite 401 tasses de café avec $(1 + i/100)$ morceaux de sucres avec $i = 0, \dots, 400$. Le buveur de café est évidemment indifférent entre deux tasses i et $i + 1$ cependant il va préférer la tasse $i = 0$ à la tasse $i = 400$. La non transitivité est due à l'utilisation de seuils de discrimination en deçà desquels l'être humain n'est pas capable de faire de différence ; une différence de quelques milligrammes de sucre n'est en effet pas perceptible par le buveur de café. L'utilisation de tels seuils induit une structure de quasi-ordre sur les objets.

Tversky [109] aborde la propriété de non transitivité

Exemple 41 (Tversky, l'URSS et le Canada) Tversky [109] parle de relation de ressemblance et donne l'exemple des pays étant dans la mouvance de l'ancienne URSS. Si l'on admet assez facilement que les pays de l'ex-bloc de l'est ressemblent à l'URSS, il ne vient pas naturellement de dire que l'URSS ressemble à Cuba. On peut trouver deux raisons à cette non symétrie de la relation. Tout d'abord, on peut considérer que la relation est dotée d'un sens. En effet, c'est l'URSS qui a influencé l'évolution politique des états du bloc de l'est et non le contraire : la relation va donc être orientée. Par ailleurs, la signification que l'on donne au verbe ressembler dans les deux assertions Cuba ressemble à l'URSS et l'URSS ressemble à Cuba n'est pas nécessairement la même, la première va faire référence au régime politique, alors que la seconde ne sait pas trop à quoi se référer. L'exemple est encore plus frappant si l'on considère les assertions Cuba ressemble à l'URSS et le Canada ressemble à l'URSS. Dans la première, il est toujours fait référence au régime politique et, dans la seconde, il est fait référence à l'aspect géographique, au relief et au climat.

Annexe B

Exemples d'affectation des employés

logiciel	développement	mobilité	vente	affectation
16	15	9	13	0,50
13	11	9	9	0,17
12	13	10	10	0,33
16	14	12	12	0,50
19	12	10	12	0,06
14	13	7	11	0,33
13	19	10	13	0,06
14	13	10	12	0,50
16	12	13	9	0,22
11	17	12	10	0,08
15	14	10	11	1,00
19	18	9	12	0,06
16	12	11	10	0,50
13	14	9	12	0,50
17	15	7	10	0,33
12	15	7	13	0,07
15	16	12	9	0,75
13	17	8	9	0,25
14	16	9	13	0,50
13	15	7	10	0,33
12	16	11	8	0,33
14	18	9	9	0,50
14	15	8	9	0,75
15	16	9	10	1,00
15	19	9	13	0,11
13	14	12	11	0,50
13	16	10	11	0,75
13	14	7	10	0,33
15	14	11	7	0,50
14	16	12	8	0,50
14	12	9	12	0,33
16	17	7	12	0,17
16	13	8	9	0,50
15	16	12	10	0,75
18	14	12	11	0,33
14	13	13	10	0,33
12	15	10	9	0,50
13	14	10	12	0,50
12	18	10	8	0,11
19	13	8	9	0,08

TAB. B.1 – *Performances des ingénieurs*

logiciel	développement	mobilité	vente	affectation
9	13	16	15	0,50
9	9	13	11	0,33
10	10	12	13	0,50
12	12	16	14	0,22
10	12	19	12	0,75
7	11	14	13	0,08
10	13	13	19	0,50
10	12	14	13	0,33
13	9	16	12	0,75
12	10	11	17	0,33
10	11	15	14	0,75
9	12	19	18	0,50
11	10	16	12	0,75
9	12	13	14	0,50
7	10	17	15	0,25
7	13	12	15	0,17
12	9	15	16	0,22
8	9	13	17	0,50
9	13	14	16	0,11
7	10	13	15	0,75
9	12	13	15	0,50
12	9	16	12	0,33
7	13	19	12	0,00
10	12	14	13	0,50
9	13	19	18	0,04
7	10	16	12	0,22
7	10	13	11	0,06
12	12	15	16	0,50
8	9	11	17	0,08
7	11	14	16	0,50
10	10	15	14	1,00
10	13	12	15	0,22
12	10	13	17	0,25
9	12	14	13	0,50
10	11	16	15	1,00
9	9	13	19	0,17
13	9	13	14	0,33
10	12	17	15	0,50
9	13	12	13	0,11
11	10	16	14	1,00

TAB. B.2 – *Performances des commerciaux*

logiciel	développement	mobilité	vente	affectation
16	14	12	12	0,22
13	15	14	16	0,75
18	13	15	17	0,17
12	14	15	16	0,50
12	16	15	17	0,33
15	13	19	15	0,17
13	17	13	14	0,25
19	16	16	17	0,17
14	14	15	14	1,00
16	12	13	14	0,33
13	14	17	18	0,17
15	15	16	14	1,00
16	18	12	14	0,22
14	12	15	13	0,33
15	13	18	14	0,33
16	17	13	18	0,17
17	18	14	13	0,17
14	17	16	13	0,50
15	16	14	13	0,75
18	14	14	15	0,50
15	17	15	14	0,75
13	13	17	14	0,25
18	16	12	19	0,04
15	11	17	18	0,06
12	14	13	12	0,11
13	16	14	15	0,75
12	15	16	16	0,50
16	15	14	14	1,00
14	14	13	16	0,75
14	19	12	13	0,06
13	16	15	15	0,75
19	13	14	16	0,17
14	13	19	16	0,17
14	12	13	13	0,17
13	19	18	16	0,06
12	16	12	18	0,07
15	13	11	14	0,17
15	12	15	14	0,50
14	14	16	17	0,75
16	13	15	13	0,50

TAB. B.3 – *Performances des polyvalents*

ingénieurs		commerciaux		inge.-com.	
init.	perturb.	init.	perturb.	init.	perturb.
0,50	0,60	0,50	0,60	0,22	0,30
0,17	0,10	0,33	0,30	0,75	0,80
0,33	0,40	0,50	0,40	0,17	0,30
0,50	0,40	0,22	0,30	0,50	0,40
0,06	0,20	0,75	0,90	0,33	0,40
0,33	0,30	0,08	0,20	0,17	0,10
0,06	0,10	0,50	0,60	0,25	0,30
0,50	0,60	0,33	0,40	0,17	0,10
0,22	0,30	0,75	0,70	1,00	0,90
0,08	0,20	0,33	0,30	0,33	0,20
1,00	0,90	0,75	0,80	0,17	0,30
0,06	0,20	0,50	0,40	1,00	0,90
0,50	0,40	0,75	0,60	0,22	0,30
0,50	0,60	0,50	0,60	0,33	0,40
0,33	0,50	0,25	0,30	0,33	0,30
0,07	0,20	0,17	0,30	0,17	0,30
0,75	0,60	0,22	0,30	0,17	0,10
0,25	0,30	0,50	0,70	0,50	0,40
0,50	0,40	0,11	0,20	0,75	0,60
0,33	0,40	0,75	0,90	0,50	0,70
0,33	0,40	0,50	0,40	0,75	0,70
0,50	0,40	0,33	0,40	0,25	0,10
0,75	0,60	0,00	0,10	0,04	0,20
1,00	0,90	0,50	0,60	0,06	0,10
0,11	0,20	0,04	0,10	0,11	0,20
0,50	0,40	0,22	0,30	0,75	0,60
0,75	0,80	0,06	0,20	0,50	0,40
0,33	0,40	0,50	0,40	1,00	0,90
0,50	0,60	0,08	0,20	0,75	0,80
0,50	0,40	0,50	0,40	0,06	0,20
0,33	0,20	1,00	0,90	0,75	0,90
0,17	0,10	0,22	0,10	0,17	0,30
0,50	0,60	0,25	0,30	0,17	0,10
0,75	0,60	0,50	0,60	0,17	0,30
0,33	0,20	1,00	0,90	0,06	0,20
0,33	0,30	0,17	0,10	0,07	0,20
0,50	0,40	0,33	0,20	0,17	0,10
0,50	0,60	0,50	0,40	0,50	0,40
0,11	0,20	0,11	0,20	0,75	0,80
0,08	0,20	1,00	0,80	0,50	0,60

TAB. B.4 – *Perturbations apportées aux coefficients d'appartenance des points de Z*

sous-ensembles	ingénieurs		commerciaux		ing.-com.	
	init.	calculé	init.	calculé	init.	calculé
apprent.	0,50	0,50	0,50	0,51	0,22	0,22
	0,17	0,20	0,33	0,37	0,75	0,80
	0,33	0,41	0,50	0,48	0,17	0,16
	0,50	0,39	0,22	0,20	0,50	0,53
	0,06	0,08	0,75	0,73	0,33	0,36
	0,33	0,27	0,08	0,08	0,17	0,16
	0,06	0,08	0,50	0,55	0,25	0,27
	0,50	0,52	0,33	0,33	0,17	0,18
	0,22	0,18	0,75	0,77	1,00	1,00
	0,08	0,08	0,33	0,33	0,33	0,31
	1,00	1,00	0,75	0,72	0,17	0,20
	0,06	0,08	0,50	0,50	1,00	1,00
	0,50	0,51	0,75	0,73	0,22	0,21
	0,50	0,60	0,50	0,52	0,33	0,29
	0,33	0,32	0,25	0,27	0,33	0,31
	0,07	0,07	0,17	0,15	0,17	0,13
	0,75	0,64	0,22	0,22	0,17	0,16
	0,25	0,25	0,50	0,50	0,50	0,43
	0,50	0,50	0,11	0,10	0,75	0,73
	0,33	0,32	0,75	0,77	0,50	0,53
test	0,33	0,40	0,50	0,55	0,75	0,70
	0,50	0,51	0,33	0,30	0,25	0,27
	0,75	0,64	0,00	0,00	0,04	0,04
	1,00	1,00	0,50	0,50	0,06	0,04
	0,11	0,12	0,04	0,04	0,11	0,13
	0,50	0,48	0,22	0,20	0,75	0,80
	0,75	0,84	0,06	0,05	0,50	0,53
	0,33	0,32	0,50	0,49	1,00	1,00
	0,50	0,50	0,08	0,08	0,75	0,77
	0,50	0,39	0,50	0,48	0,06	0,04
	0,33	0,35	1,00	1,00	0,75	0,80
	0,17	0,10	0,22	0,24	0,17	0,17
	0,50	0,40	0,25	0,23	0,17	0,16
	0,75	0,64	0,50	0,50	0,17	0,13
	0,33	0,32	1,00	1,00	0,06	0,06
	0,33	0,27	0,17	0,17	0,07	0,09
	0,50	0,56	0,33	0,33	0,17	0,16
	0,50	0,60	0,50	0,55	0,50	0,47
	0,11	0,16	0,11	0,12	0,75	0,73
	0,08	0,08	1,00	1,00	0,50	0,43

TAB. B.5 – Affectations effectuées à l'aide des poids construits à partir des coefficients dégradés

sous-ensembles	ingénieurs		commerciaux		ing.-com.	
	init.	calculé	init.	calculé	init.	calculé
apprent.	0,50	0,36	0,50	0,49	0,22	0,11
	0,17	0,25	0,33	0,39	0,75	0,75
	0,33	0,50	0,50	0,58	0,17	0,35
	0,50	0,26	0,22	0,18	0,50	0,50
	0,06	0,06	0,75	0,54	0,33	0,35
	0,33	0,38	0,08	0,09	0,17	0,16
	0,06	0,06	0,50	0,59	0,25	0,22
	0,50	0,39	0,33	0,49	0,17	0,18
	0,22	0,25	0,75	0,73	1,00	1,00
	0,08	0,15	0,33	0,18	0,33	0,31
	1,00	1,00	0,75	0,87	0,17	0,00
	0,06	0,06	0,50	0,73	1,00	1,00
	0,50	0,57	0,75	0,54	0,22	0,21
	0,50	0,43	0,50	0,57	0,33	0,52
	0,33	0,40	0,25	0,46	0,33	0,31
	0,07	0,04	0,17	0,14	0,17	0,17
	0,75	0,71	0,22	0,18	0,17	0,35
	0,25	0,46	0,50	0,73	0,50	0,78
	0,50	0,36	0,11	0,06	0,75	0,78
	0,33	0,40	0,75	0,73	0,50	0,50
test	0,33	0,29	0,50	0,59	0,75	1,00
	0,50	0,57	0,33	0,28	0,25	0,22
	0,75	0,71	0,00	0,00	0,04	0,00
	1,00	1,00	0,50	0,27	0,06	0,06
	0,11	0,09	0,04	0,02	0,11	0,00
	0,50	0,61	0,22	0,18	0,75	0,75
	0,75	0,89	0,06	0,06	0,50	0,50
	0,33	0,40	0,50	0,60	1,00	1,00
	0,50	0,36	0,08	0,09	0,75	0,47
	0,50	0,26	0,50	0,58	0,06	0,06
	0,33	0,26	1,00	1,00	0,75	0,75
	0,17	0,07	0,22	0,26	0,17	0,25
	0,50	0,57	0,25	0,27	0,17	0,16
	0,75	0,71	0,50	0,27	0,17	0,17
	0,33	0,40	1,00	1,00	0,06	0,05
	0,33	0,38	0,17	0,13	0,07	0,00
	0,50	0,60	0,33	0,49	0,17	0,16
	0,50	0,43	0,50	0,59	0,50	0,67
	0,11	0,13	0,11	0,06	0,75	0,78
	0,08	0,15	1,00	1,00	0,50	0,78

TAB. B.6 – Affectations effectuées à l'aide des poids construits à partir des coefficients nets

Annexe C

Méthode de nuées dynamiques

Introduction

Il s'agit d'une méthode de *clustering* due à E. Diday [28], [30] et [29]. Cette méthode a introduit à sa sortie une nouvelle problématique en classification de données, ... elle vise à *optimiser un critère qui exprime l'adéquation en une classification des objets et un mode de représentation des classes. Le problème d'optimisation se pose alors en termes de "recherche **simultanée** de la classification et de la représentation des classes de cette classification parmi un ensemble de classifications et de représentations possibles qui optimisent le critère."*¹. Ainsi il s'agit d'une méthode générale capable de considérer plusieurs modes de représentation, voir tableau C.1.

Mode de représentation	Méthode
centres de gravité	centres mobiles
loi de proba.	décomposition de mélanges
distances	distances adaptatives
regression	regression typologique
axe factoriel	analyse factorielle typologique
axe discriminant	analyse discriminante typologique
variables	segmentation typologique
individus	tableaux de distances
courbes	lissage typologique

TAB. C.1 – *Différents modes de représentation dans les nuées dynamiques*

1. dans [29] pages 67 et 68

Principe de la méthode

Structure et espace de représentation

La méthode des nuées dynamiques repose sur les notions de structure et d'espace de représentation.

Définition 27 Soit Ω un ensemble, Ω est muni d'une structure de représentation si on lui associe :

- d'une part un ensemble $\mathbf{L}$ et une application D de $\mathbf{P}(\Omega) \times \mathbf{L}$ dans $\mathbb{R}^+$,
- et d'autre part un ensemble $\mathbf{L}_k$ et une application W de $\mathbf{P}_k \times \mathbf{L}_k$ dans $\mathbb{R}^+$.

Les espaces $\mathbf{L}$ et $\mathbf{L}_k$ sont appelés respectivement espace de représentation d'une classe et espace de représentation d'une partition.

où $\mathbf{P}(\Omega)$ est l'ensemble des parties de Ω et $\mathbf{P}_k$ l'ensemble des partitions $\mathcal{P} = (\mathcal{P}_1, \dots, \mathcal{P}_k)$ en k classes de Ω .

D est une mesure de dissemblance entre une classe et une représentation de classe.

W est une "critère" mesurant l'adéquation entre une partition et une représentation d'une partition.

Fonction de représentation et d'affectation

Parallèlement on définit deux fonctions qui f et g qui vont servir respectivement à affecter le objets et à calculer les représentations des classes. Ainsi :

Fonction de représentation : La fonction de représentation g va servir à trouver une représentation des classes, c'est une fonction de $\mathbf{P}_k$ à valeurs dans $\mathbf{L}_k$. Dans le cas de la construction de prototypes (*cf. chap. 5*), c'est la fonction qui va calculer les prototypes des catégories.

Fonction d'affectation : La fonction d'affectation f va servir à trouver une partition, c'est une fonction de $\mathbf{L}_k$ à valeurs dans $\mathbf{P}_k$. Dans le cas de la construction de prototypes, cette fonction a pour objet d'affecter les points d'apprentissage aux *clusters*.

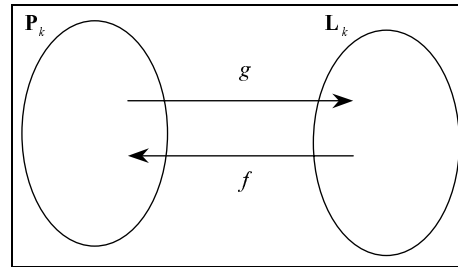


FIG. C.1 – Les fonctions f et g

Les étapes de la méthode

La méthode des nuées dynamiques consiste à trouver un couple $(\mathcal{P}^*, \mathcal{L}^*) \in \mathbf{P}_k \times \mathbf{L}_k$ qui minimise le critère W , ainsi on peut écrire :

$$W(\mathcal{P}^*, \mathcal{L}^*) = \min\{W(\mathcal{P}, \mathcal{L}) \mid \mathcal{P} \in \mathbf{P}_k, \mathcal{L} \in \mathbf{L}_k\} \quad (\text{C.1})$$

Afin de poursuivre cet objectif, la méthode va comme suit :

1. choisir un espace de représentation $\mathbf{L}_k$,
2. définir un critère $W : \mathbf{P}_k \times \mathbf{L}_k \rightarrow \mathbb{R}^+$, ce critère doit permettre de mesurer l'adéquation entre toute partition $\mathcal{P} \in \mathbf{P}_k$ et toute représentation $\mathcal{L} \in \mathbf{L}_k$ de cette partition,
3. définir un problème d'optimisation qui consiste à chercher simultanément la partition $\mathcal{P} \in \mathbf{P}_k$ et une représentation $\mathcal{L}$ de cette partition de façon que $\mathcal{P}$ et $\mathcal{L}$ aient la meilleure adéquation au sens du critère W ,
4. construire un algorithme qui consiste à utiliser de manière itérative les fonctions g et f (affectation et représentation). L'algorithme est initialisé à l'aide d'une classification $\mathcal{P}^0 \in \mathbf{P}_k$ et d'une représentation $\mathcal{L}^0 \in \mathbf{L}_k$ estimées ou tirées au hasard.

Rappel des Notations

Aide multicritère à la décision

$A = \{a, b, c, \dots\}$ ensemble de actions

$E = E_1 \times \dots \times E_n$ espace des critères

g_j fonction critère définie sur A à valeurs dans E_j

$F = \{g_1, \dots, g_n\}$ ou $F = \{1, \dots, n\}$ famille cohérente de critères

$g_j(a) = x_j$ performance de l'action a sur le critère j

$g(a) = x = (g_1(a), \dots, g_n(a)) = (x_1, \dots, x_n)$ image de a dans l'espace des critères

$x_I y_J$ point de E construit à partir des performances du point x sur les critères de I et du points y sur les critères de J avec $F = I \cup J$

q_j seuil d'indifférence du critère j

p_j seuil de préférence du critère j

v_j^+ seuil de veto supérieur du critère j

v_j^- seuil de veto inférieur du critère j

C_j^H concordance avec la relation H sur le critère j

D_j^H discordance avec la relation H sur le critère j

C^H concordance globale avec la relation H

D^H discordance globale la relation H

S_j relation de surclassement sur le critère j

S relation de surclassement globale

P_j ou $\succ_j$ relation de préférence sur le critère j

P ou $\succ$ relation de préférence globale

Q ou $\succsim$ relation de préférence faible

I_j ou $\sim_j$ relation d'indifférence sur le critère j

I ou $\sim$ relation d'indifférence globale

Δ relation de dominance

R relation d'incomparabilité
 $P.\alpha$ problématique du choix
 $P.\beta$ problématique du tri
 $P.\gamma$ problématique du rangement
 $P.\delta$ problématique de la description

Apprentissage et classification

$Z = \{z_1, z_2, \dots\}$ ensemble des points d'apprentissage
 $\{C_1, \dots, C_q\}$ ensemble des catégories
 $Y = Y_1 \cup \dots \cup Y_q$ ensemble des points de références
 $Y_t = Y_t^+ \cup Y_t^-$ ensemble des points de référence centraux de la catégorie C_t
 Y_t^+ ensemble des points de référence centraux positifs de la catégorie C_t
 Y_t^- ensemble des points de référence excluants/négatifs de la catégorie C_t
 $Y_t = Y_t^{sup} \cup Y_t^{inf}$ ensemble des points de référence limites de la catégorie C_t
 Y_t^{sup} ensemble des points de référence limites supérieurs de la catégorie C_t
 Y_t^{inf} ensemble des points de référence limites inférieurs de la catégorie C_t
 $\mu_t(a)$ degré d'affectation de l'action a à la catégorie C_t
 $\mu_t^o(y)$ degré d'affectation du point de référence y à C_t

Construction de paramètres

s solution/individu
 $QR(C_t/s)$ qualité de représentation de C_t par s
 $NR(Y_t)$ non redondance des points de référence
 χ opérateur de croisement
 ϖ opérateur de mutation
CS contraintes structurelles
CE contraintes spécialisées

Table des figures

1.1	<i>Les différentes problématique en A.D.</i>	16
2.1	<i>Arbre de décision de la procédure trichotomique de segmentation</i>	30
2.2	<i>Fonction de "goodness" et de "badness" dans nTOMIC</i>	31
2.3	<i>Partition de l'espace (d, D)</i>	32
2.4	<i>Procédures d'affectation de Electre Tri</i>	32
2.5	<i>Affectation optimiste et pessimiste de Electre Tri</i>	33
2.6	<i>Frontières construites par les méthodes de discrimination linéaire, logistique et quadratique</i>	36
2.7	<i>Modèle à deux couches</i>	37
2.8	<i>Perceptron Multi Couches à une couche cachée</i>	38
2.9	<i>Graphe de diagnostic de cancer</i>	39
2.10	<i>Affectation d'un point par les 3 plus proches voisins et par la fenêtre de Parzen de rayon h</i>	42
2.11	<i>Arbre de décision pour la classification d'objets volants</i>	44
3.1	<i>Préférence faible stricte et graduelle</i>	52
3.2	<i>Valeurs de l'indifférence sur le critère vitesse</i>	53
3.3	<i>Appartenance aux grands et aux géants</i>	54
3.4	<i>Coefficient d'appartenance aux coalitions concordante et discordante</i>	57
3.5	<i>Préordre de préférence sur les catégories de cuisson</i>	64
3.6	<i>Filtrage par indifférence avec catégories monopprofil et multiprofil</i>	70
3.7	<i>Filtrage par préférence et combiné</i>	71
3.8	<i>Algorithme des k plus indifférents prototypes</i>	72
3.9	<i>Arbre de calcul</i>	75
4.1	<i>Construction interactive des poids</i>	85
4.2	<i>Structure d'un RN pour l'apprentissage de coefficients d'importance</i>	90

5.1	<i>Algorithme des nuées dynamiques</i>	95
5.2	<i>Calcul de clusters par nuées dynamiques</i>	96
5.3	<i>Qualité de Représentation et non-redondance</i>	99
5.4	<i>Points Z_4 et de Y_4</i>	99
5.5	<i>Procédure générale de l'algorithme génétique</i>	107
5.6	<i>Procédure de croisement</i>	108
5.7	<i>Procédure de mutation</i>	108
5.8	<i>Fonction d'évaluation</i>	110
6.1		116
6.2		116
6.3		117
6.4		118
6.5		119
6.6		119
6.7		120
6.8		120
6.9		121
6.10		121
6.11		122
6.12		123
6.13		123
A.1	<i>Les six cavaliers</i>	129
A.2	<i>Graphe de déplacement des six cavaliers</i>	130
C.1	<i>Les fonctions f et g</i>	139

Liste des tableaux

1.1	<i>Évaluations des actions a, b et c</i>	12
1.2	<i>Performances</i>	13
2.1	<i>Table de décision pour le diagnostic de la grippe</i>	43
3.1	<i>Tableau comparatif des méthodes de tri multicritère</i>	50
3.2	<i>Discordance et indifférence en fonction de k</i>	59
4.1	<i>Profils des catégories</i>	87
4.2	<i>Poids construits avec les coefficients dégradés de l'ensemble d'apprentissage</i>	88
4.3	<i>Différence entre affectations initiales et calculées avec les poids construits</i>	89
4.4	<i>Poids construits avec les coefficients nets de l'ensemble d'apprentissage</i>	89
4.5	<i>Différence entre affectations initiales et calculées avec les poids construits à partir des coefficients nets</i>	89
5.1	<i>Prototypes de C_4 construits par nuées dynamiques</i>	100
5.2	<i>Affectation avec Z_t et Y_t</i>	100
5.3	<i>Prototypes de C_1</i>	109
5.4	<i>Fonction d'évaluation</i>	109
B.1	<i>Performances des ingénieurs</i>	132
B.2	<i>Performances des commerciaux</i>	133
B.3	<i>Performances des polyvalents</i>	134
B.4	<i>Perturbations apportées aux coefficients d'appartenance des points de Z</i>	135
B.5	<i>Affectations effectuées à l'aide des poids construits à partir des coefficients dégradés</i>	136
B.6	<i>Affectations effectuées à l'aide des poids construits à partir des coefficients nets</i>	137
C.1	<i>Différents modes de représentation dans les nuées dynamiques</i>	138

Bibliographie

- [1] A.K. AGRAWALA, éditeur. *Machine Recognition of Patterns*. IEEE, New York, 1977.
- [2] E.W. ARMSTRONG. « The determinateness of the utility function ». *Economic Journal*, (49):453–467, 1939.
- [3] C. ARONDEL et P. GIRARDIN. « Sorting cropping systems on the basis of their impact on groundwater quality ». *European Journal of Operational Research*, (à paraître).
- [4] K.J. ARROW. « Rational choice function and ordering ». *Econometrica*, (26):121–127, 1959.
- [5] K.J. ARROW. *Social Choice and Individual Values*. John Wiley & Sons, 1963.
- [6] F. BADRAN, S. THIRIA, et F. FOGELMAN-SOULIE. Étude du comportement des réseaux multicouches - comparaison avec l'analyse discriminante. Dans Y. KADRATOFF E. DIDAY, éditeur, *Induction symbolique et numérique à partir de données*, pages 259–282. Cépadués-Édition, 1991.
- [7] N. BELACEL. « Méthodes de Classification Multicritère : Méthodologie et Application au Diagnostic Médical ». Thèse de doctorat, Université Libre de Bruxelles, 1999.
- [8] M. BÉREAU et B. DUBUISSON. « A Fuzzy extended k -nearest neighbor rule ». *Fuzzy Sets and Systems*, (44):17–32, 1991.
- [9] J.C. BEZDEK et P.F. CASTELAZ. « Prototype Classification and Features Selection with Fuzzy Sets ». *IEEE, Trans. Syst., Man, Cybern.*, 7(2):87–92, 1977.
- [10] J.C. BEZDEK, S.K. CHUAH, et D. LEEP. « Generalized k -nearest neighbor rules ». *Fuzzy Sets and Systems*, 3(18):237–256, 1986.
- [11] J.C. BEZDEK, R. EHRLICH, et W. FULL. « FCM: the fuzzy c-means clustering algorithm ». *Computers & Geosciences*, 10(2):191–203, 1984.
- [12] J. C. BORDA. *Mémoire sur les élections au scrutin*. Mémoire de l'Académie des Sciences, 1791.
- [13] B. BOUCHON-MEUNIER, C. MARSALA, et M. RAMDANI. « Inductive Learning and Fuzziness ». *Scientia Iranica*, 2(4), 1996.
- [14] B. BOUCHON-MEUNIER, M. RIFQI, et S. BOTHOREL. « Toward general measures of comparison of objects ». *Fuzzy Sets and Systems*, (84):143–153, 1996.
- [15] D. BOUYSSOU. « Modelling inaccurate determination, uncertainty, imprecision using multiple criteria ». Dans *Proceedings of the VIIIth International Conference on MCDM*. A.G. Lockett, 1991.
- [16] J.S. BREEZE, D. HECKERMAN, et K. ROMMELSE. « Decision-Theoric Troubleshooting ». *Communications of the ACM*, 38(3):49–57, 1995.

- [17] M. CAYROL, H. FARRENY, et H. PRADE. « Fuzzy Pattern Matching ». *Kybernetes*, 11:103–116, 1982.
- [18] CLUB CRIN LOGIQUE FLOUE. *Évaluation Subjective, Méthodes Applications et enjeux*. Ecrin, Paris, 1997.
- [19] CNIL. « Les initiatives de la CNIL ». <http://www.cnil.fr/>. section : Droits et Obligations - Aide à la Décision.
- [20] G.F. COOPER. « *NESTOR : A Computer Based Medical Diagnostic Aid that Integrates Causal and Probabilistic Knowledge* ». thèse de doctorat, Stanford California University, E.U., 1984.
- [21] J.J. CORNAERT et X. DAFPE. « Guide d'achat 2000 ». *Le Moniteur Automobile*, (hors-série), 1999.
- [22] T.M. COVER et P.E. HART. « Nearest neighbor pattern classification ». *IEEE Trans. Info. Theory*, (13):21–27, 1967.
- [23] S.L. CROWFORD et R.M. FUNG. « Constructor : A system for the Induction of Probabilistic Models ». Dans *Proceedings Eight National Conference on Artificial Intelligence*, volume 1, pages 762–769, Boston, 1990. American Association for Artificial Intelligence.
- [24] S.P. CURRAM et J. MINGERS. « Neural Networks, Decision Tree Induction and Discriminant Analysis : a Empirical Comparision ». *J. Opl. Res. Soc.*, 45(4):440–450, 1994.
- [25] B.V. DASARATHY, éditeur. *Nearest Neighbor (NN) Norms: NN Pattern Classification Techniques*. IEEE Computer Society Press, Los Alamitos, CA, 1991.
- [26] M. J. A. N. C. Marquis de CONDORCET. *Essai sur l'application de l'analyse à la probabilité des décisions rendues à la pluralité des voix*. Imprimerie Royale, 1785.
- [27] K.A. DE JONG et W.M. SPEARS. « Learning Concept Classification Rules Using Genetic Algorithms ». Dans *IJCAI*, pages 651–656, 1993.
- [28] E. DIDAY. « Une nouvelle méthode en classification automatique et reconnaissance des formes ». *Rev. de stat. appl.*, 2(19), 1971.
- [29] E. DIDAY, G. CELEUX, G. GOVAERT, Y. LECHEVALIER, et H. RALAMBONDRAIN. *Classification automatique de données, environnement statistique et informatique*. Dunod Informatique, Paris, 1989.
- [30] E. DIDAY, J. LEMAIRE, J. POUGET, et F. TESTU. « *Éléments d'analyse de données* », Chapitre 2, pages 116–135. Dunod, 1982.
- [31] D. DUBOIS. « *Modèles Mathématiques de l'imprécis et de l'incertain en vue d'applications aux techniques d'aide à la décision* ». Thèse d'état, Université Scientifique et médicale de grenoble, Institut Nationale polytechnique de Grenoble, 1983.
- [32] D. DUBOIS et H. PRADE. « Unfair coins and necessity measures: towards a possibilistic interpretation of histograms ». *Fuzzy Sets and Systems*, 10:15–20, 1983.
- [33] D. FAHLOUL, G. TRYSTRAM, et A. DUQUÉNOY. « Mesures, modélisation et interprétation de la coloration des biscuits pendant la cuisson ». Dans *Actes des journées AGORAL*, pages 329–334, 1994.
- [34] P.C. FISHBURN. « *Relations binaires* », Chapitre 3, pages 23–32. Mouton Gauthier-Villars, Paris-La Haye, 1973.

- [35] P.C. FISHBURN. « A survey of multiattribute/multicriterion evaluation theory ». Dans *Conference on Multiple Criteria Problem Solving - Theory, Methodology and Solving* -, Buffalo, 1977.
- [36] R.A. FISHER. « The use of multiples measurements in taxonomic problems ». *Anal. of eugenics*, (7):179–188, 1936.
- [37] E. FIX et J.L. HODGES. « Discriminatory analysis, non parametric discrimination: Consistency properties ». rapport technique no. 4, US Air Force School of Aviation Medicine, Randolph Field, Texas, 1951. [publié dans Agrawala 1977 [1], Silverman et Jones 1989 [104], Dasarathy 1991 [25]].
- [38] E. FIX et J.L. HODGES. « Discriminatory Analysis: Small Sample Performance ». Report no. 11, US Air Force School of Aviation Medicine, Randolph Field, Texas, 1952.
- [39] L. FOULLOY, E. BENOIT, L. KHODJA, et T. TALOU. « Fuzzy techniques for coffee flavour classification ». Dans *IPMU 96*, pages 709–714, 1996.
- [40] G.W. GATES. « The Reduced Nearest Neighbor Rule ». *IEEE Transactions on Information Theory*, pages 431–433, 1972.
- [41] R. GERRITSEN et E. BRAND. « Exclusive Ore Inc. ». <http://www.xore.com/>.
- [42] D. GOLDBERG. « Real-coded genetic algorithms, virtual alphabets, and blocking ». *Complex systems*, 2(5):139–168, 1991.
- [43] D.E. GOLDBERG. *Genetic Algorithms in Search, Optimization & Machine Learning*. Addison-Wesley, Massachusetts, 1989.
- [44] M. GRABISCH, H.T. NGUYEN, et E.A. WALKER. Subjective Multicriteria Evaluation. Dans *Fundamentals of uncertainty calculi with applications to fuzzy inference*, Series C: Theory and decision library, Chapitre 8, pages 213–260. Kluwer Academic Publishers, 1995.
- [45] J.W. GRZYMALA-BUSSE. « *Managing uncertainty in expert systems* », pages 162–172. Engineering and Computer Science. Kluwer Academic Publishers, Dordrecht, 1991.
- [46] D. HACKERMAN et M.P. WELLMAN. « Bayesian Networks ». *Communications of the ACM*, 38(3):27–30, 1995.
- [47] D.J. HAND. *Discrimination and Classification*. John Wiley, 1981.
- [48] P. E. HART. « The Condensed Nearest Neighbor Rule ». *IEEE Transactions on Information Theory (corresp.)*, pages 515–516, 1968.
- [49] A. HATCHUEL et B. WEIL. *L'Expert et le Système, suivi de quatre histoires de systèmes experts*. Economica, 1992.
- [50] F. HERRERA. « Combination Fuzzy Logic-Genetic Algorithms Bibliography ». <http://decsai.ugr.es/herrera/fl-ga.html>.
- [51] J.H. HOLLAND. *Adaptation in natural and artificial systems*. Ann Arbor: The University of Michigan Press, 1975.
- [52] R.A. HOWARD et J.E. MATHESON. Influence diagrams. Dans *Readings on the Principles and Applications of Decision Analysis*, volume 2. R.A. Howard and J.E. Matheson Eds., Menlo Park, California, strategic decision groups édition, 1981.
- [53] ICS UCI. « UCI Machine Learning Repository Content Summary ». <http://www.ics.uci.edu/mlearn/MLSummary.html>.

- [54] E. JACQUET-LAGRÈZE. « PREFCALC: évaluation et décision multicritère ». *Revue de l'utilisateur de IBM-PC*, (3):38–55, 1984.
- [55] M. KABRISKY, M.E. OXLEY, S.K. ROGERS, D.W. RUCK, et B.W. SUTER. « The Multilayer Perceptron as an Approximation to a Bayes Optimal Discriminant Function ». *IEEE Trans. on Neural Networks*, 1(4):296–298, 1990.
- [56] A. KAUFMANN. *Introduction à la théorie des sous-ensembles flous (tome 1)*. Masson, Paris, 1973.
- [57] J. M. KELLER, D. SUBHANGKASEN, et K. UNKLESBAY. « An approximate reasoning technique for recognition in color images of beef steak ». *Int. J. General Systems*, 16:331–342, 1990.
- [58] J.M. KELLER, M.R. GRAY, et J.A. jr GIVENS. « A Fuzzy K -Nearest Neighbor Algorithm ». *IEEE Trans. Syst., Man, Cybern.*, 15(4):580–585, 1985.
- [59] J.M. KELLER et R. KRISHNAPURAM. « The Possibilistic C -Means Algorithm: Insights and Recommendation ». *IEEE Trans. on Fuzzy Syst.*, 4(3):385–393, 1996.
- [60] J.S. KELLY. « Social Choice Bibliography ». *Social Choice and Welfare*, (8):97–169, 1991.
- [61] J. KITTLER. « A method for determining k -nearest neighbours ». *Kybernetes*, 7:313–315, 1978.
- [62] J. LANG et P. PERNY. « Décision et IA ». *Bulletin de l'AFIA*, (31):21–40, 1997.
- [63] O.I. LARICHEV et H.M. MOSHKOVICH. « An Approach to Ordinal Classification Problems ». *Int. Trans. Opl. Res.*, 1(3):375–385, 1994.
- [64] J.L. LAURIÈRE. *Résolution de problèmes par l'homme et la machine*. Éditions Eyrolles, Paris, 1985.
- [65] F. LECAMP. « Algorithmes génétiques pour l'apprentissage de poids dans les méthodes de tri multicritère: conception et expérimentation ». mémoire de dea, LAMSADE, Université Paris IX Dauphine, 1997.
- [66] R.D. LUCE. « Semiorders and theory of utility discrimination ». *Econometrica*, (24):178–191, 1956.
- [67] R. MASSAGLIA et A. OSTANELLO. N-tomic: A Support System for Multicriteria Segmentation Problems. Dans *Multiple Criteria Decision Support*, volume 356 de *Lecture Notes in Economics and Mathematical Systems*, pages 167–174. P. Korhonen A. Lewandowski J. Wallenius (Eds.), springer-verlag édition, 1989.
- [68] W.S. MCCULLOCH et W. PITTS. « A logical calculus of the ideas immanent in various activity ». *Bulletin of Math. Biophysics*, (5):115–133, 1943.
- [69] D. MICHIE, D.J. SPIEGELHALTER, et C.C. TAYLOR. *Machine learning, neural and statistical classification*. Artificial Intelligence. Ellis Horwood, 1994.
- [70] J. MOSCAROLA. « Aide à la décision en présence de critères multiples fondée sur une procédure trichotomique, Méthode et Applications ». thèse de doctorat, Université Paris IX Dauphine, France, 1977.
- [71] J. MOSCAROLA et B. ROY. « Procédure automatique d'examen de dossiers fondés sur une segmentation trichotomique en présence des critères multiples ». *R.A.I.R.O. Recherche opérationnelle/Operations Research*, 11(2):145–173, 1977.
- [72] M. NARANJO, M. RICHETIN, et G. RIVES. « Algorithme rapide pour la détermination des k plus proches voisins ». *R.A.I.R.O. Informatique/Computer Science*, 14(4):369–378, 1980.

- [73] S.K. PAL et D.D. MAJUMDER. « Fuzzy Sets and Decisionmaking in Vowel and Speaker Recognition ». *IEEE Transactions on Systems, Man, and Cybernetics*, SMC-7(8):625–629, 1977.
- [74] Yoh-Han PAO. *Adaptive Pattern Recognition and Neural Networks*. Addison-Wesley Publishing Compagny, 1989.
- [75] E. PARZEN. « On estimation of a probability density function and mode ». *Ann. Math. Stat.*, (33):1065–1076, 1962.
- [76] Z. PAWLAK. « Rough Sets ». *Int. J. of Computer and Information sciences*, 11(5), 1982.
- [77] Z. PAWLAK. *Rough sets theoretical aspects of reasoning about data*. Theory and Decision Library. Kluwer Academic Publishers, 1991.
- [78] J. PEARL. *Probabilistic Reasoning in Intelligent Systems : Networks of Plausible Inference*. Morgan Kaufmann, San Mateo, California, 1988.
- [79] P. PERNY. « Méthode, agrégation et exploitation de préférences floues dans une problématique de rangement ». thèse de doctorat, Université Paris IX Dauphine, France, 1992.
- [80] P. PERNY. « Multicriteria Filtering Methods Based on Concordance/Non-Discordance principles ». *Annals of Operational Research*, 80:137–167, 1998. Special issue on Preference Modelling.
- [81] P. PERNY. « Multicriteria Filtering Methods Based on Concordance/Non-Discordance principles ». *Annals of Operations Research*, 80:137–167, 1998. Special issue on Preference Modelling.
- [82] P. PERNY et B. ROY. « The use of fuzzy outranking relations in preference modeling ». *Fuzzy Sets and Systems*, (49):33–53, 1992.
- [83] P. PERNY et A. TSOUKIÀS. « On the continuous extension of a four valued logic for preference modelling ». Dans *IPMU'98*, pages 302–309, Paris, 1998.
- [84] P. PERNY et J.D. ZUCKER. « Collaborative Filtering Methods based on Fuzzy Preference Relations ». Dans *EUROFUSE-SIC'99*, pages 279–285, Budapest, 1999.
- [85] N. PERROT. « Application de la logique floue à la conduite des fours de cuisson en biscuiterie ». Mémoire de dea en génie des procédés, option agro-alimentaire, ENSIA, 1994.
- [86] H. POINCARÉ. *La science et l'hypothèse*. Flammarion, 1902.
- [87] J.C. POMEROL. « Artificial intelligence and human decision making ». *European Journal of Operational Research*, (99):3–25, 1997.
- [88] J.C. POMEROL et S. BARBA-ROMERO. *Choix multicritère dans l'entreprise, principe et pratique*. Hemès, 1993.
- [89] J.R. QUINLAN. « Induction of decision trees ». *Machine Learning*, (1):81–106, 1986.
- [90] J.R. QUINLAN. *C4.5 Programs for machine learning*. Morgan Kaufmann Publishers, San Mateo, California, 1993.
- [91] H. RAÏFFA. *Decision Analysis*. Addison Wesley, Massachussets, USA, 1970.
- [92] M. RIFQI. « Mesures de comparaison, typicalité et classification d'objets flous : théorie et pratique ». thèse de doctorat, Université Paris 6 Pierre et Marie Curie, 1996.
- [93] B. ROY. « Classement et choix en présence de points de vue multiples (la méthode Electre) ». *Cahiers du Centre d'Études de Recherche Opérationnelle*, (8):57–75, 1968.

- [94] B. ROY. « *Algèbre moderne et théorie des graphes orientées vers les sciences économiques et sociales* », Chapitre 10. Dunod, 1969-70.
- [95] B. ROY. « Electre III: un algorithme de classements fondé sur une représentation floue des préférences en présence de critères multiples ». *Cahiers du CERO*, 20(1), 1978.
- [96] B. ROY. Acceptance, rejection, delay for additional information presentation of a decision aid approach. Dans Alperovich / de DOMBAL / GREMY, éditeur, *Evaluation of efficiency of medical action*, pages 73–82. North Holland Publishing Company, 1979.
- [97] B. ROY. *Méthodologie multicritère d'aide à la décision*. Economica, Paris, 1985.
- [98] B. ROY. Main sources of inaccurate determination, uncertainty and imprecision in decision models. Dans B. Munier M. SHAKUN, éditeur, *Compromise, Negotiation and Group Decision*, pages 43–62. D. Reidel Publishing Company, 1988.
- [99] B. ROY. « Science de la décision ou science de l'aide à la décision? ». *Revue internationale de systémique*, 6(5):497–529, 1992.
- [100] B. ROY et D. BOUYSSOU. *Aide Multicritère à la Décision*. Economica, Paris, 1993.
- [101] D.E. RUMELHART, G.E. HINTON, et R.J. WILLIAMS. Learning internal representations by error propagation. Dans *Parallel distributed processing*, volume 1, Chapitre 8, pages 318–362. MIT Press, Cambridge, 1986.
- [102] G. SAPORTA. *Probabilités, analyse de données et statistiques*. Technip, Paris, 1990.
- [103] S. SEN et L. KNIGHT. « A genetic prototype learner ». Dans *IJCAI*, pages 725–731, Montreal, 1995.
- [104] B. W. SILVERMAN et M. C. JONES. « E. Fix and J. L. Hodges (1951): An important contribution to nonparametric discriminant analysis and density estimation ». *International Statistical Review*, 57:233–247, 1989.
- [105] D.B. SKALAK. « Using a Genetic Algorithm to Learn Prototypes for Case Retrieval and Classification ». Dans *AAAI-93 Case-Based Reasoning Workshop*, pages 64–69, 1993.
- [106] R. SŁOWIŃSKI, éditeur. « *Intelligent Decision Support, Handbook of Applications and Advances of the Rough Sets Theory* ». D: System Theory, Knowledge Engineering and Problem Solving. Kluwer Academic Publishers, 1992.
- [107] R. SŁOWIŃSKI, éditeur. « *Fuzzy sets in decision analysis, operations research and statistics* », Chapitre 1-5, pages 1–176. The handbook of fuzzy sets series. Kluwer Academic Publishers, 1998.
- [108] A. TVERSKY. « Intransitivity of preferences ». *Psychological Review*, 1(76):31–48, 1969.
- [109] A. TVERSKY. « Features of Similarity ». *Psychological Review*, 4(84):327–352, 1977.
- [110] D. VANDERPOOTEN. « *L'approche interactive dans l'aide multicritère à la décision* ». Thèse de doctorat, Université Paris IX-Dauphine, 1990.
- [111] J.C. VANSNICK. « De Borda et Condorcet à l'agrégation multicritère ». Cahier du Lamsade 70, LAMSADE, Université Paris IX - Dauphine, 1986.
- [112] P. VINCKE. *L'aide multicritère à la décision*. Éditions de l'Université de Bruxelles, Bruxelles, 1989.
- [113] YU WEI. « *Aide multicritère à la décision dans le cadre de la problématique du tri* ». thèse de doctorat, Université Paris IX Dauphine, France, 1992.

- [114] S.M. WEISS et C.A. KULIKOWSKI. *Computer Systems that Learn, Classification and Prediction Methods from Statistics, Neural Nets, Machine Learning and Expert Systems*. Morgan Kaufmann Publishers, Inc., San Mateo, California, 1991.
- [115] WRIGHT. « Correlation and Causation ». *J. Agric. Res.*, (20):557–585, 1921.
- [116] R.R. YAGER. « On ordered weighted averaging aggregation operators in multicriteria decision making ». *IEEE Trans. on Systems, Man and Cybernetics*, (18):183–190, 1988.
- [117] R.R. YAGER. « Families of OWA operators ». *Fuzzy Sets and Systems*, (59):125–148, 1993.
- [118] M.S. YANG et C.T. CHEN. « On strong consistency of the fuzzy generalized nearest neighbor rule ». *Fuzzy Sets and Systems*, (60):273–281, 1993.
- [119] L.A. ZADEH. « Fuzzy Sets ». *Information and Control*, (8):338–353, 1965.

Vu : le Président
M.....

Vu les suffragants
MM.....

Vu et permis d'imprimer :

le vice-Président du Conseil Scientifique chargé de la Recherche de l'Université PARIS IX DAUPHINE.

RÉSUMÉ

Cette thèse traite des problèmes et des méthodes de classification dans le cadre de l'aide multicritère à la décision. Un problème de classification consiste en l'affectation d'objets à des catégories prédéfinies (nous parlerons aussi de problème d'affectation). Par ailleurs, le cadre de l'aide à la décision dans lequel nous nous plaçons nous amène à tenir compte des préférences d'un décideur. L'approche que nous proposons se fonde sur l'utilisation d'une relation d'indifférence comme reflet des préférences du décideur. Le principe d'affectation consiste à évaluer l'appartenance d'un objet à une catégorie à partir de l'évaluation de l'indifférence l'objet à classer et des objets typiques de la catégorie (points de référence).

Nous présentons tout d'abord une définition générale des méthodes de classification multicritère. Puis nous proposons deux méthodes d'affectation (PIP et k -PIP) qui vont principalement différer par la nature des points de référence à considérer (prototypes ou points exemples). Ces méthodes possèdent l'avantage de procurer au décideur, non seulement un résultat, mais aussi explication de celui-ci : c'est là un des aspects fondamentaux de l'aide à la décision.

Nos méthodes considèrent des catégories décrites par des points de référence. Donner les points de référence de ces catégories n'est cependant pas toujours une tâche aisée pour un décideur. C'est pourquoi nous proposons deux procédures qui permettent de construire les prototypes des catégories à partir d'ensembles de points d'apprentissage. D'autre part, et de même que pour les points de référence, un décideur peut éprouver certaines difficultés à fixer certains paramètres préférentiels (seuils, poids). Nous proposons à cet effet une procédure permettant de construire à partir de points exemples les coefficients d'importance relative des différents critères (poids) entrant en compte dans le processus d'aide à la décision relatif au problème de classification.

Mots-clés : *aide à la décision, multicritère, classification, procédure d'affectation, construction de paramètres.*

Evaluation and Multicriteria Classification Systems for Decision Aid: model's and assignment procedures learning

ABSTRACT

This thesis deals with classification problems and methods for MultiCriteria Decision Aid (MCDA). Classification -or assignment- problems concern objects' assignment to predefined categories. For dealing with the MCDA concepts, Decision-Maker's (DM) preferences have to be taken into consideration. Our approach is based on an indifference relation considered as an indication of DM preferences. The classification principle lies in estimating the membership of an object to categories by valuing the indifference between this object and typical objects of the categories called reference points.

This thesis is organized as follows. First, we present a general definition of Multicriteria Classification Methods. Then, we propose a couple of new multicriteria classification methods: the Most Indifferent Prototype and the k -More Indifferent Prototypes, respectively PIP and k -PIP in French. Their reference points (prototypes or examples) mainly distinguish these methods. Beyond their assignment function, the advantage is to provide an understandable explanation of the assignment, which is a crucial aspect of decision aid.

Our methods use categories that are characterized by reference points, but finding their determination may not be easy for the DM. This is the reason why we propose two procedures to construct categories' reference points on the basis of a set of learning points. As for the reference points, it could not be easy for the DM to state preferential parameters such as thresholds and weights. With this end of view we propose a procedure for aiding the DM to set the coefficient that measure the relative of criteria (weights), these coefficients being a natural issue of our classification DA process.

Keywords: *decision aid, multicriteria, classification, assignment procedures, parameters construction.*