



Extraction automatique de connaissances pour la décision multicritère

Michel Plantié

► To cite this version:

Michel Plantié. Extraction automatique de connaissances pour la décision multicritère. Interface homme-machine [cs.HC]. Ecole Nationale Supérieure des Mines de Saint-Etienne, 2006. Français. NNT : 2006EMSE0018 . tel-00353770

HAL Id: tel-00353770

<https://theses.hal.science/tel-00353770>

Submitted on 16 Jan 2009

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.

N° d'ordre : 411 I

THESE
présentée par

Michel Plantié

Pour obtenir le grade de Docteur
de l'Ecole Nationale Supérieure des Mines de Saint-Etienne
et de l'Université Jean Monnet de Saint-Etienne

Spécialité : Informatique

*Extraction automatique de connaissances
pour la décision multicritère*

Soutenue à Nîmes..... le 29 septembre 2006

Membres du jury

Président :
M. Ali KHENCHAF

Professeur / ENSIETA Brest

Rapporteurs :
Mme Camille ROSENTHAL-SABROUX
M. Joël QUINQUETON

Professeur / Université Paris-Dauphine
Professeur / Université Paul Valéry, Montpellier III

Examineurs :
Mme Danièle HERIN
M. Elie SANCHEZ

Professeur / Université Montpellier II
MCU-PH HDR / Université de la Méditerranée, Marseille

Directeur de thèse :
M. Jacky Montmain

Professeur / Ecole des Mines d'Ales

Invité :
M. David PEARSON

Professeur / Université Jean Monnet, Saint Etienne

● **Spécialités doctorales :**
SCIENCES ET GENIE DES MATERIAUX

MECANIQUE ET INGENIERIE
GENIE DES PROCEDES
SCIENCES DE LA TERRE
SCIENCES ET GENIE DE L'ENVIRONNEMENT

MATHEMATIQUES APPLIQUEES
INFORMATIQUE
IMAGE, VISION, SIGNAL
GENIE INDUSTRIEL
MICROELECTRONIQUE

Responsables :
J. DRIVER Directeur de recherche – Centre SMS
A. VAUTRIN Professeur – Centre SMS
G. THOMAS Professeur – Centre SPIN
B. GUY Maître de recherche
J. BOURGOIS Professeur – Centre SITE

E. TOUBOUL Ingénieur
O. BOISSIER Professeur – Centre G2I
JC. PINOLI Professeur – Centre CIS
P. BURLAT Professeur – Centre G2I
Ph. COLLOT Professeur – Centre CMP

● **Enseignants-chercheurs et chercheurs autorisés à diriger des thèses de doctorat** (titulaires d'un doctorat d'Etat ou d'une HDR)

BENABEN	Patrick	PR 2	Sciences & Génie des Matériaux	SMS
BERNACHE-ASSOLANT	Didier	PR 1	Génie des Procédés	CIS
BIGOT	Jean-Pierre	MR	Génie des Procédés	SPIN
BILAL	Essaïd	MR	Sciences de la Terre	SPIN
BOISSIER	Olivier	PR 2	Informatique	G2I
BOUDAREL	Marie-Reine	MA	Sciences de l'inform. & com.	DF
BOURGOIS	Jacques	PR 1	Sciences & Génie de l'Environnement	SITE
BRODHAG	Christian	MR	Sciences & Génie de l'Environnement	SITE
BURLAT	Patrick	PR 2	Génie industriel	G2I
COLLOT	Philippe	PR 1	Microélectronique	CMP
COURNIL	Michel	PR 1	Génie des Procédés	SPIN
DAUZERE-PERES	Stéphane	PR 1	Génie industriel	CMP
DARRIEULAT	Michel	ICM	Sciences & Génie des Matériaux	SMS
DECHOMETS	Roland	PR 2	Sciences & Génie de l'Environnement	SITE
DELAFOSSSE	David	PR 2	Sciences & Génie des Matériaux	SMS
DOLGUI	Alexandre	PR 1	Informatique	G2I
DRAPIER	Sylvain	PR 2	Mécanique & Ingénierie	CIS
DRIVER	Julian	DR	Sciences & Génie des Matériaux	SMS
FOREST	Bernard	PR 1	Sciences & Génie des Matériaux	SMS
FORMISYN	Pascal	PR 1	Sciences & Génie de l'Environnement	SITE
FORTUNIER	Roland	PR 1	Sciences & Génie des Matériaux	CMP
FRACZKIEWICZ	Anna	MR	Sciences & Génie des Matériaux	SMS
GARCIA	Daniel	CR	Génie des Procédés	SPIN
GIRARDOT	Jean-Jacques	MR	Informatique	G2I
GOEURIOT	Dominique	MR	Sciences & Génie des Matériaux	SMS
GOEURIOT	Patrice	MR	Sciences & Génie des Matériaux	SMS
GRAILLOT	Didier	DR	Sciences & Génie de l'Environnement	SITE
GROSSEAU	Philippe	MR	Génie des Procédés	SPIN
GRUY	Frédéric	MR	Génie des Procédés	SPIN
GUILHOT	Bernard	DR	Génie des Procédés	CIS
GUY	Bernard	MR	Sciences de la Terre	SPIN
GUYONNET	René	DR	Génie des Procédés	SPIN
HERRI	Jean-Michel	PR 2	Génie des Procédés	SPIN
JOYE	Marc	Ing. (Gemplus)	Microélectronique	CMP
KLÖCKER	Helmut	CR	Sciences & Génie des Matériaux	SMS
LAFOREST	Valérie	CR	Sciences & Génie de l'Environnement	SITE
LE COZE	Jean	PR 1	Sciences & Génie des Matériaux	SMS
LI	Jean-Michel	EC (CCI MP)	Microélectronique	CMP
LONDICHE	Henry	MR	Sciences & Génie de l'Environnement	SITE
MOLIMARD	Jérôme	MA	Sciences & Génie des Matériaux	SMS
MONTHEILLET	Frank	DR 1 CNRS	Sciences & Génie des Matériaux	SMS
PERIER-CAMBY	Laurent	MA1	Génie des Procédés	SPIN
PIJOLAT	Christophe	PR 1	Génie des Procédés	SPIN
PIJOLAT	Michèle	PR 1	Génie des Procédés	SPIN
PINOLI	Jean-Charles	PR 1	Image, Vision, Signal	CIS
SOUSTELLE	Michel	PR 1	Génie des Procédés	SPIN
STOLARZ	Jacques	CR	Sciences & Génie des Matériaux	SMS
THOMAS	Gérard	PR 1	Génie des Procédés	SPIN
TRAN MINH	Cahn	MR	Génie des Procédés	SPIN
VALDIVIESO	Françoise	CR	Génie des Procédés	SPIN
VAUTRIN	Alain	PR 1	Mécanique & Ingénierie	SMS
VIRICELLE	Jean-Paul	CR	Génie des procédés	SPIN
WOLSKI	Krzysztof	CR	Sciences & Génie des Matériaux	SMS
XIE	Xiaolan	PR 1	Génie industriel	CIS

Glossaire :

PR 1 Professeur 1^{ère} catégorie
PR 2 Professeur 2^{ème} catégorie
MA(MDC)Maître assistant
l'Environnement
DR 1 Directeur de recherche
Ing. Ingénieur
MR(DR2)Maître de recherche
CR Chargé de recherche
EC Enseignant-chercheur
ICM Ingénieur en chef des mines

Centres :

SMS Sciences des Matériaux et des Structures
SPIN Sciences des Processus Industriels et Naturels
SITE Sciences Information et Technologies pour
G2I Génie Industriel et Informatique
CMP Centre de Microélectronique de Provence
CIS Centre Ingénierie et Santé

Remerciements...

Je tiens à remercier ici les personnes exceptionnelles qui m'ont aidé dans ce travail de recherche. J'ai reçu pendant ce parcours : aide, encouragements et inspiration de bien des personnes qui y ont tenu une place importante. Il m'est impossible de les citer toutes nommément, toutefois certaines contributions furent telles que je tiens à leur rendre un hommage tout particulier.

- Jacky Montmain, professeur à l'Ecole des Mines d'Alès, qui a encadré mes travaux de thèse et envers qui je tiens à exprimer ma très profonde gratitude. Il a guidé, orienté, critiqué, amélioré, et surtout soutenu ces travaux. Grâce à sa grande compétence, son professionnalisme, il m'a éclairé sur de nombreux domaines de la recherche scientifique.*
- Gérard Dray, m'a souvent conseillé, et apporté sa parfaite maîtrise de multiples disciplines scientifiques ainsi que ses nombreuses marques d'estime et de soutien. Ce fut un plaisir de participer à ses côtés au défi DEFT 2005.*
- Pascal Poncelet, professeur et directeur adjoint du laboratoire LGI2P pour ses conseils toujours précieux.*
- Alexandre Meimouni, qui m'a apporté une aide notable et une grande maîtrise pour le développement du prototype, ainsi que de nombreuses remarques scientifiques pertinentes pendant nos discussions animées qui ont fait progresser ce travail de recherche.*

Je remercie tous les membres du jury d'avoir bien voulu examiner mes travaux de thèse et d'avoir participé au jury de cette thèse :

- Monsieur Ali Khenchaf,*
- Madame Camille Rosenthal-Sabroux, Monsieur Joël Quinqueton,*
- Madame Danièle Herin, Monsieur Elie Sanchez,*
- Monsieur David Pearson.*

Je remercie également :

- Sylvie Ranwez pour son aide et sa courtoisie,*
- Stefan Janaki, Christelle Urtado, Thomas Lambolais, Michel Vasquez, pour leurs contributions.*
- Rachid Arezki, Abdenour Mokrane, et Fabien Jalabert doctorants du laboratoire, pour leurs conseils et les discussions pleines de créativité que nous avons eues ensemble.*
- l'Ecole des Mines d'Alès et ses dirigeants qui m'ont permis de réaliser ce travail de recherche.*
- Sylvie Cruvellier pour sa patience, son aide remarquable et pour son vaste savoir-faire dans les démarches administratives.*
- Les membres du laboratoire LGI2P, enseignants, chercheurs ou étudiants, que j'ai sollicités à maintes reprises. Ils ont contribué à ce que mon travail se déroule dans une atmosphère agréable.*

Je remercie infiniment ma femme Christine qui m'a soutenu et encouragé pendant ces années.

Table des matières

Introduction générale.....	15
Chapitre I. Conception d'un SIADG et fiabilité humaine.....	19
1. Cognition et décision.....	20
1.1. Le processus de la décision	20
1.2. Le kaléidoscope de la décision et l'apport des sciences humaines	21
1.3. Les pièges de la raison	28
2. Fiabilité humaine et aide à la décision	32
2.1. L'erreur humaine.....	33
2.2. Instanciation à l'activité de décision en organisation.....	40
Chapitre II. Le rôle de l'information dans le contrôle d'un processus décisionnel.....	45
1. Introduction	46
2. Le modèle S.T.I de H.A. Simon.....	47
2.1. Rationalité substantive versus rationalité procédurale	47
2.2. Acquisition et traitement de l'information	48
2.3. Synthèse	50
3. L'interprétation cybernétique du modèle S.T.I.....	51
3.1. Boucles cognitives et dynamique du processus décisionnel	51
3.2. Boucle de rétroaction et contrôle de la dynamique du processus décisionnel	52
4. Conclusion.....	68
Chapitre III. Gestion des connaissances et décision	71
1. Introduction	72
2. La course à l'information à l'ère du WWW.....	72
2.1. La croissance exponentielle de l'information	72
2.2. Un besoin avéré ou suscité ?	73
2.3. La surcharge cognitive	74
2.4. Indexer pour mieux gérer	74
3. La notion de connaissance.....	75
3.1. Nature de la connaissance	75
3.2. Une typologie des connaissances pour une organisation	77
4. La gestion des connaissances	78
5. Le processus d'indexation : de l'information à la connaissance.....	79
5.1. Indexation manuelle	80
5.2. Indexation automatique	81
5.3. Comparaison des deux approches	81
6. Connaissance actionnable	81
7. Instanciation de ces notions à notre problématique	82
Chapitre IV. Représentations de corpus documentaires et techniques de classification...	85
1. La représentation de corpus documentaires	86
1.1. Le mot comme unité linguistique.....	87
1.2. La phrase comme unité linguistique.....	87
1.3. La technique des « n-grammes ».....	87
1.4. Traitement préliminaire.....	88
1.5. Différents types de représentations textuelles.....	89
1.6. Conclusion.....	92
2. Les techniques de classification	92
2.1. Le représentant moyen d'une classe.....	94
2.2. Les classifieurs	94
2.3. Classifieur de Bayes	95

2.4. Autres classifieurs en traitement automatiques du langage	97
2.5. Extension de la sémantique distributionnelle.....	97
2.6. Techniques de segmentation automatique de texte	97
2.7. Validation	98
2.8. Mesure de performance	100
3. Conclusion.....	102
Chapitre V. Extraction automatisée de CAs dans un processus d'évaluation multicritère	103
1. Description générale de la chaîne de traitement.....	104
2. Corpus documentaire et base d'apprentissage	106
3. Représentation vectorielle du corpus de documents	107
3.1. Index et représentation vectorielle	108
3.2. Constitution de l'ensemble d'apprentissage de chaque classifieur	109
3.3. Réduction de l'index dédié à un classifieur	111
3.4. Ajout des synonymes	112
4. La classification des textes	113
4.1. le classifieur filtrage des textes évaluatifs et narratifs.....	113
4.2. Le classifieur « classification critère »	114
4.3. Attribution d'un score au commentaire d'une CA	115
4.4. Score d'une CA	117
5. Quelques justifications	117
6. Conclusion.....	119
Chapitre VI. Application : un outil d'aide à la décision pour un gestionnaire de vidéo-club	121
1. Gestionnaire d'un vidéo-club	122
1.1. Liste des films candidats	124
1.2. Critères d'évaluation	125
2. Collecte et évaluation des critiques	125
2.1. Constitution de la base de connaissance en vue de l'apprentissage.....	125
2.2. Calcul de l'index global	125
2.3. Écrans de saisie d'une nouvelle critique	125
3. Agrégation multicritère, justification et identification	132
3.1. Présentation générale de l'IHM.....	133
3.2. Utilisation des grilles multicritères	134
3.3. La grille Film/Date	135
3.4. La grille Fim/Critère.....	139
3.5. La grille Critère/Date	140
3.6. Report de nouvelles critiques sur la grille.....	141
4. Evaluation du risque décisionnel.....	142
4.1. L'IHM de la fonction « risque ».....	142
4.2. Calcul du risque.....	143
4.3. Signal de contrôle.....	144
4.4. Calcul du nombre de critiques minimum pour modifier le classement.....	145
5. Argumentation.....	146
5.1. Argumentation en absolu	146
5.2. Argumentation en relatif	149
6. Scénario d'évolution dans le temps et stratégies de contrôle.....	152
6.1. Scénario CONFIRMER.....	153
6.2. Scénario EProuver.....	154
6.3. Scénario LIBRE ou BOUCLE OUVERTE.....	156
7. Conclusion.....	157
Conclusion générale	159

Bibliographie	163
Annexes	173
1. Calcul d'un modèle de classification par les arbres de décision	173
1.1. Modèles de classification et arbres décisionnels	173
1.2. Pouvoir de décision et quantité d'information	174
1.3. Autre mesure du meilleur candidat	175
2. Calcul d'un modèle de classification par la méthode des SVM.....	175

Liste des figures

Figure I-1: Stratégie politique ou choix du consommateur, la décision semble partout présente dans l'Histoire et dans les petites histoires de la vie de tous les jours	21
Figure I-2: Décision et stratégie de concurrence.....	21
Figure I-3 : Approche organisationnelle de la décision	22
Figure I-4 : Approche politique de la décision.....	22
Figure I-5 : Approche psychologique de la décision	22
Figure I-6 : Le décideur n'est pas monolithique	23
Figure I-7 : Le dirigeant d'une entreprise a une connaissance limitée de la situation	24
Figure I-8 : Limites de la pensée stratégique	26
Figure I-9: Schéma de principe du modèle GEMS	34
Figure I-10: Erreurs au niveau SB	35
Figure I-11: Erreurs au niveau RB.....	36
Figure I-12: Erreurs au niveau KB.....	38
Figure I-13: L'aide à la décision collective sous l'angle de l'erreur humaine.....	44
Figure I-14 : Conception du SIADG et fiabilité humaine.....	44
Figure II.1: boucles cognitives dans le modèle IDCR de Simon	52
Figure II.2: Un exemple de grilles d'évaluation par les CAs.....	54
Figure II.3: Le processus de <i>vote</i>	54
Figure II.4: Suivi de l'évaluation globale des solutions envisagées dans le temps.....	55
Figure II.5: Découpage en étiquettes symboliques des contributions partielles	62
Figure II.6: Contrôle de l'information décisionnelle ou pilotage par le risque.....	67
Figure III.1 : Les processus de transformation des connaissances de l'entreprise.....	77
Figure III.2: Connaissance Actionnable (CA).....	82
Figure IV-1 : Processus de validation par le test.....	99
Figure IV-2 : Processus de validation croisée.....	99
Figure IV-3 : Processus de Bootstrap.....	100
Figure V-1: Processus général.....	105
Figure V-2 : Détail du processus général	105
Figure V-3 : Structure de la base de connaissance pour trois critères.....	107
Figure V-4 : Processus de calcul des index.....	108
Figure V-5 : Traitements liés à l'apprentissage	108
Figure V-6 : Exemple de la structure de table de correspondance entre lemmes de l'index et les synonymes.	113
Figure V-7 : Diagramme d'activités générique du calcul d'un classifieur	113
Figure V-8 : Diagramme d'activités du classifieur « filtrage ».....	114
Figure V-9 : Diagramme d'activités du classifieur « classification critère » (cas de 3 critères)	114
Figure V-10 : Evaluation hiérarchique des phrases	116
Figure V-11 : Diagramme d'activités du classifieur « classification P/N ».....	116
Figure V-12 : Diagramme d'activités du classifieur « classification N/TN ».....	117
Figure VI.1 : Le processus de décision multicritère du gestionnaire de vidéo-club	123
Figure VI.2 Ecran de saisie d'une nouvelle critique	126
Figure VI.3 : Grille multicritère Film/Critère	133
Figure VI.4 : Notes associées aux couleurs de la grille	134
Figure VI.5 : Les grilles multi-critères et l'explication.....	135
Figure VI.6 : Grille film/date	135
Figure VI.7 : Une vue de l'ACP sur l'espace initial des évaluations partielles.	138

Figure VI.8 : La méthode du « clustering hiérarchique » sur la base test des clients du vidéo-club.	139
Figure VI.9 : La grille Film/critère au 15 mai 2004.....	140
Figure VI.10 : La grille Critère/date.....	140
Figure VI.11 : Affichage des CA pour une case	141
Figure VI.12 : Affichage d'une CA suite à une demande d'explication.....	141
Figure VI.13 : La grille Film/Date au 29 mai 2004	141
Figure VI.14 : Ecran « risque » au 15 mai 2004	142
Figure VI.15 : Ecran risque après adjonction de critiques positives sur <i>Le pacte des loups</i> .	143
Figure VI.16 : exemple de grille risque affichée.....	144
Figure VI.17 : Ecran d'une explication en absolu.....	147
Figure VI.18 : Explication en absolu et paramètre pourcentage d'explication.....	148
Figure VI.19 : Ecran d'explication et paramètre tolérance	149
Figure VI.20 : écran d'explication en relatif.	150
Figure VI.21 : écran d'explication en relatif lorsque $k\%$ est faible.....	151
Figure VI.22 : Explication en relatif avec une tolérance plus grande	152
Figure VI.23 : Grille Film/Date arrêtée au 15 mai 2004.....	152
Figure VI.24 : Risque décisionnel associé à chaque film candidat.....	153
Figure VI.25 : Situation décisionnelle au 22 mai 2004 –Scénario CONFIRMER	154
Figure VI.26 : Situation décisionnelle au 22 mai 2004 –Scénario EPROUVER.....	155
Figure VI.27 : Situation décisionnelle au 29 mai 2004 –Scénario EPROUVER.....	155
Figure VI.28 : Situation décisionnelle au 22 mai 2004 –Scénario LIBRE ou BO.....	156
Figure VI.29 : Situation décisionnelle au 29 mai 2004 –Scénario LIBRE ou BO.....	157
Figure VI.30 : exemple d'arbre de décision	173

Liste des tableaux

Tableau I-1 : Synthèse des approches de la décision	23
Tableau I-2 : La théorie de Bayes est-elle utilisée dans nos choix de tous les jours ?.....	29
Tableau I-3 : Classifications d'éléments bibliographiques de la décision vue par les SHS et notre typologie des erreurs humaines	41
Tableau II-1 : Exemple de classement et notion d'effort à produire.....	64
Tableau IV-1 : les quatre possibilités d'un classifieur	100

Introduction générale

Décision et incertitude épistémique

De nos jours l'explosion des technologies de l'information a fait de l'Internet un outil, incontournable tant au niveau personnel que professionnel occupant une place prépondérante dans notre société et notre vie. Le WWW nous offre un monde de l'information prodigieux sans limite, ce qui pourrait apparaître comme une véritable opportunité, mais d'un autre côté, c'est l'ouverture à un monde inexplorable et incontrôlable en termes de quantité et de fiabilité des informations délivrées, où l'ensemble des choix possibles est quasi infini... Le problème pour le décideur—au sens commun du terme—en quête de l'« information pertinente » n'est plus l'accès à cette information qui tend à être uniformément distribuée, mais de savoir la trouver dans d'imposants flux informationnels extrêmement « bruités ». S'aventurer dans le monde de l'information sur des sujets, où ses connaissances sont limitées, est un pari risqué. Leurres et illusions. Pourquoi prendre pour argent comptant, une information scientifique qui est passée à travers tout processus de vérification et de validation et cela, qui plus est, en toute légitimité ? Comment choisir un téléphone mobile ? Un hôtel ? Tout est sur la toile, à chacun de savoir consulter et de se faire son opinion.

La quête de l'information pour prendre une décision est classiquement un processus itératif. Si l'on reste sur l'exemple du web, la plupart du temps, l'évaluation, la comparaison et le classement des alternatives qui s'offrent à nous va reposer sur une première consultation de quelques sites de recommandation (CIAO, Leguide, etc.) ou de comparateurs (Kelkoo, etc.), mais cette recherche ne saurait être exhaustive et systématique. Néanmoins, en général, elle permet d'avoir une catégorisation grossière des alternatives. La poursuite de la quête de l'information peut s'affiner et consiste alors à ne s'enquérir que des points particuliers du problème sur lesquels l'état de nos connaissances ne permet pas de définir une situation

décidable de façon fiable. Pour parvenir à cette situation, il nous faut combler un déficit informationnel qui, tant qu'il subsiste, induit une incertitude épistémique, i.e. un risque décisionnel lié à l'état insuffisant de nos connaissances à un instant donné. Par la suite, notre investigation se focalise alors sur des dimensions de l'évaluation des alternatives bien spécifiques : celles pour lesquelles une sélection fiable ne saurait être envisagée tant les valeurs des options sont proches. Plus notre quête du savoir avance, plus l'incertitude épistémique est supposée diminuer et la situation décidable se rapprocher. L'enjeu de ce processus cognitif est de savoir conduire sa recherche d'information avec perspicacité : sur quels aspects du problème sera-t-il le plus judicieux d'acquérir des compléments d'information pour gérer au mieux l'entropie de la décision ?

Certes, le terme d'entropie est ici choisi par analogie avec la théorie de l'information, mais on peut lui prêter une définition plus psychologique : elle est la mesure de la richesse en information qu'un message contient pour le décideur (le *récepteur*) ; l'entropie est fonction du rapport des réponses possibles et connues avant et après que l'on a reçu l'information. Elle correspond au degré d'indétermination dans la communication...

Ce qu'il est important de retenir au travers de ces remarques, c'est le caractère dynamique de ce processus cognitif et la nécessité de bien conduire son investigation pour arriver à une conclusion au plus vite. Gérer l'information, c'est contrôler la dynamique décisionnelle. Savoir discerner l'information utile dans le temps, c'est être capable d'accélérer la dynamique de décision. On parlera ainsi de la conduite du processus de décision ou du contrôle de sa dynamique par analogie avec la régulation d'un système technique [Akharraz, 2004].

Avant d'exposer comment l'on souhaite apporter quelques éléments de réponse à cette problématique, nous pensons qu'il est utile de justifier la terminologie automatique de notre discours. On oublie trop souvent que l'émergence de la décision comme domaine d'étude scientifique remonte aux années 1943 et 1948, moment où apparaissent plusieurs courants de recherche parallèle dont la cybernétique comme science de la communication et de la commande dans les systèmes naturels et artificiels créée par N. Wiener en 1948 et vite associée aux problématiques de la recherche opérationnelle.

Pour Waliser [Waliser, 1977], l'étude des modifications que l'on peut apporter à un système débouche sur une typologie des rôles des modèles. On peut être ainsi amené à distinguer quatre fonctions principales d'un modèle :

- une fonction cognitive : le modèle sert à représenter les relations qui existent entre variables d'entrée et variables de sortie du système ;
- une fonction prévisionnelle : le modèle sert à prévoir comment évolueront les variables de sortie du système, en fonction de l'évolution probable des variables externes et d'hypothèses de fixation des variables de commande ;
- une fonction normative : le modèle sert à représenter les relations souhaitables entre variables d'entrée et de sortie du système ;
- une fonction décisionnelle : le modèle sert à déterminer comment fixer les variables de commande pour atteindre les objectifs que l'on s'est fixés sur les variables de sortie, compte tenu de l'évolution probable des variables externes.

Au sens large un modèle cognitif a pour fonction de fournir une représentation plus ou moins conforme d'un système existant, mettant en évidence certaines de ses propriétés et permettant éventuellement d'en déduire d'autres (modèle explicatif). Au sens restreint, un modèle cognitif doit fournir une relation aussi bonne que possible entre les entrées et les sorties du

système et en particulier, préciser l'influence relative des diverses variables d'entrée (modèle descriptif).

Au sens large du terme, un modèle prévisionnel a pour fonction, à partir de la connaissance d'un système dans des situations données, d'inférer son comportement dans des situations non encore observées (modèle de simulation). Au sens restreint, un modèle prévisionnel doit permettre, à partir de la connaissance des variables d'entrée et la relation entrée-sortie du système, d'évaluer la valeur des sorties dans le futur (modèle de prévision).

Au sens large du terme, un modèle normatif a pour fonction de fournir une représentation plus ou moins idéale d'un système à créer, mettant en évidence certaines propriétés souhaitables (modèle prescriptif). Au sens restreint, un modèle normatif doit fournir une épure d'un système reliant en fonction de certaines propriétés des entrées et des sorties (modèle constructif).

Au sens large du terme, un modèle décisionnel a pour fonction de fournir à un décideur des informations lui permettant d'éclairer une décision visant à modifier le système (modèle de décision). Au sens restreint, un modèle décisionnel doit permettre de fournir les valeurs optimales des variables de commande au regard de certains objectifs, compte tenu des variables d'entrée (modèle d'optimisation).

Cette typologie de Waliser confère au modèle décisionnel une définition qui soutient parfaitement notre point de vue : la « vanne de l'information », le système d'information, est l'organe de commande qui permet de contrôler la dynamique d'une décision. D'une phase d'apprentissage, sélective et pertinente, dépend la dynamique pour atteindre une situation décidable.

Informatique décisionnelle et automatisation cognitive

L'étude du processus de la décision doit donc inclure un véritable processus de traitement de l'information en particulier lorsque la situation n'est pas complètement « mathématisable », ce qui est généralement le cas lorsque la décision relève d'une organisation ou d'un quelconque collectif. H.A. Simon, qui a élaboré la théorie économique de la rationalité limitée, présentée en 1947, dans *Administration et processus de décision* [Simon, 1983], précise à cet effet que la difficulté consiste à traiter l'information entre autre parce qu'elle est trop abondante. Pour cela, il met en avant les outils informatiques, qu'il appelle des « prothèses de l'homme » au sens où ils aident ce dernier à poser plus rationnellement les problèmes, filtrer les informations et simuler et planifier l'action qui devra suivre. Il édicte alors quelques principes pour la conception de ces outils et précise que l'essentiel est de comprendre la manière dont les décisions sont prises dans l'organisation, soulever les questions auxquelles l'information va répondre, adopter une approche arborescente et modulaire des problèmes. Dans le cadre de la décision organisationnelle, le vocable « système de traitement de l'information » (S.T.I) permet de désigner commodément la lignée des modèles issus de la pensée de H.A.Simon [Simon, 1997].

Prise au pied de la lettre, cette approche des systèmes d'information comme prothèses de l'homme conduit aujourd'hui à la naissance d'une « informatique décisionnelle », dont le mot d'ordre principal est : "fournir à tout utilisateur reconnu et autorisé, les informations nécessaires à son travail". Ce slogan fait naître une nouvelle informatique, intégrante, orientée vers les utilisateurs et les centres de décision des organisations. Tout utilisateur de l'entreprise ayant à prendre des décisions doit pouvoir accéder en temps réel aux données de l'entreprise, doit pouvoir traiter ces données, extraire l'information pertinente de ces données pour prendre les "meilleures" décisions.

La question que l'on peut alors se poser est : jusqu'où peut-on pousser l'informatisation du processus décisionnel ? Si l'on s'en réfère au contrôle de processus industriels, l'informatisation des contrôles et de certaines actions peut être poussée jusqu'à l'automatisation cognitive, c'est-à-dire le remplacement de l'opérateur pour certaines fonctions de décision [Després, 1991]. Si on reconnaît aujourd'hui de plus en plus clairement que, contrairement à ce qui était prévu, l'automatisation (au sens de la commande) a pu parfois rendre les tâches des opérateurs plus difficiles, la même chose peut être imaginée pour l'automatisation cognitive [Bainbridge, 1991].

Cette thèse, sans prendre parti, aborde ce sujet délicat qu'est l'automatisation cognitive. Elle propose la mise en place d'une chaîne informatique complète pour supporter chacune des étapes de la décision. Elle traite en particulier de l'automatisation de la phase d'apprentissage en faisant de la connaissance actionnable—la connaissance utile à l'action—une entité informatique manipulable par des algorithmes.

Organisation du document

Le plan du manuscrit est le suivant.

Au chapitre I, nous adoptons le point de vue pragmatique du fiabiliste humain, analysant l'erreur humaine tout au long du processus de décision, pour en déduire les aides informatiques que l'on peut légitimement attendre d'un système interactif d'aide à la décision (SIAD). En pratique, nous passons en revue l'apport des Sciences Humaines à l'étude de la décision sous la perspective d'analyse du fiabiliste. Nous procédons alors à une catégorisation des erreurs humaines dans la prise de décision, ce qui nous permet de déduire un ensemble de fonctionnalités informatiques nécessaires à la conception d'un système d'aide à la décision.

Au second chapitre, nous proposons un SIAD. D'abord, il est fait mention des principales caractéristiques du modèle de la décision de H.A. Simon, avant d'en rappeler l'interprétation cybernétique proposée par Akharraz dans sa thèse et sur laquelle s'édifie notre SIAD. Des travaux de ce dernier, nous reprenons les définitions formelles de justification des choix et de risque décisionnel dans un processus de décision multicritère pour les étendre à des opérateurs d'agrégation plus discriminants et génériques. Ce modèle théorique se veut soutenir la faisabilité de notre SIAD qui ouvrirait les portes de l'automatisation cognitive.

A ce stade de la chaîne de traitement de l'information, la phase d'apprentissage pour la décision est « dirigée », mais exige néanmoins l'intervention de l'homme. La suite du manuscrit est consacrée à l'automatisation de cet apprentissage. Le troisième chapitre est un chapitre bibliographique. Il introduit les notions de connaissances, de gestion des connaissances avant de conclure sur les concepts fondamentaux de *connaissance actionnable* et d'*indexation automatique* sur lesquels reposent les modèles et outils développés dans la suite de nos travaux. Le chapitre IV propose une synthèse rapide des techniques d'apprentissage les plus éprouvées en vue de l'extraction automatique de connaissances actionnables pour le processus de décision. Il détaille plus spécifiquement le classifieur de type Naïve Bayes multinomial qui sera l'outil de base retenu dans la suite de nos expérimentations. Le chapitre V décline l'ensemble des notions et techniques introduites dans les chapitres III et IV à la problématique spécifique d'extraction automatique des connaissances dans un processus d'évaluation multicritère. La phase d'apprentissage ainsi informatisée, notre modèle franchit alors un pas de plus dans la voie délicate de l'automatisation cognitive. Le chapitre VI reprend et illustre le processus informatisé dans sa globalité sur un exemple d'application. Il s'agit de doter le gérant d'un vidéoclub d'un outil informatique pour l'aider à constituer et gérer son stock.

Chapitre I.

Conception d'un SIADG et fiabilité humaine

« Heyam Dukham Anagatam »
*Eviter les erreurs (le danger) avant
qu'elles (il) ne surviennent.*
Yoga Sutra, (I, 16)
Patanjali

1. Cognition et décision.....	20
1.1. Le processus de la décision	20
1.2. Le kaléidoscope de la décision et l'apport des sciences humaines	21
1.3. Les pièges de la raison	28
1.3.1. La logique des erreurs	28
1.3.2. Décider dans l'urgence	29
1.3.3. Entre stratégie consciente et force aveugle	30
1.3.4. Décision raisonnée ou raisonnable	31
2. Fiabilité humaine et aide à la décision	32
2.1. L'erreur humaine	33
2.1.1. La typologie de Reason	33
2.1.2. La typologie de Nicolet	38
2.1.3. La typologie de de Keyser.....	39
2.1.4. Caractéristiques de ces typologies.....	40
2.2. Instanciation à l'activité de décision en organisation.....	40

Parce que l'on s'intéresse plus spécifiquement à la décision dans une organisation au sens large du terme, on parlera de décision de groupes. Avant de présenter le modèle de la décision de groupes que nous avons adopté pour notre étude et le Système Interactif d'Aide à la Décision de Groupe (SIADG) qui le supporte, nous proposons quelques éléments de justification quant aux fonctionnalités de ce SIADG. A cet effet, nous introduisons de manière originale la conception d'un SIADG en nous appuyant sur un argumentaire relevant d'une analyse de fiabilité humaine.

Nous pensons d'abord qu'il est utile de regarder l'éclairage que d'autres approches que les méthodes mathématiques proposent sur la décision. Nous en profitons pour brièvement lister les apports fondamentaux des Sciences Humaines et Sociales (SHS) à l'étude de la décision. La décision y est toujours décrite comme un processus éminemment cognitif. Ce constat opéré, cette liaison forte entre cognition et décision exige donc que l'on prête une attention particulière à l'erreur humaine tout au long du processus décisionnel. Quelles sont les erreurs, les biais cognitifs potentiels qui sont susceptibles de venir entacher la décision ? Quels enseignements en tirer pour doter le SIADG des fonctionnalités les plus appropriées pour soutenir l'activité cognitive des acteurs de la décision dans chacune des phases de celle-ci ? Nous nous sommes donc livrés à cette analyse bibliographique de la décision vue par les SHS parce que c'est chez les politologues, les socio économistes, les psychosociologues, etc., que sont répertoriés les biais et pièges qui font de la décision une activité particulièrement sensible à l'erreur humaine.

Nous proposons alors une synthèse sur les typologies de l'erreur humaine, puis nous réinterprétons, fusionnons et déclinons celles-ci relativement à la décision de groupe et aux facteurs d'erreurs que l'analyse bibliographique précédente a mis en lumière. C'est de cette classification que nous déduisons la conception du SIADG, les fonctions attendues d'un SIADG : les apports informationnels et représentationnels nécessaires à l'évitement des biais cognitifs qui menacent le bon déroulement du processus décisionnel dans chacune de ces phases sont mis en évidence et il est alors possible de spécifier une instrumentation pertinente pour accompagner et assister les acteurs de la décision dans ce processus hautement cognitif par définition.

1. Cognition et décision

Pour aborder la décision sous l'angle de la cognition, il est utile de s'attarder déjà sur les différents sens que l'on prête communément au terme de décision : acte ou processus.

Dans une décision finale se mêlent raisonnements et intuitions, calculs et appréciations globales, volonté et doutes, hypothèses et certitudes, expériences et rêveries, débats et compromis. Comment qualifier cette connaissance action dont l'usage est si courant ? Nous proposons dans cette section quelques éléments concourant à éclairer l'énigme de la décision sous l'angle des Sciences Humaines...

1.1. Le processus de la décision

La décision est au cœur d'une des grandes énigmes de la pensée et de l'action. Elle semble partout présente dans l'Histoire et dans les petites histoires de la vie de tous les jours. L'examen de la réalité des décisions montre que décider ne correspond pas à une phase précise, clairement identifiable où tombe le couperet. La langue française emploie l'expression « prendre une décision » un peu comme si la décision était un objet identifiable. Le point d'arrivée est confondu avec le processus. La langue anglaise, qui emploie

l'expression de « decision making process » (la fabrication de la décision) rend mieux compte de ce fait. La décision est un cheminement : elle se construit, se négocie, suit des voies sinueuses au cours du temps. Elle pourrait sembler insaisissable.

Depuis longtemps, les sciences de l'homme tentent de percer les mystères de la décision, d'en mettre en évidence les logiques, d'y déceler des régularités. Les Sciences de l'Homme et de la Société apportent un éclairage à la compréhension de situations de décisions, bien loin du calcul décisionnel de la recherche opérationnelle, de la théorie des jeux ou encore de la cybernétique.

1.2. Le kaléidoscope de la décision et l'apport des sciences humaines

Il n'existe clairement pas de science de la décision. Il existe simplement plusieurs entrées, plusieurs points de vue, plusieurs lectures. Enfin, il est légitime de distinguer quelques grands domaines qui constituent des centres d'intérêt particulier pour des études sur la décision : le domaine politique [Allison, 1971], celui des entreprises et des organisations, celui de la vie quotidienne (Figure I-1).



TABLEAU COMPARATIF CAMÉSCOPES NUMÉRIQUES										
MODELE	Canon MV300i	Canon MV30i	Canon MV3i MC	Canon MV3i	Canon MV3i MC	Canon MV300	Canon MV430 i	Canon MV3i	Canon XM-1	
SORTIE	01/04/00	01/04/00	01/11/00	01/09/00	01/11/00	01/04/00	01/04/00	01/04/01	01/09/00	01/10/99
Prix	8990 FFHTC 5844.48 FFHT	8990 FFHTC 6102.01 FFHT	8290 FFHTC 7767.56 FFHT	8490 FFHTC 7098.66 FFHT	8290 FFHTC 6803.68 FFHT	8290 FFHTC 5259.2 FFHT	8290 FFHTC 4789.3 FFHT	7990 FFHTC 6880.6 FFHT	8290 FFHTC 7767.56 FFHT	9990 FFHTC 13359.66 FFHT
ZOOM	10 x optique	12 x optique	10 x optique	10 x optique	10 x optique	10 x optique	12 x optique	10 x optique	10 x optique	20 x optique
Capteur	CCD 540 000 pixels 1/4" filter	CCD 490 000 Pixels 1/4" progressive scan, FV8 filter	800 000 pixels 1/4" progressive scan	800 000 pixels 1/4" progressive scan	800 000 pixels grand scan, fill, FV8 filter	CCD 540 000 pixels 1/4" filter	CCD 490 000 Pixels 1/4" progressive scan, FV8 filter	0.8cm 540 000 pixels (340 000 effctifs) filter	800 000 pixels 1/4" progressive scan	9 CCD 520 Kpixels
Stockage	cassette mini DV	mini DV	mini DV	mini DV	miniDV, carte MMC 3 Mb	cassette mini DV	cassette mini DV	cassette Mini DV	mini DV	cassette miniDV
Sensibilité	2.5 lux	1.5 lux	5 lux au 1/25 sec	5 lux au 1/25 sec, 2.5lux mode low light	5 lux au 1/25 sec	2.5 lux en vitesse lente	1.5 lux	2 lux (mode haute lumière)	5 lux au 1/25 sec	
Affichage	écran LCD 6.3 cm 112 Kpixels	écran LCD 6.3 cm 200 Kpixels	LCD 6.3 cm couleur 200 Kpixels	LCD 6.5 cm couleur 200 Kpixels	LCD 6.3 cm couleur 200 Kpixels	écran LCD 6.3 cm 112 Kpixels	écran LCD 9 cm 200 Kpixels	écran LCD 9 cm 200 couleur TFT LCD 112 Kpixels	LCD 6.3 cm couleur 200 Kpixels	8.35 cm 122 Kpixels LCD

Figure I-1: Stratégie politique ou choix du consommateur, la décision semble partout présente dans l'Histoire et dans les petites histoires de la vie de tous les jours

Les théories de la décision pourraient être déclinées en cinq grands types d'approches qui privilégient selon le cas : la rationalité de l'action, les contraintes organisationnelles, le poids des compromis, la psychologie du décideur ou encore le contexte de l'action.

Selon le premier modèle rationnel, la décision prise doit être le résultat d'un choix comparatif entre les diverses solutions possibles. Le décideur et ses conseillers doivent mesurer avec soin les risques et les issues probables de chaque formule, peser leurs avantages et inconvénients pour retenir finalement celle qui représente le meilleur rapport « coût/efficacité ». Cette analyse, en termes de calcul rationnel, postule l'existence d'un acteur unique qui agirait en vertu de préférences hiérarchisées en fonction de la meilleure utilité (Figure I-2).



Figure I-2: Décision et stratégie de concurrence

Ce séduisant modèle théorique fait abstraction des aspects organisationnels souvent implicites dans un processus décisionnel. En effet, la panoplie des solutions théoriquement possibles est plus large que celles réellement envisagées par une approche purement rationnelle. Les propositions émanent de scénarios préétablis par des membres bien identifiés de

l'organisation. De plus les informations disponibles sont généralement incomplètes, imprécises ou contradictoires !

L'idée de cette approche organisationnelle est que le décideur ne possède pas en fait une connaissance totale de la situation, d'où le terme de « rationalité limitée » cher à H.A. Simon [Simon, 1997]. Les limitations dans la connaissance des faits et hypothèses proviennent des contraintes de l'organisation qui sélectionne ou favorise tel ou tel scénario en fonction de ses intérêts (Figure I-3).



Figure I-3 : Approche organisationnelle de la décision

Une approche dite politique de la décision consiste à mener une analyse qui privilégie le jeu des acteurs, leur capacité à manœuvrer, à produire des coalitions. Le choix fait l'objet de jeux de pouvoirs entre les parties concernées. Une solution de compromis entre les positions radicales est généralement adoptée (Figure I-4).



Figure I-4 : Approche politique de la décision

Les approches psychologiques s'intéressent selon l'angle de vue adopté, aux processus mentaux mis à l'œuvre lors d'une décision individuelle ou bien à la personnalité des décideurs (Figure I-5). La psychologie cognitive envisage la décision comme un cas particulier des stratégies de résolution de problèmes. Elle explore des rationalités subjectives employées par les individus en situation. Beaucoup d'expériences concernent les biais cognitifs c'est-à-dire les raisonnements en apparence corrects qui conduisent à de fausses conclusions. Heureusement, les heuristiques (c'est-à-dire les méthodes de résolution de problèmes approximatives) employées dans la vie courante sont souvent fiables.

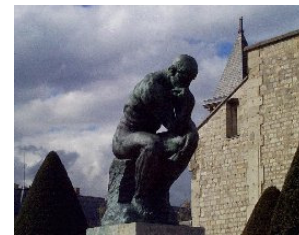


Figure I-5 : Approche psychologique de la décision

L'autre voie explorée par la psychologie de la décision concerne la personnalité des décideurs. Les spécialistes américains de management ont décrit toute une panoplie de profils de personnalité ou styles décisionnels qui correspondent à autant de styles de leadership : les intuitifs, les rationnels, les ingénieux, les rigides, les autoritaires, les innovants, etc. Alors que l'intuitif aime suivre son inspiration, le rationnel s'informe, calcule, soupèse, s'arme de prudence. Le premier est créatif, le second plus conformiste.

On peut encore ranger les études de psychosociologie parmi ces analyses psychologiques. On souligne alors comment les interactions entre participants vont contribuer à polariser la décision soit dans le sens d'une radicalisation, soit dans le sens du conformisme.

Les modèles dits composites désignent une panoplie d'autres approches dont le seul trait commun est de combiner plusieurs aspects des modèles précédents.

Le nombre d'approches de la décision reflète bien sûr la complexité du phénomène mais aussi la diversité des domaines d'application. On ne décide pas de la même façon dans un ministère ou dans une PME, dans une assemblée générale de grévistes ou dans une firme multinationale. Selon le contexte d'action, la dimension politique, psychologique, organisationnelle, le poids de l'environnement,... seront plus ou moins prépondérants [Stratégor, 1988] ([Tableau I-1](#)).

<i>Types d'approches</i>	<i>Dimensions mises en lumière</i>
Approche rationnelle	• définition précise du problème • liste systématique des choix possibles • calcul de la meilleure utilité
Approche organisationnelle	• rationalité limitée du décideur • contraintes de l'organisation • choix d'une solution satisfaisante
Approche politique	• importance des négociations et compromis entre les groupes concernés
Approche psychologique	• stratégie mentale de résolution de problèmes • profil de personnalité des décideurs et styles décisionnels
Approche composite	• intègre plusieurs dimensions des modèles précédents • étude des phases, des niveaux hiérarchiques et des contextes de la décision

Tableau I-1 : Synthèse des approches de la décision

Quelles caractéristiques du processus de décision nous enseignent ces quelques lignes ? D'abord dans le domaine du politique au sens large du terme, c'est la distribution du pouvoir entre les hommes politiques, technocrates et société civile qui constitue la trame des acquis de la science politique.

Le décideur n'est pas monolithique : les actions ou mesures résultantes sont élaborées à travers un processus long et enchevêtré mobilisant de nombreux acteurs ; il n'y a pas de décision, un décideur, mais une série de stratégies et de compromis entre les points de vue, entre des groupes qui ne partagent pas la même solution ([Figure I-6](#)). Dans sa théorie du surcode [Sfez, 1992], L. Sfez montre que derrière l'image trompeuse d'une décision consciente et unifiée, il y a en fait une multiplicité de rationalités différentes qui s'imbriquent, se superposent se confrontent.



Figure I-6 : Le décideur n'est pas monolithique

Le décideur in fine sert de détonateur à l'élaboration de nouvelles politiques au sens large du terme. Même si c'est lui qui signe et tranche en dernier ressort, il s'appuie sur des solutions concrètes dont la formulation détaillée, techniquement constituée, n'est pas son œuvre. Des groupes tiers fournissent ces solutions, élaborent les options et assurent leur légitimité opérationnelle. Les mesures prises en définitive ne sont pas nécessairement des réponses à des exigences générales de la société, de l'entreprise mais le produit d'activités intermédiaires d'experts et de conseillers [Thoenig, 1987].

L'argumentaire est un incontournable exercice de rhétorique du monde politique. A ce propos, Henri Nallet, ancien ministre de l'agriculture puis de la justice, déclarait lors de sa première rencontre avec les représentants de l'institution judiciaire : « Je ne demanderai au premier ministre une « rallonge » budgétaire que si nous sommes capables de justifier tout franc supplémentaire par une amélioration correspondante dans le fonctionnement du service. Si on ne peut pas justifier, je ne demanderai rien. ».

Sur le plan de la psychologie sociale, un premier constat est que lors d'une prise de décision, les groupes peuvent adopter des choix plus risqués que les individus ; le groupe est alors supposé favoriser la dilution des responsabilités, la réduction de l'engagement des individus. On peut penser au contraire, que dans certaines circonstances, les individus s'engagent plus au sein d'un groupe : il a été en effet observé que la décision en groupe provoque un renforcement des options initiales quelles qu'elles soient, on parle de processus de polarisation.

Un autre point de la décision de groupe concerne la nécessité de parvenir à un consensus [Doise *et al.*, 1992] : l'engagement des individus s'en trouve augmenté car, s'il n'y a pas nécessité d'aboutir à une décision commune, les membres d'un groupe débattent de leurs opinions personnelles, mais sans chercher à retravailler cognitivement les opinions exprimées par les autres. L'objectif d'un consensus ne doit pas être de supprimer les conflits, mais de les tolérer. Plus les options individuelles sont variées, plus la confrontation est importante et plus la reformulation du problème peut devenir explicite. La recherche d'un compromis par le groupe incite les individus à reconsidérer leurs positions sur le problème. Lorsque les membres du groupe se sont mis d'accord sur une conclusion, on constate qu'ils la maintiennent ensuite lorsqu'on les interroge individuellement. La possibilité de laisser s'exprimer les opinions divergentes favorise donc la cohésion du groupe. Par ailleurs, il faut noter que l'innovation vient souvent des points de vue minoritaires dissidents : ainsi, pour être véritablement efficaces, les mécanismes de décision doivent favoriser la contradiction.

Pour ce qui est de la décision stratégique d'entreprise, compte tenu des remarques précédentes, on peut affirmer qu'aujourd'hui l'action d'un dirigeant d'entreprise ne correspond pas à l'image mythique du grand patron, clairvoyant et omniprésent [Koenig, 1990; March, 1991; Stratégor, 1988]. Sa connaissance de l'environnement est limitée, il se heurte à une organisation souvent rétive à ses choix. L'élaboration d'une stratégie émergente doit tenir compte de ces contraintes (Figure I-7).



Figure I-7 : Le dirigeant d'une entreprise a une connaissance limitée de la situation

Le dirigeant a bien souvent été représenté comme le grand stratège, à la tête d'un arsenal de prévisions, scénarios, calculs, méthodes, plans, programmes, projets, qu'une organisation disciplinée et efficace élabore sous sa direction, et dont il conduit fermement mais sereinement la réalisation. Représentation mythique ou modèle normatif ? L'important est que se révèle une trame commune, une sorte de modèle de base, implicite, de la décision stratégique, qui comporte trois temps successifs :

- l'anticipation : les décisions procèdent de l'état futur de l'environnement ;
- le choix : le décideur est le dirigeant qui exprime sa volonté, fruit de son analyse ;

- la mise en œuvre : le choix arrêté par le dirigeant est réalisé par l'entreprise conformément à sa volonté.

Malheureusement, de nombreuses études empiriques ont montré que les processus de décision stratégiques dans les entreprises s'écartent sensiblement de ce modèle rationnel et linéaire. Tout d'abord l'environnement décourage les anticipations. Ensuite, l'entreprise est bien moins docile qu'on ne le croit, pour exécuter les décisions du dirigeant. Enfin, le dirigeant lui-même ne ressemble guère à ce décideur aux capacités exceptionnelles.

C'est d'un environnement incertain que naissent et disparaissent les menaces qui peuvent affecter la stratégie de l'entreprise. La stratégie implique donc de mettre l'entreprise en adéquation avec son environnement, et tout d'abord avec le ou les marchés visés. Dans ce domaine, la fiabilité des prévisions est plus que douteuse et certaines entreprises estiment que des prévisions fausses sont plus nuisibles que l'absence de prévisions, car elles produisent de fausses certitudes, qui limitent les capacités d'intelligence et de réaction des managers. L'environnement stratégique d'une entreprise est aussi un univers d'acteurs—concurrents, fournisseurs, distributeurs —parties prenantes de toutes sortes. Les alliances de plus en plus nombreuses que les firmes passent entre elles sont souvent décrites comme des moyens de réduire l'incertitude générée par la multiplicité des acteurs, en stabilisant leurs relations.

Pour ce qui est de l'organisation rétive, les moyens humains et organisationnels qui composent l'entreprise sont censés servir les desseins stratégiques élaborés par le dirigeant. La formulation de la stratégie serait l'apanage du dirigeant, et sa mise en œuvre serait la mission de l'ensemble de l'entreprise. Cette distinction si naturelle occulte le rôle considérable de l'organisation dans la réalisation de la stratégie. L'organisation se montre bien souvent rétive et cela de plusieurs manières.

- Les grandes entreprises sont dotées de structures et de systèmes de gestion qui déterminent largement leur fonctionnement. Ces dispositifs constituent souvent une force d'inertie non négligeable. Pour nombre de grands groupes, il a fallu des années pour cesser de mesurer la performance de l'entreprise en quantité produite et en coût par unité produite et centrer leur comportement sur d'autres critères, tels que l'adaptation des produits aux demandes du marché, la qualité des produits et du service, la diversification de l'offre, etc.
- Les structures produisent leur propre dynamique, mais les acteurs, individus et groupes, qui peuplent ces structures, s'en servent également pour satisfaire leurs propres intérêts. Les décisions stratégiques, par essence, désignent des enjeux importants. Notamment, elles répartissent les ressources humaines et financières au sein de l'organisation. Les actions stratégiques de la direction générale sont décodées en ces termes par les acteurs, qui de plus nourrissent souvent des projets spécifiques. Des processus politiques (luttons, conflits, négociations, alliances, ruses, etc.) s'engagent autour de la formulation, de la sélection, et de la réalisation de ces actions et de ces projets. Il n'est pas toujours facile de distinguer clairement les intérêts particuliers et l'intérêt de l'entreprise.
- A titre d'exemple, la conception d'une automobile dans une grande firme implique une myriade de micro-décisions prises au niveau des services techniques. Les choix sont fondés sur le compromis, l'autorégulation, l'ajustement spontané. Les conflits potentiels sont tranchés par le sommet à des dates clés qui scandent les phases de la réalisation.

Revenons sur les limites de la pensée stratégique du dirigeant. Le décideur ne domine ni l'environnement auquel il est confronté ni l'organisation qu'il dirige (Figure I-8).

Il n'est pas non plus, ainsi que le mythe le dépeint, idéalement lucide et rationnel. La liste est longue des biais cognitifs auxquels les dirigeants sont exposés. Il faut convenir que les problèmes stratégiques, par leur importance, leurs ambiguïtés, leur singularité, favorisent les erreurs et les désillusions.

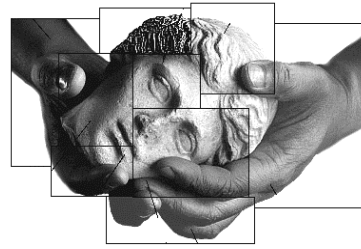


Figure I-8 : Limites de la pensée stratégique

Le dirigeant est confronté à un dilemme : en privilégiant les approches méthodiques, traitant des données objectives et autant que possible quantifiées, il perd les informations qualitatives, impalpables mais essentielles, car un problème stratégique est toujours une configuration complexe de faits et d'interprétations ; mais en faisant confiance à l'expérience, aux analogies, à sa perception, à l'intuition, il risque d'être victime de sa subjectivité. Le dirigeant n'aborde pas les situations d'un œil vierge. Il a notamment des représentations de ce qui est important ou non, de ce qu'il convient de faire ou non. Il s'appuie sur des règles et principes qui sont pour lui des évidences, mais qui pour l'observateur apparaissent comme des croyances, c'est-à-dire des manières de penser dont la validité n'est pas démontrée.

On peut ainsi s'interroger sur la capacité des dirigeants à contrôler la stratégie ? Entre les incertitudes générées par l'environnement, les dérives organisationnelles, les croyances des dirigeants, y a-t-il donc une chance de voir se former des décisions stratégiques sensées et cohérentes ? Il est vrai que ces tensions peuvent empêcher la formation d'une stratégie véritable mais dans beaucoup de cas, elles peuvent être sinon résolues, du moins contrôlées.

Souligner que les processus qui fabriquent la stratégie s'écartent souvent de la rationalité mise en scène par le dirigeant décideur ne signifie pas que le dirigeant soit impuissant. Il peut développer des moyens indirects de contrôle et d'influence allant dans le sens souhaité. Le dirigeant ne manque pas d'atouts du fait de sa position privilégiée à l'interface entre les unités de l'entreprise et en contact avec l'environnement. Il peut se poser en arbitre régulateur des conflits internes et favoriser les projets et les acteurs qui sont compatibles avec ses propres préférences et ses propres ambitions. Ainsi, par la mise en place de structures et de systèmes de gestion, ils définissent des règles du jeu qui peuvent induire une dynamique.

La véritable question semble pourtant être ailleurs. La configuration des structures et la répartition du pouvoir dans l'entreprise sont généralement des parties d'un système plus ou moins cohérent qui comprend également un ensemble de croyances portant sur la stratégie, largement partagé par le groupe des dirigeants, et souvent aussi par les autres niveaux du personnel. On parle alors de paradigme stratégique, par analogie avec les paradigmes scientifiques. Le paradigme décrit l'environnement, l'organisation dans l'environnement, et ce qui marche et ce qui ne marche pas pour survivre et prospérer dans cet environnement.

Le paradigme stratégique constitue une logique dominante qui détermine largement les résultats de la plupart des décisions. La question centrale est la capacité des dirigeants à en prendre conscience et à en jouer de manière à profiter de l'élan, tout en évitant les écueils et les dérives. Dans cette perspective, les événements stratégiques, faillite d'un gros client, entrée d'un nouveau concurrent puissant, apparition d'une nouvelle technologie de substitution, constituent des mises à l'épreuve du paradigme stratégique, des tests de sa capacité à maintenir l'entreprise sur une trajectoire favorable. Ils sont aussi pour les dirigeants l'occasion de faire évoluer le paradigme : redéfinir ce qui doit être tenu pour vrai et ce qui doit être tenu pour bon, et réorienter les forces qui façonnent la stratégie de l'entreprise.

On voit que le rôle du dirigeant, s'il n'a plus l'aura du grand décideur, est cependant fondamental, déterminant et... difficile ! Sa situation est assez inconfortable. On attend de lui une lecture claire d'un environnement incertain et souvent ambigu. Tenu d'incarner une volonté, il est pourtant dépendant, pour l'action concrète, d'une foule de relais—hommes, unités, procédures, systèmes. Entre les choix énoncés et les résultats promis et attendus, des années entières peuvent s'écouler, pendant lesquelles les orientations prises ne manqueront pas d'être contestées ou détournées, à la faveur d'événements extérieurs qui sembleront remettre en question son projet. La réalité des décisions stratégiques est donc bien éloignée de l'image mythique. Les réussites stratégiques se construisent rarement à l'aide d'intuitions fulgurantes ou de plans rigoureux. La capacité à produire une stratégie émergente à partir d'un flux de décisions et d'événements apparemment disjoints est une des premières conditions de survie pour les entreprises où le dirigeant est devenu un arbitre plus qu'un décideur mythique.

Toujours dans les stratégies d'entreprise, on peut signaler un modèle particulier, dit modèle de la poubelle (garbage can model [March, 1988]). Deux américains définissent à la fin des années 60, deux nouvelles notions. La première est celle d'anarchies organisées dont les universités sont selon, eux, un parfait exemple. Cette expression désigne les organisations :

- sans objectifs cohérents et partagés par tous ;
- où le processus de production relève d'une technologie complexe, peu matérialisable ;
- dont les membres participent de façon active aux prises de décision.

La seconde est le modèle de la poubelle qui décrit le style de décision propre aux anarchies organisées. Ce modèle remet en cause les théories où les décisions résultent d'une confrontation entre des objectifs identifiés, des solutions disponibles et leurs conséquences (modèle rationnel) et les théories où les décisions sont le résultat d'une négociation entre des groupes aux intérêts divergents (modèle politique). Dans les anarchies organisées, des choix sont à la recherche de problèmes, des questions cherchent des opportunités pour décider, des solutions cherchent des questions auxquelles elles pourraient être une réponse et des décideurs cherchent du travail...

Des décisions se produisent quand les flux de problèmes, de solutions, de participants et d'opportunité de choix se rencontrent. Toute prise de décision est ainsi assimilable à une poubelle où des types de problèmes et de solutions sont déchargés par les participants dès qu'ils sont générés et qui, se rencontrant, font émerger un choix.

La simulation informatique de ce modèle, a priori chaotique, ne fait pourtant apparaître que trois styles de choix possibles :

- les décisions par inattention ;
- par déplacement des problèmes ;
- et par résolution des problèmes.

Les deux premiers étant plus fréquents que le dernier. Cette émergence d'ordre dans les processus de décision anarchique a inspiré des travaux sur la capacité d'apprentissage des anarchies organisées. En particulier, on retrouve cette notion d'émergence élaborée dans le cadre de la technique des agents logiciels.

Les apports du modèle de la poubelle sont nombreux. Premièrement, il invite à ne pas surestimer la rationalité des acteurs, à ne pas surévaluer leur capacité à appréhender le processus où ils interviennent. De même, il incite les décideurs à rester modestes et à ne pas recourir à des outils managériaux trop ambitieux (planification, rationalisation...). Mieux vaut travailler à la marge et utiliser une technologie de l'absurde : traiter les objectifs comme des hypothèses, considérer que sa perception intuitive correspond à la réalité, que son expérience a valeur de théorie...

Ce modèle permet également de mieux prendre en compte les aléas : le timing est un élément décisif de l'analyse, tout comme les événements qui se produisent au cours d'une décision.

Ainsi, la présence (ou l'absence) d'une personne à une réunion peut modifier le sens de la décision.

Enfin, il permet de raisonner en dissociant les intentions des actions, les causes des effets. On peut ainsi trouver des problèmes parce qu'on détient une solution. On peut prendre une décision, non parce qu'on a un problème, mais parce qu'on a l'opportunité de faire un choix. On peut opter pour une solution même si elle ne répond pas au problème soulevé.

Stimulant et ludique, ce modèle donne un éclairage moins rationnel et linéaire et, par conséquent, plus contingent et aléatoire. Mais il comporte une philosophie du décideur souvent critiquée : inconstant dans ses préférences, infidèle à ses intentions, il est aussi incapable d'action rationnelle. A tel point qu'il disparaît de la scène et n'est pas central dans l'analyse du processus. Par crainte de sur-rationaliser, ce modèle tend aussi à sous-estimer tous modes de régulation, toutes règles du jeu formelles ou informelles pouvant circonscrire le comportement des acteurs [Musselin, 1990].

1.3. Les pièges de la raison

1.3.1. La logique des erreurs

L'objectif de la théorie mathématique de la décision était initialement d'expliquer comment un homme rationnel doit prendre des décisions dans l'incertitude. Le décideur construit un arbre de décision pour projeter les conséquences des différents choix possibles quant à la probabilité et à l'utilité des événements correspondants ; il décide entre les choix selon une règle de maximisation de la valeur espérée ou de l'utilité attendue. Cette première orientation s'est révélée incapable de prédire les décisions effectivement prises par des individus et donc a fortiori de décrire les processus mentaux qu'ils mettent en œuvre pour décider. On a expliqué cette incapacité par le fait que d'une part le système humain de traitement de l'information est limité, d'autre part l'individu n'estime pas et ne traite pas les probabilités selon les principes des théories probabilistes (Tableau I-2).

Le choix rationnel suppose la connaissance, l'examen et l'évaluation de toutes les options possibles, de tous les états de la nature actuels et futurs, des différentes issues possibles. Une fois tous ces aspects estimés sur plusieurs dimensions, il reste à faire la sélection présentant l'utilité maximale. Un tel processus requiert à la fois une grande disponibilité d'informations et une grande capacité de calcul. Kahneman, Slovic et Tversky ont mis quant à eux en évidence les propriétés de prise en compte et de traitement des probabilités par les individus [Kahneman *et al.*, 1982].

Dans la mesure où le type de phénomène d'écart à une norme illustré dans le Tableau I-2 est très général, l'idée s'est imposée qu'il traduit une propriété de l'activité mentale (cognitive). Il a été désigné sous le terme de biais cognitif. Les principales propriétés de l'activité mentale dans la décision ont été inférées à partir d'un inventaire des principaux biais cognitifs [Caverni *et al.*, 1990].

Les biais que subit le décideur sont de quatre types :

- l'acquisition de l'information ;
- le traitement de l'information ;
- l'expression de la réponse ;
- l'information qu'il reçoit en retour.

L'acquisition d'informations est biaisée dans plusieurs circonstances : lors de l'évocation de données en mémoire, lors de la sélection de données dans l'environnement, et en fonction du mode sous lequel les données se présentent. Ainsi la fréquence des événements qui font l'objet d'une certaine publicité est généralement surestimée par rapport à celle des événements dont on parle peu. La fréquence absolue des événements est préférée à leur fréquence relative. Les informations concrètes sont plus prégnantes en mémoire que les informations abstraites.

L'ordre dans lequel les informations sont présentées peut produire des effets dits de primauté ou de récence. Les premières ou les dernières informations induisent la décision finale.

Illustration	Théorie bayésienne
Un taxi est impliqué, la nuit, dans un accident de la circulation. Deux compagnies de taxi exercent dans la ville, la verte et la bleue. 85% des taxis de la ville sont verts et 15% sont bleus. Un témoin a identifié le taxi de l'accident comme étant bleu. Le tribunal a testé la capacité du témoin à identifier les taxis la nuit. Lorsqu'on lui a présenté un échantillon équiparti de taxis, le témoin a correctement identifié les taxis dans 80% des cas. Quelle est la probabilité que le taxi impliqué dans l'accident soit bleu ? Une très forte majorité des personnes soumises à ce problème répond que la probabilité que le taxi soit bleu est de 80%. La réponse exacte, selon la théorie bayésienne des probabilités est 41%. En général cette dernière estimation surprend et ne paraît pas intuitivement plausible. Les réponses observées traduisent le fait que l'individu néglige les probabilités a priori (ici 85% et 15%), qui donnent le cadre général de l'estimation, pour ne considérer que l'information spécifique (ici la fiabilité du témoin).	B1 : le taxi est bleu B2 : le taxi est vert A : le témoin a vu un taxi bleu $P(A/B_k)$ est la probabilité que le témoin voit un taxi bleu la nuit quand celui-ci est B_k. $P(A/B1) = .8$ $P(A/B2) = .2$ $P(B1) = .15$ $P(B2) = .85$ La probabilité $p(A)$ est la probabilité de voir un taxi bleu : $P(A) = p(A/B1).p(B1) + p(A/B2).p(B2) = 0.29$ Enfin, on applique la règle de Bayes pour calculer la probabilité que le taxi soit bleu quand le témoin affirme qu'il est bleu : $P(B1/A) = \frac{p(A/B1).p(B1)}{p(A)} = \frac{0.12}{0.29} = 0.41$

Tableau I-2 : La théorie de Bayes est-elle utilisée dans nos choix de tous les jours ?

L'individu apprend-il à décider ? Cela supposerait qu'il puisse reconnaître des relations stables entre événements dans l'information qu'il reçoit en retour de ces décisions. Or plusieurs biais interfèrent avec la possibilité d'apprentissage. Le fait que certains événements ne se produisent jamais est un obstacle pour estimer les conséquences de certaines décisions. D'autre part certaines décisions sont d'une fréquence trop faible pour qu'un individu puisse tirer profit d'éventuelles conséquences qu'il aurait observées. Enfin, lorsque des événements consécutifs à une décision sont accessibles, une tendance est de les attribuer systématiquement à l'action du décideur (illusion du contrôle).

Les biais mis en évidence par la psychologie cognitive [Hogart, 1987] attestent que le décideur ne fonctionne pas selon les principes des théories formelles. Toutefois les explications fournies par la psychologie cognitive sur les processus en œuvre dans la décision sont encore limitées : perception sélective, traitement séquentiel, capacités de calcul et de mémorisation limitées, sensibilité aux caractéristiques de la tâche, action sur l'environnement... Au lieu d'être algorithmique l'activité mentale du décideur serait heuristique : les procédures mises en œuvre seraient habituellement efficaces mais n'auraient pas de justifications rigoureuses.

Nous retiendrons que les processus psychologiques de la décision individuelle ne correspondent qu'imparfaitement au modèle d'un calculateur froid. Même lorsque l'individu se veut rationnel, de nombreuses erreurs logiques viennent piéger son raisonnement.

1.3.2. Décider dans l'urgence

Ajoutons que la décision doit parfois être prise dans l'urgence. En situation de crise, la mécanique de la décision se grippe parce que :

- la crise peut provoquer des dommages importants, et dans certains cas menacer la survie de l'organisation ;
- la réponse à apporter au problème n'est pas planifiée ;
- une attention immédiate est nécessaire ;

- l'événement est tel qu'il exige la mise en œuvre de toutes les ressources de l'organisation ;
- l'ensemble des caractéristiques de la crise va déclencher un stress important chez le décideur.

Les dommages que la crise peut provoquer, la complexité de la situation, et la rapidité avec laquelle il faut réagir créent une tension qui va peser sur la prise de décision. L'entrée en crise se matérialise souvent par un choc qui peut paralyser toute réaction du décideur. Le problème est tellement démesuré et on semble avoir tellement peu de prise sur les événements que l'on ne sait que faire. Passé ce cap, c'est en réfléchissant aux façons de répondre à la crise que l'on réalise la complexité de la situation. La résolution du problème ne va-t-elle pas entraîner l'aggravation d'un autre (la crise peut se révéler à la conjonction de défaillances techniques, humaines et organisationnelles).

La décision dans les situations de crise donne la limite des processus classiques. La pression empêche une évaluation sérieuse des possibilités de réponse et de leurs conséquences. Le stress initial provoqué par l'étendue des dommages est renforcé par la pression temporelle, l'incertitude et la complexité de la situation. Ainsi, ce qui peut aider le décideur, ce n'est pas tant la préparation de scénarios ou de procédures, que la prise de conscience qu'une crise peut survenir dans son entreprise, et donc de la préparation matérielle et psychologique à l'affronter [Gilbert, 1992; Lagadec, 1988].

1.3.3. Entre stratégie consciente et force aveugle

La décision peut être vue sous l'angle du déterminisme social ou de l'autonomie individuelle, des pulsions inconscientes ou du raisonnement lucide. Et si la diversité des regards tenaient moins aux conflits théoriques qu'à l'échelle d'observation adoptée ?

La décision en situation collective porte sur l'étude des processus collectifs, c'est-à-dire sur les interactions et les transactions sociales qui concourent à la production d'une décision. Les acteurs ont de « bonnes » raisons de décider ou de faire pression en faveur de telle ou telle solution. La décision est donc la résultante d'une somme d'interactions dont la logique peut parfois échapper aux acteurs eux-mêmes. Une clarification du débat sur la décision se trouve dans l'existence d'une différence d'échelles d'observation, l'obligation de découper dans la réalité, l'intégration d'un effet de situation dans le découpage de la réalité, la pertinence de la méthode des itinéraires pour reconstituer un processus de décision [Desjeux *et al.*, 1988].

La première échelle est l'échelle macro-sociale. C'est celle de la société. Elle permet de montrer comment les individus ont incorporé les modèles culturels, les codes, les styles de vie de leur groupe d'appartenance. Dans ce type d'approches macro-sociales, la décision individuelle disparaît au profit des régularités sociales. A cette échelle, la tendance est plutôt de sous-valoriser les calculs de l'acteur.

La deuxième échelle est interactionniste. L'observation porte sur les liens concrets entre les acteurs. Elle montre comment les décisions individuelles sont contingentes, fruit de négociations sociales ; comment le discours, les enjeux sociaux, les intérêts, le sens, l'émotionnel organisent le champ décisionnel. A cette échelle, la raison et/ou l'émotion de l'acteur ont tendance à être survalorisées.

La troisième échelle est micro-individuelle. Elle permet de comprendre les arbitrages par lesquels un individu raisonne ses choix. La méthode consiste alors à faire reconstruire par les décideurs les qualités qu'ils cherchent, puis à leur faire élucider les signes par lesquels ils reconnaissent cette qualité.

Ainsi, il est illusoire de croire que même le consommateur qui achète son poulet au supermarché procède à une décision individuelle, c'est la situation qui est individuelle. La

décision, elle, est tout autant qu'en organisation la résultante d'un processus dans le temps. L'acteur subit des influences sociales liées à son origine sociale, aux médias ou interactions familiales.

1.3.4. Décision raisonnée ou raisonnable

B.Jarrosion explique dans [Jarrosion, 1994] les raisons pour lesquelles la décision ne peut être ramenée à la norme du calcul :

- le calcul néglige la complexité, c'est-à-dire l'existence de critères d'évaluation multiples et non commensurables ;
- le calcul néglige le sujet connaissant qui interprète l'information, qui se construit en construisant ses croyances ;
- le calcul néglige que le décideur existe en tant qu'être autonome.

Ce n'est pas l'information brute qui sert à décider mais plutôt le sens qu'a cette information. Une information ne dit rien de ce qu'il faut faire, seule son interprétation l'indique. Son interprétation, autrement dit le sens qu'elle prend ou pas pour une personne. Le sens de l'information n'est jamais contenu dans l'information elle-même.

[Sfez, 1992] explique qu'il est classique de penser que la décision suit un schéma causal et linéaire à cinq temps :

- Formulation d'un désir et conception d'un projet y répondant ;
- Prise d'information ;
- Délibération ;
- Décision proprement dite ;
- Exécution.

Le politologue Sfez rejoint l'économiste Simon et explique alors que ce schéma linéaire rencontre les limites de toute linéarité : l'impossibilité de considérer des causalités enchevêtrées. L'exécution, par exemple, ne va-t-elle pas changer la conception même du projet, le but recherché ? D'une manière générale, chacune des phases modifie l'objectif initial. La délibération conduit à rechercher de nouvelles informations, etc. Le schéma est simple mais il ne correspond pas à la réalité de causalités enchevêtrées. Le simple et le compliqué se comprennent analytiquement, mais le réel n'est ni simple, ni compliqué puisqu'il est complexe.

Il convient de distinguer l'évidence et la pertinence. L'évidence relève d'une cohérence interne, d'une compatibilité avec une structure logique ; la pertinence relève d'une cohérence externe, d'une compatibilité avec une situation extérieure. Il y a les solutions logiques (évidence) et les solutions efficaces (pertinence). Dans la mesure où l'on peut regarder le monde avec différentes logiques, les solutions logiques et les solutions efficaces ne sont pas toujours les mêmes. La décision est un processus d'interactions entre le décideur et le monde. Il faut donc qu'elle soit en accord avec le monde plutôt qu'avec sa logique interne. Elle doit être pertinente plutôt qu'évidente.

Dans l'idée de pertinence apparaît la notion de contexte. Une action est pertinente par rapport à un contexte, qu'elle soit logique ou pas, évidente ou pas. La décision évidente est raisonnée, la décision pertinente est raisonnable. Le schéma causal linéaire est évident mais il n'est pas pertinent, raisonné mais pas raisonnable. Certaines causes sont obscures et ambiguës, soit qu'elles existent de façon souterraine soit qu'elles ne possèdent pas l'intensité que l'évidence leur attribue.

Jarrosion explique encore à propos du pouvoir qu'il existe deux façons d'être : s'opposer ou devenir. La première façon est réaliste et enfantine, tout au moins dans la vision du pouvoir

qu'elle met en œuvre, la deuxième est idéaliste et adulte. Ces deux façons d'être relèvent de la décision. Pour s'opposer comme pour mettre ses projets en œuvre, il faut décider. L'être est au cœur de la décision. En premier lieu, la décision a pour fonction de permettre à l'acteur d'agir. Le système agit à travers ses acteurs. Il faut donc que les acteurs aient l'impression (vraie ou fausse) d'être libres, de prendre des décisions. Ces décisions étant prises dans la perspective de projets, la décision est une façon d'être un devenant. En second lieu, la décision a pour fonction de permettre à l'acteur de supporter le monde. Si l'acteur n'opposait pas à chaque instant sa liberté de décider à la pesanteur des déterminismes, la vie ne lui serait pas agréable. La décision est une façon d'être en s'opposant au monde. On supporte plus facilement ce monde qui n'est pas tel qu'on le voudrait si on peut à chaque instant projeter sa liberté sur l'imperfection du monde, le redressant tout en distinguant ses scories.

2. Fiabilité humaine et aide à la décision

Dans cette section, nous allons reprendre quelques typologies de l'erreur humaine avant de proposer la nôtre instanciée au processus de la décision. Les classes d'erreurs font référence aux biais et défaillances que l'analyse bibliographique de la section précédente a évoqués. Cette réflexion nous permettra de proposer un schéma fonctionnel de premier niveau pour notre SIADG.

Pour qu'il y ait vraiment une résolution collective d'une situation de décision par le couple homme/machine, il est nécessaire que la méthodologie de conception du système d'aide à la décision tienne compte de l'aspect de coopération avec l'homme et intègre une analyse des activités de celui-ci tout au long du processus décisionnel. Si l'individu peut être considéré comme une composante d'adaptation et de réactivité de l'organisation face aux événements ou perturbations internes et externes à celle-ci, d'un autre côté, ses performances sont sujettes à d'importantes variations, qui peuvent avoir une incidence sur le processus et remettre en cause la mission qui lui a été confiée. Pour prendre la mesure de ce risque, il faut considérer les acteurs d'un processus décisionnel comme des facteurs de risque par la variabilité de leurs performances. Cette variabilité est à l'origine de la défaillance (inaptitude à assurer une fonction requise), qui peut conduire à l'erreur humaine (les actions effectives ne correspondent pas aux intentions) [Benkhannouche, 1996; Reason, 1993]. C'est justement en analysant ces erreurs humaines dans le cadre de l'activité de la décision, leurs manifestations et leurs conséquences, leurs liens avec les tâches confiées au décideur que l'on peut proposer des aides pertinentes pour l'aide à la décision.

Reprenons le schéma linéaire en cinq étapes de la décision (§ 1.3.4) pour illustrer nos propos. Les cinq phases étaient la formulation d'un désir et conception d'un projet y répondant, la prise d'information, la délibération, la décision proprement dite et l'exécution. L'absence d'une ou plusieurs de ces phases correspond à une déviation que l'on peut étiqueter :

- s'il manque l'exécution, on parlera de velléité ;
- si l'on passe directement de la conception à la décision, on parlera d'impulsivité ou d'instinct ;
- si l'on ne saute que la prise d'informations, on parlera de paresse, de manque de rigueur ;
- si l'on saute la délibération, on parlera de légèreté ou de précipitation ;
- si l'on s'arrête à la délibération, il s'agira plutôt d'intellectualisme.

Si la déviation n'est pas intentionnelle, alors il s'agit d'une défaillance humaine individuelle ou collective selon le cas qui peut conduire à l'erreur malheureuse, la mauvaise décision, l'action manquée.

2.1. L'erreur humaine

L'erreur est plus particulièrement l'objet d'étude de la fiabilité humaine, préoccupation qui remonte aux années soixante et qui visait à déterminer les risques liés aux erreurs humaines dans des activités telles que la conduite d'installations industrielles. La fiabilité humaine est à présent devenue la science des défaillances humaines [Villemeur, 1988], dont l'objet est de prédire et de prévenir l'erreur humaine dans le but d'optimiser la fiabilité globale et la productivité du processus que l'homme cherche à contrôler. Même si ces études sont classiquement le fruit d'analyses fiabilistes dédiées à la conduite de processus industriels, nous en proposons une brève synthèse car nous verrons plus loin dans cette section que la transposition de certains résultats à l'activité d'un collectif d'acteurs confronté à une décision complexe est assez naturelle et utile à la conception de notre SIADG. Nous considérerons alors que la décision est un ensemble de processus, principalement cognitifs—apprentissage, sélection, délibération, argumentation, etc.—, leur conduite relève d'une activité managériale.

En conduite d'ateliers, aborder l'étude de l'erreur humaine s'avère délicat, car l'analyse de l'erreur peut conduire à des interprétations divergentes selon que l'on considère les écarts par rapport à la tâche telle que la définit le prescripteur, par rapport à la tâche telle que la redéfinit l'opérateur, ou encore par rapport à l'activité réelle de l'opérateur [Cellier, 1990]. En effet, l'écart à la prescription n'est pas nécessairement à interpréter en terme d'erreur, à forte connotation négative ; elle peut être également bénéfique pour le processus, par la capacité de l'homme à produire des comportements qui favorisent une “récupération innovante mais efficace” [Poyet, 1990]. Par ailleurs, l'erreur humaine ne peut être réduite à une activité individuelle isolée, elle est la combinaison de facteurs à la fois internes et externes (conditions techniques et sociologiques, contexte).

La fiabilité humaine, en tant que discipline de l'ergonomie, s'intéresse aux conditions d'apparition de l'erreur et aux moyens de les prévenir ; la psychologie cognitive s'intéresse plus particulièrement aux mécanismes cognitifs de l'erreur. La première discipline s'intéresse davantage à des indices, des traces de l'erreur ainsi qu'aux moyens d'y remédier, tandis que la seconde vise à en comprendre et en expliquer la genèse.

Une des premières tentatives de classification des erreurs humaines date du début des années 80 [Rasmussen *et al.*, 1981]. Depuis, plusieurs typologies ont été proposées, nous en passerons trois en revue : celle de Reason en psychologie cognitive, celle de Nicolet en ergonomie et celle de de Keyser en psychologie du travail. Nous pensons qu'elles sont suffisamment représentatives, chacune dans un domaine particulier, pour être l'objet d'une étude comparative.

2.1.1. La typologie de Reason

Reason a introduit un système générique de modélisation de l'erreur (GEMS pour Generic Error-Modelling System), qui s'inspire du modèle de raisonnement de Rasmussen. Ce dernier distingue fondamentalement trois types de comportements cognitifs [Reason, 1993] : le comportement basé sur les automatismes (SB), le comportement basé sur les règles (RB) et le comportement basé sur les connaissances déclaratives (KB).

Ce modèle est général, c'est-à-dire que Reason s'intéresse à l'explication psychologique de l'apparition de l'erreur. Ce modèle est donc applicable à bon nombre des activités humaines.

L'hypothèse fondamentale dans le modèle GEMS est que le comportement SB est celui qui est prioritairement sollicité dans une tâche, puis le comportement RB dès la détection d'un problème dans l'exécution, puis le comportement KB si le problème n'est toujours pas résolu (Figure I-9). Cette typologie de comportements sert de base à la typologie d'erreurs de Reason : les erreurs associées au comportement SB sont dues à des actions non conscientes

(les lapsus et les ratés), tandis que les erreurs associées aux comportements RB et KB sont dues à des actions délibérées (les fautes). Par analogie avec la fiabilité technique, Reason parle de modes de défaillance qui donnent les indices observables par lesquels une erreur est décelée. Chacun de ces modes de défaillance a une ou plusieurs causes à caractère psychologique.

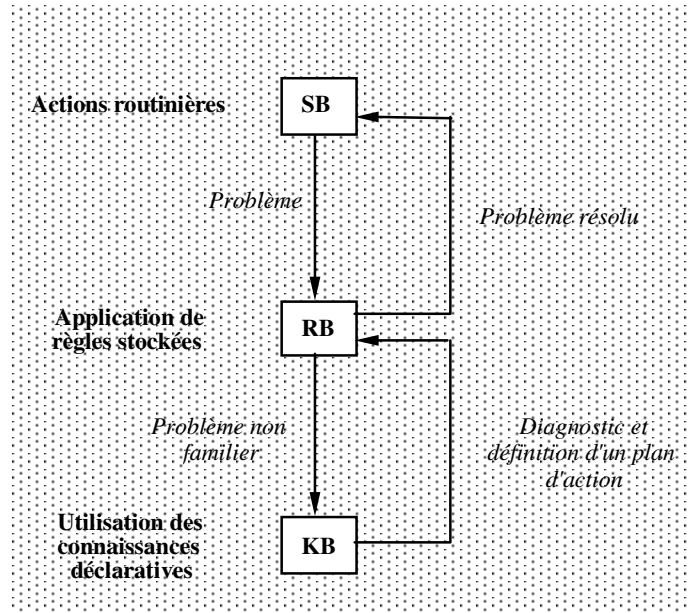


Figure I-9: Schéma de principe du modèle GEMS

Les modes de défaillance recensés par Reason du type SB sont l'inattention et l'attention excessive (Figure I-10). Pour l'inattention, les causes sont :

- une perturbation ou une interruption dans le séquençage d'actions en cours qui favorise les oublis ;
- une "intentionnalité réduite", c'est-à-dire qu'un délai entre la formation d'une intention et son exécution effective peut conduire à oublier l'objectif de la séquence d'action en cours ;
- une "confusion perceptive", qui traduit l'idée que des caractéristiques peuvent être perçues de façon très grossière quand celles-ci sont devenues familières (l'effort impliqué dans la reconnaissance est minimal) ; des caractéristiques approximativement similaires d'une action peuvent ainsi conduire à faire exécuter une action non souhaitée.

Dans le cas d'une attention excessive, la vérification à contretemps peut conduire à perturber l'action en cours par l'attention qui se fixe sur un point particulier, alors que cela n'est pas prévu ; cela survient pour les tâches fortement automatisées et peut conduire à des difficultés de repositionnement par rapport à la tâche en cours.

Type de comportement	Modes de défaillance	Causes psychologiques
SB	<i>Attention excessive</i>	Vérification à contretemps
	<i>Inattention</i>	Perturbation Intentionnalité réduite Confusion perceptive

Figure I-10: Erreurs au niveau SB

Les modes de défaillances pour le comportement RB sont une application erronée de bonnes règles et l'application de mauvaises règles (Selon Reason, "une bonne règle est celle qui s'est révélée utile dans une situation particulière") (Figure I-11).

L'application erronée de bonnes règles a pour causes :

- une "règle forte mais fausse", lorsque le sujet est face à une situation non reconnue et qu'il applique une règle qui s'est jusqu'ici avérée utile faute d'en avoir une adaptée ;
- une situation peut être ambiguë, c'est-à-dire perçue à travers des "signes" et des "contresignes" ; certaines caractéristiques peuvent confirmer une règle alors que d'autres peuvent venir la contredire, il peut également arriver que ces contresignes soient rejetés car ne correspondent pas à l'image que se fait le sujet ;
- une capacité de traitement limitée, qui implique que l'ensemble des informations disponibles, qui peuvent conduire à se faire une représentation complète de la situation, ne sont pas toutes traitées ;
- une compétition entre règles, qui indique une préférence pour certains signes au détriment d'autres ; plus une règle est forte, plus elle aura tendance à être considérée comme la plus vraisemblable ;
- un "conservatisme cognitif", qui conduit le sujet à ne pas changer les règles qu'il a établies au profit de règles plus élégantes et plus simples.

L'application de mauvaises règles a pour causes :

- une représentation partielle de l'espace-problème, qui conduit à des actions sur certains aspects et à négliger les autres (exemple du conducteur débutant qui éprouve des difficultés à contrôler la direction et le passage de vitesses) ;
- une représentation incorrecte de l'espace-problème, c'est-à-dire une représentation simpliste des propriétés physiques du monde ;
- de mauvaises habitudes, qui consistent à employer des règles erronées (erreur de calculs, cf **Tableau I-2**), maladroit, inélégant ou déconseillé (violation consciente de règles de sécurité) pour atteindre un objectif ; ces règles peuvent avoir certains avantages mais sont potentiellement dangereuses ;
- une confiance excessive en ses capacités : le sujet ne respecte pas certaines consignes, de sécurité notamment, car il est convaincu qu'il aura toujours suffisamment de temps pour faire face aux situations redoutées le cas échéant (c'est le cas de l'automobiliste qui ne respecte pas les distances de sécurité recommandées sur autoroute).

Type de comportement	Modes de défaillance	Causes psychologiques
RB	<i>Mauvaise application de règles correctes</i>	Règles fortes mais fausses Ambiguïté des situations Capacité limitée de traitement Compétition entre règles "Conservatisme cognitif"
	<i>Application de mauvaises règles</i>	Représentation partielle de l'espace-problème Représentation incorrecte de l'espace-problème Mauvaises habitudes Règles déconseillées

Figure I-11: Erreurs au niveau RB

Les défaillances du niveau KB se produisent quand le sujet met en œuvre son attention sur un problème qui n'a pu être résolu au niveau RB (Figure I-12).

Les modes de défaillance peuvent être :

- une récapitulation biaisée ou "illusion du pointage", qui indique que le plan d'action est évalué par rapport à certains aspects de la situation et non de manière systématique et rigoureuse sur l'ensemble des données disponibles ;
- une résistance au changement, qui indique qu'un plan d'action initialement établi est maintenu malgré les nouvelles données qui indiquent que l'objectif poursuivi n'est pas accessible de cette façon ;
- une limitation de l'espace de travail, ou principe du "premier entré/ premier sorti" selon lequel l'interprétation des caractéristiques du problème est réalisée par rapport à un modèle mental (la charge de travail est différente selon la façon dont le problème est posé) ;
- "hors de la vue, hors de l'esprit", ou "heuristique d'accessibilité", selon laquelle la priorité est donnée aux faits qui viennent facilement à l'esprit ;
- une difficulté à maîtriser la complexité, qui est liée à :
 - une perception limitée des aspects temporels, c'est-à-dire une mauvaise maîtrise des processus présentant un temps de réponse important (les actions sont alors réactives, c'est-à-dire dirigées par les données, plutôt qu'anticipatives) ;
 - une difficulté à traiter les évolutions exponentielles, c'est-à-dire qu'il y a sous-estimation systématique de l'importance de certains phénomènes ;
 - une forte tendance à raisonner selon des séquences linéaires, c'est-à-dire que les effets des actions ne sont pas considérés sur l'ensemble du processus ;

- un "vagabondage thématique et enkystement", qui apparaissent dès qu'un problème ne trouve pas de solution ; le vagabondage est le passage d'un aspect à un autre sans démarche précise, l'enkystement est une focalisation excessive sur un aspect particulier du problème (De Keyser emploie le terme de fixation [De_Keyser *et al.*, 1990a]) ;
- un biais de confirmation, qui correspond à un phénomène de rejet d'ambiguïté, c'est-à-dire qu'une préférence est accordée aux interprétations disponibles malgré l'arrivée d'informations qui viennent la remettre en question ;
- une difficulté à se représenter la causalité, en raison d'une simplification excessive des phénomènes causaux, qui conduit à juger sur des critères de similarité avec des cas antérieurs, donc à ne pas considérer l'ensemble des cas possibles ;
- une difficulté à diagnostiquer, directement liée à la conduite de raisonnement ; l'erreur peut provenir d'une imbrication néfaste entre assimilation des symptômes et génération des scénarii qui ont pu conduire aux symptômes observés ;
- une mauvaise sélectivité, qui indique que trop d'attention est accordée à des caractéristiques non pertinentes, c'est-à-dire que l'attention est absorbée par des aspects psychologiquement attirants.

Selon Reason, deux types d'erreurs doivent être envisagés :

- les erreurs actives concernent directement l'activité des hommes, elles présentent la caractéristique d'être récupérables, en ce sens que l'homme peut parvenir à corriger ses erreurs avant que leur effet ne se fasse sentir ;
- les erreurs latentes sont, par exemple, les erreurs de conception du processus que l'homme cherche à contrôler ; leur caractéristique est donc leur permanence. Ce sont des défauts insidieux dont on peut difficilement quantifier l'impact.

Les expériences menées en psychologie montrent que l'erreur active est présente dans tous les aspects de l'activité humaine (posture, parole, actions, résolution de problème). Des travaux expérimentaux ont également montré que la grande majorité des erreurs sont corrigées avant leur effet sur le processus. Reason cite trois modes de détection de l'erreur : par un écart entre l'intention et le résultat observé (analogie avec la boucle de rétroaction en automatique), blocage des actions par les automatismes (ce sont généralement les sécurités mises en place en conception ou les règles d'intégrité de l'organisation), et la détection par un tiers (contremaître ou supérieur hiérarchique).

Selon Reason, la non détection est susceptible de conduire à des conséquences graves, et l'étude des accidents du type *Three Mile Island* ont permis de montrer que certains processus cognitifs sont plus particulièrement responsables de la non détection d'erreurs par l'homme.

Les travaux expérimentaux relatés par Reason ont un caractère expérimental fort, puisque les résultats ont été obtenus dans des conditions et des contextes différents de ceux rencontrés in situ, et les conclusions tirées n'ont pas une portée générale. Cependant, le résultat qu'il faut en tirer est le suivant : vouloir réduire les erreurs humaines est illusoire, il faut plutôt s'attacher à se prémunir contre les conséquences possibles d'erreurs difficilement corrigées par les hommes, sans une intervention auxiliaire.

Type de comportement	Modes de défaillance	Causes psychologiques
KB	<i>Récapitulation biaisée</i>	"Illusion du pointage"
	<i>Résistance au changement</i>	Résistance au changement
	<i>Limitation de l'espace de travail</i>	Principe du "premier entré/ premier sorti"
	<i>Hors de la vue, hors de l'esprit</i>	"Heuristique d'accessibilité "
	<i>Difficultés avec la complexité</i>	Perception des aspects temporels limitée Difficulté à traiter les "évolutions exponentielles" Forte tendance à raisonner selon des séquences linéaires Vagabondage thématique et enkystement
	<i>Biais de confirmation</i>	Rejet d'ambiguïté
	<i>Difficultés avec la causalité</i>	Simplification excessive de la causalité
	<i>Difficultés de diagnostic</i>	Conduite de raisonnement
	<i>Sélectivité</i>	Attention du sujet

Figure I-12: Erreurs au niveau KB

2.1.2. La typologie de Nicolet

Nicolet a une approche plus spécifiquement dédiée à la conduite d'ateliers de production. Il propose une typologie basée sur une approche ergonomique de l'erreur et établie à partir des étapes d'une tâche ; avec ce découpage, l'auteur distingue sept types d'erreurs [Nicolet *et al.*, 1985] :

- les erreurs de perception ont des causes multiples, par exemple l'information n'est pas correctement perçue, le signal est fugitif et l'opérateur n'a pas le temps de le percevoir ou de le mémoriser, l'information est masquée par un écran, l'information est noyée au milieu d'un grand nombre d'autres, etc. ;
- les erreurs de décodage relèvent d'une mauvaise interprétation des indices sensoriels ;
- le non respect d'une procédure est un acte conscient qui peut ne pas se révéler néfaste pour le bon fonctionnement du processus, mais il est le signe de mauvaises habitudes de la part des opérateurs ;
- les erreurs de communication interviennent lors de l'échange d'informations entre agents ; une information peut être mal transmise ou mal comprise par l'opérateur qui reçoit ;
- les décisions peuvent ne pas être prises en temps voulu : une fois le diagnostic établi, l'opérateur repousse la décision en attendant un complément d'information ;

- les actions mal dosées ou mal adaptées relèvent d'une réalisation incorrecte d'une action requise, du non accomplissement d'une action nécessaire, de l'exécution d'une action impromptue (non prévue dans la séquence en cours), de l'inversion de l'ordre d'exécution de deux actions ;
- les erreurs de représentation ou de modèle traduisent les écarts qui peuvent exister entre les certitudes de l'opérateur et la situation réelle.

2.1.3. La typologie de de Keyser

De Keyser propose une approche davantage orientée vers la psychologie du travail. L'auteur distingue quatre grandes classes d'erreurs [De_Keyser, 1990b] : erreurs d'évaluation, de coordination, d'intégration et de planification.

Les erreurs d'évaluation sont de trois types :

- les erreurs de filtrage¹ correspondent à des informations que l'opérateur ne prend pas en compte, par exemple un changement, car il se fonde davantage sur une faible probabilité de dérèglement d'un paramètre et ne le consulte pas ou car il néglige de l'information ;
- les erreurs dues à une carence ergonomique traduisent les situations mal perçues car l'opérateur n'a pas les moyens d'évaluer l'apparition ni la durée d'un événement, il n'existe pas de vue synoptique des événements du système ;
- les erreurs de mode résultent d'une mauvaise mise en correspondance de quatre modes d'appréhension du temps par l'opérateur : intériorisation des événements réguliers, mode logique (mise en correspondance d'événements et de séquences d'action), mode causal (événements produits par le processus) et mode de l'horloge (les délais et contraintes de production).

Les erreurs de coordination dans l'équipe de conduite sont dues, par exemple, à une charge de travail trop importante (qui peut entraîner des délais et oublis de transmission d'informations), à une mauvaise évaluation temporelle (mauvaise connaissance de l'état général du système), à une difficulté de joindre les personnes à contacter, ou bien à des facteurs psychologiques ou organisationnels (par exemple rétention d'information et conflits).

Les erreurs d'intégration concernent spécifiquement les décisions d'actions d'ajustement des actions de l'équipe en fonction de retards ou d'avances au niveau de la production. Elles se caractérisent par une évaluation correcte de la situation, mais à une erreur dans la prise de décision, par la difficulté de l'opérateur à suivre les différents événements qui constituent le processus, ou encore par des difficultés à intégrer les actions de l'équipe avec les événements du processus.

Enfin, les erreurs de planification apparaissent lorsqu'un écart trop important entre déroulement temporel prévu et déroulement actuel requiert de la part de l'opérateur, seul ou en accord avec sa hiérarchie, la prise de décisions bouleversant la planification initiale. Ces décisions portent généralement sur des modifications de l'ordre des événements ou de leur durée. Elles sont caractérisées par une mauvaise évaluation de l'état temporel du système, une mauvaise connaissance des contraintes fixes et des zones de flexibilité, et une mauvaise connaissance des liens de causalité entre événements ou actions.

¹Toute l'information de l'environnement de travail n'est pas prise en compte par l'opérateur et il ne retient que ce qui est important pour lui. L'erreur provient du fait qu'il néglige certaines informations au profit d'autres qu'il considère relever davantage de l'état de marche du système.

2.1.4. Caractéristiques de ces typologies

Les typologies de Nicolet et de de Keyser sont orientées supervision de processus de fabrication. Celle de de Keyser est principalement centrée sur la description des erreurs cognitives (évaluation, intégration, planification), alors que celle de Nicolet décrit plus spécifiquement les erreurs commises au niveau des étapes d'une tâche (perception, décodage, représentation, décision, définition de procédures, action, communication). La typologie de Reason s'intéresse plus particulièrement aux mécanismes de production de l'erreur et propose une explication psychologique à leur occurrence ; elle a une portée plus générale que les deux autres et n'est pas spécifique à la conduite de processus industriels.

Ces typologies répondent, en fait, à des objectifs spécifiques [Benkhannouche, 1996] :

- pour Nicolet, il s'agit d'avoir une typologie utilisable par l'ergonome, c'est-à-dire qui mette l'accent sur les traces de l'erreur, de façon à permettre de proposer des solutions ergonomiques ;
- pour de Keyser, la typologie d'erreurs doit prendre en compte la dimension temporelle de l'activité de conduite ;
- pour Reason, il s'agit de recenser les causes psychologiques des erreurs commises, en conduite de processus industriel comme dans les situations de la vie quotidienne.

2.2. Instanciation à l'activité de décision en organisation

Le découpage du processus de la décision selon trois dimensions de l'activité humaine, collective, cognitive et opératoire permet de regrouper les erreurs humaines.

La dimension collective ou sociale tient compte des erreurs dues à des problèmes de coordination entre les membres de l'organisation impliquée dans le processus décisionnel.

La dimension cognitive, qui regroupe les aspects ajustement, adaptation et anticipation dans le processus décisionnel est exposée aux erreurs d'évaluation (mauvaise interprétation des observations ; informations que le décideur ne prend pas en compte parce qu'elles sont associées à des événements dont la probabilité d'occurrence est faible, à des valeurs dont une variation significative lui semble improbable : le décideur ne va donc pas s'en enquérir et négliger ni plus ni moins cette information ; identification déficiente de la causalité des événements), erreurs d'intégration (interprétation correcte de la situation mais action non pertinente, mal jaugée, effets inopportuns ; incohérence entre le temps de réponse des mesures mises en place et la réalité de la situation ; incapacité à évaluer le temps de réponse de l'action entreprise) et erreurs de planification (mauvaise anticipation des évolutions et mauvais choix de la stratégie d'action ; incapacité à estimer la durée d'une crise, les délais et contraintes liés à l'urgence ; mauvaise évaluation de la nature de la situation (situation de routine, de maîtrise, de crise, d'exception)).

Enfin, la dimension opératoire, qui tient compte des actions, des procédures... Les erreurs qui peuvent être commises ici relèvent spécifiquement d'une mauvaise réalisation des plans d'actions.

Cette classification résulte essentiellement de la fusion des typologies proposées par Nicolet—importance de la représentation et de la compréhension de la situation de décision—et de Keyser—complexité de l'interprétation de la situation liée aux aspects dynamiques du processus décisionnel. Ces deux points essentiels sont à l'origine des fonctionnalités de notre SIADG. Le [Tableau I-3](#) propose de classer les éléments bibliographiques évoqués dans la section précédente concernant la décision vue par les SHS selon notre typologie d'erreurs. Sur la base de cette typologie d'erreurs, on peut alors se faire une idée plus précise des aides informatisées à spécifier ainsi que de leur « spectre d'action » pour l'aide à la décision.

Erreurs de Coordination	• Répartition des rôles et tâches mal définies • Problèmes relationnels • Organisation rétive qui sélectionne ou favorise tel ou tel scénario en fonction de ses intérêts • Il n'y a pas de décision, un décideur, mais une série de stratégies et de compromis entre les points de vue, entre des groupes qui ne partagent pas la même solution
Erreurs d'évaluation	• Disparité, dispersion et divergence des opinions difficilement contrôlables (conflit de préférences) • Multiplicité de rationalités différentes qui s'imbriquent, se superposent, se confrontent (affrontement de logiques décisionnelles) • Incertitudes épistémiques, état des connaissances jugé insatisfaisant • Principe de la rationalité limitée • Les informations disponibles sont généralement incomplètes, imprécises ou contradictoires • Même lorsqu'il y a un exécutif unique, il s'appuie sur des solutions concrètes dont la formulation détaillée, techniquement constituée, n'est pas son œuvre. Des groupes tiers fournissent ces solutions, élaborent les options et assurent leur légitimité opérationnelle • L'argumentaire est un incontournable exercice de rhétorique pour légitimer une décision engageant un collectif d'acteurs • La liste est longue des biais cognitifs auxquels les dirigeants sont exposés. Les problèmes stratégiques, par leur importance, leurs ambiguïtés, leur singularité, favorisent les erreurs et les désillusions • Le dirigeant n'aborde pas les situations d'un œil vierge. Il a notamment des représentations de ce qui est important ou non, de ce qu'il convient de faire ou non. Il s'appuie sur des règles et principes qui sont pour lui des évidences, mais qui pour l'observateur apparaissent comme des croyances, c'est-à-dire des manières de penser dont la validité n'est pas démontrée • Un problème stratégique est toujours une configuration complexe de faits et d'interprétations • Les biais cognitifs que subit le décideur sont de quatre types : l'acquisition de l'information, le traitement de l'information, l'expression de la réponse, l'information qu'il reçoit en retour • Perception sélective du décideur, traitement séquentiel, capacités de calcul et de mémorisation limitées, sensibilité aux caractéristiques de la tâche, action sur l'environnement... Au lieu d'être algorithmique l'activité mentale du décideur est heuristique : les procédures mises en œuvre sont habituellement efficaces mais n'ont pas nécessairement de justifications rigoureuses • La décision ne peut être ramenée à la norme du calcul parce que le calcul néglige la complexité, c'est-à-dire l'existence de critères d'évaluation multiples ; le calcul néglige le sujet connaissant qui interprète l'information • Ce n'est pas l'information brute qui sert à décider mais plutôt le sens qu'a cette information. Une information ne dit rien de ce qu'il faut faire, seule son interprétation l'indique. Son interprétation, autrement dit le sens qu'elle prend ou pas pour une personne. Le sens de l'information n'est jamais contenu dans l'information elle-même
Erreurs d'intégration	• Incertitudes générées par l'environnement quant aux temps de réponse des actions envisagées • La décision peut être un processus long et enchevêtré (boucles cognitives, causalités enchevêtrées) • Dans une décision de groupe, s'il n'y a pas nécessité d'aboutir à une décision commune (compromis), les membres d'un groupe débattent de leurs opinions personnelles, mais sans chercher à retravailler cognitivement les opinions exprimées par les autres, ce qui peut pénaliser la dynamique d'obtention d'un consensus • L'objectif d'un consensus ne doit pas être de supprimer les conflits, mais de les tolérer. Plus les options individuelles sont variées, plus la confrontation est importante et plus la reformulation du problème peut devenir explicite. La dynamique de consensus est fonction de l'intelligence de cette phase de délibération • La décision en organisation n'obéit pas à un modèle rationnel linéaire, de nombreuses itérations sont nécessaires entre les phases d'information, de représentation et de sélection • L'ordre dans lequel les informations sont présentées peut produire des effets dits de primauté ou de récence. Les premières ou les dernières informations induisent la décision finale
Erreurs de planification	• Incertitudes générées par l'environnement, imprédictibilité de la réaction de l'environnement • Entre les incertitudes générées par l'environnement, les dérives organisationnelles, les croyances des dirigeants, y a-t-il une chance de voir se former des décisions stratégiques sensées et cohérentes ? • Les processus qui fabriquent la stratégie s'écartent souvent de la rationalité mise en scène par le dirigeant décideur mais cela ne signifie pas que le dirigeant soit impuissant. Il peut développer des moyens indirects de contrôle et d'influence allant dans le sens souhaité. Le dirigeant peut se poser en arbitre régulateur des conflits internes et favoriser les projets et les acteurs qui sont compatibles avec ses propres préférences et ses propres ambitions. Ainsi, par la mise en place de structures et de systèmes de gestion, ils définissent des règles du jeu qui peuvent induire une dynamique • Les événements stratégiques non prévus constituent des mises à l'épreuve du paradigme stratégique, des tests de la capacité du dirigeant à maintenir l'entreprise sur une trajectoire favorable. Ils sont aussi pour les dirigeants l'occasion de faire évoluer le paradigme : redéfinir ce qui doit être tenu pour vrai et ce qui doit être tenu pour bon, et réorienter les forces qui façonnent la stratégie du groupe • La capacité à produire une stratégie émergente à partir d'un flux de décisions et d'événements apparemment disjoints est une des premières conditions de survie pour l'organisation où le dirigeant est devenu un arbitre plus qu'un décideur mythique. • Il convient de distinguer l'évidence et la pertinence. L'évidence relève d'une cohérence interne, d'une compatibilité avec une structure logique ; la pertinence relève d'une cohérence externe, d'une compatibilité avec une situation extérieure • Le décideur in fine sert de détonateur à l'élaboration de nouvelles politiques au sens large du terme
Erreurs d'exécution	• La formulation de la stratégie serait l'apanage du dirigeant et sa mise en œuvre serait la mission de l'ensemble de l'entreprise. Cette distinction si naturelle occulte le rôle considérable de l'organisation dans la réalisation de la stratégie • Mauvaise application des plans d'actions

Tableau I-3 : Classifications d'éléments bibliographiques de la décision vue par les SHS et notre typologie des erreurs humaines

Les fonctionnalités d'un SIADG que l'on peut spécifier sur la base de cette analyse peuvent être ainsi répertoriées.

Les erreurs de coordination peuvent être dues à des problèmes de communication, de répartition des rôles et des tâches, ou à des problèmes d'ordre relationnels, elles concernent de toute façon la dimension collective de la décision. Les améliorations envisageables sont d'augmenter la résolution collective de problème, de négocier les conflits et d'instaurer des relations de confiance entre les membres de l'organisation, des acteurs de la décision. Le comportement collectif pourrait être assisté par des systèmes stimulant une meilleure circulation des informations et une meilleure répartition des tâches. Ces aides relèvent de l'intelligence collective, du travail collaboratif [Penalva *et al.*, 2002]. Ainsi, par la mise en œuvre de systèmes informatisés de gestion, il est possible d'induire une dynamique et une réactivité nouvelles dans l'organisation. La délibération entre différents systèmes de valeurs est favorisée, le travail collaboratif soutenu par une plate-forme informatisée d'échanges des points de vue encourage le développement de coalitions, la négociation de compromis ou l'élaboration de consensus.

(i) Les erreurs d'évaluation peuvent être réduites par des aides destinées à améliorer la représentation de la situation et la perception de sa dynamique en particulier. Cela consiste à fournir une information compatible avec les modes de représentation et de raisonnement des acteurs de la décision. Il s'agit notamment :

- de concevoir une véritable interface homme-machine (IHM) pour le SIADG, de manière à présenter aux différents acteurs une information pertinente et dynamique par rapport à la situation et adaptée à leurs modes de raisonnement, leurs systèmes de valeurs. Les connaissances utiles à l'action doivent être facilement accessibles, être restituées de manière synthétique et cohérente avec les modes cognitifs des acteurs : le SIADG doit favoriser une lecture aussi claire que possible d'une situation incertaine et souvent ambiguë. La décision et l'apprentissage sont des processus dynamiques qui exigent donc une IHM spécifique de monitoring des indicateurs décisionnels dans le temps ;
- de permettre au décideur d'évaluer rapidement une situation et de l'aider à diagnostiquer les risques encourus, c'est-à-dire fournir une aide au diagnostic. Le SIADG doit permettre d'automatiser autant que possible le traitement de l'information afin de suppléer autant que possible l'être humain quant à la combinatoire des possibles, la multitude et la variété des opinions. Toute décision doit être accompagnée d'une estimation des risques attachés à la situation qu'ils soient dus à une déficience épistémique ou à une forte variabilité ou dispersion de l'opinion. Cette estimation ne vaut que si elle est comparée aux autres alternatives de la décision et si ce diagnostic est justifié par l'état des connaissances au moment où l'action a été retenue. L'évaluation repose sur un état des connaissances que le SIADG doit pouvoir faire valoir à tout instant pour soutenir la rhétorique décisionnelle du collectif et alimenter le débat. L'évaluation assistée de la situation doit canaliser la réflexion sur les points critiques et ainsi favoriser la dynamique de réduction des incertitudes épistémiques².

² L'incertitude épistémique est l'incertitude relative au degré de connaissance. A une date donnée, l'état des connaissances détermine la précision et la vraisemblance de l'évaluation de la situation et par suite la fiabilité de la décision qui en découle. Pour parvenir à une situation décidable, cette incertitude devra être dissipée. C'est l'objet de la conduite du processus décisionnel.

(ii) Les erreurs d'intégration révèlent des écarts entre les adaptations apportées par l'organisation concernant le bon déroulement de la planification et les contraintes temporelles imposées par la situation. Une amélioration consisterait à proposer aux décideurs les moyens de s'adapter aux événements perturbateurs qui modifient les conditions dans lesquelles les objectifs peuvent être maintenus ; il peut s'agir d'une remise en question de la perception de la situation comme d'une modification de la stratégie. L'aide à l'adaptation consiste à proposer aux acteurs de la décision des indicateurs sur la marche à suivre, c'est-à-dire du conseil d'action. Le SIADG doit donc être en mesure de proposer la dimension stratégique sur laquelle agir ou l'information à acquérir pour réduire au mieux le risque épistémique et tendre au plus tôt vers une situation décidable.

(iii) Les erreurs de planification peuvent être levées en fournissant au décideur les moyens de décider d'une intervention dans les délais imposés par la situation. Il semble que les aides correspondantes sont celles qui permettraient de fournir des stratégies basées sur le diagnostic de la situation : l'incapacité d'atteindre un niveau de risque acceptable avec la stratégie entreprise. Ces aides sont dites aides à l'anticipation. Ce type d'aide diffère notamment de l'aide à l'adaptation par le fait qu'une stratégie implique que le système d'information soit capable d'évaluer une situation au regard des objectifs et des contraintes de la planification. Le SIADG doit donc être capable d'anticiper qu'aucun consensus ne pourra être obtenu pour une stratégie définie a priori et que le modèle décisionnel—contraintes et objectifs—n'est pas viable et qu'il doit être modifié en conséquence. Cette évolutivité du système de valeurs utilisé par le SIADG dans son module d'évaluation est essentielle car elle laisse entrevoir un élément de réponse au problème qu'entre les choix énoncés et les résultats promis et attendus, le temps peut être long, et que pendant ce temps, les orientations prises ne manqueront pas d'être contestées ou détournées, à la faveur d'événements extérieurs qui sembleront remettre en question le projet.

(iv) Les erreurs d'exécution peuvent avoir des causes multiples, par exemple le non respect de certaines règles organisationnelles, auquel cas des mesures visant à supprimer les mauvaises habitudes doivent être envisagées. Ces erreurs peuvent également être dues à la mauvaise sélection ou réalisation d'une procédure. Une solution consiste alors à proposer aux différents acteurs de la décision des supports pour l'application de procédures en particulier en situation de crise (conduite à tenir en cas d'accident nucléaire, d'incendie, etc.), des outils pour la validation et la vérification du bon déroulement de la planification.

Les aides de la dimension cognitive apparaissent essentielles (Figure I-13). Il s'agit d'assister les décideurs dans leur évaluation d'une situation multidimensionnelle et multi acteurs avec un état des connaissances à un instant donné. Il faut faciliter la restitution de la logique décisionnelle appliquée, en interne pour soutenir les acteurs dans leurs modes cognitifs d'une part, à des fins argumentatives pour un public externe d'autre part. Une fois la situation correctement appréhendée, il s'agit ensuite d'assister les décideurs dans la dynamique à mettre en place pour ne pas laisser au hasard l'évolution de la situation, pour poursuivre et fiabiliser leur entreprise de manière cohérente avec leur système de valeurs d'une part, pour diagnostiquer au plus tôt toute déviance manifeste de l'opinion quant aux objectifs fixés. L'aide à l'évaluation permet

- d'estimer la valeur de chaque solution ;
- de les comparer, de les argumenter.

L'aide à l'adaptation permet de s'assurer que les actions entreprises sont cohérentes avec le système de valeurs et la dynamique de la situation.

L'aide à la planification permet de diagnostiquer le changement de nature de la situation et d'anticiper une orientation avec un autre système de valeurs.

Les aides à l'évaluation ont clairement un objectif commun : réduire le risque à prendre une décision avec un état des savoirs donné—le risque épistémique—c'est-à-dire gérer la rationalité limitée.

Le schéma fonctionnel de la [Figure I-14](#) synthétise l'ensemble de ces réflexions dans le formalisme de SADT et propose une description des fonctionnalités au plus haut niveau du SIADG qui se fonde sur le modèle décrit dans le chapitre suivant.

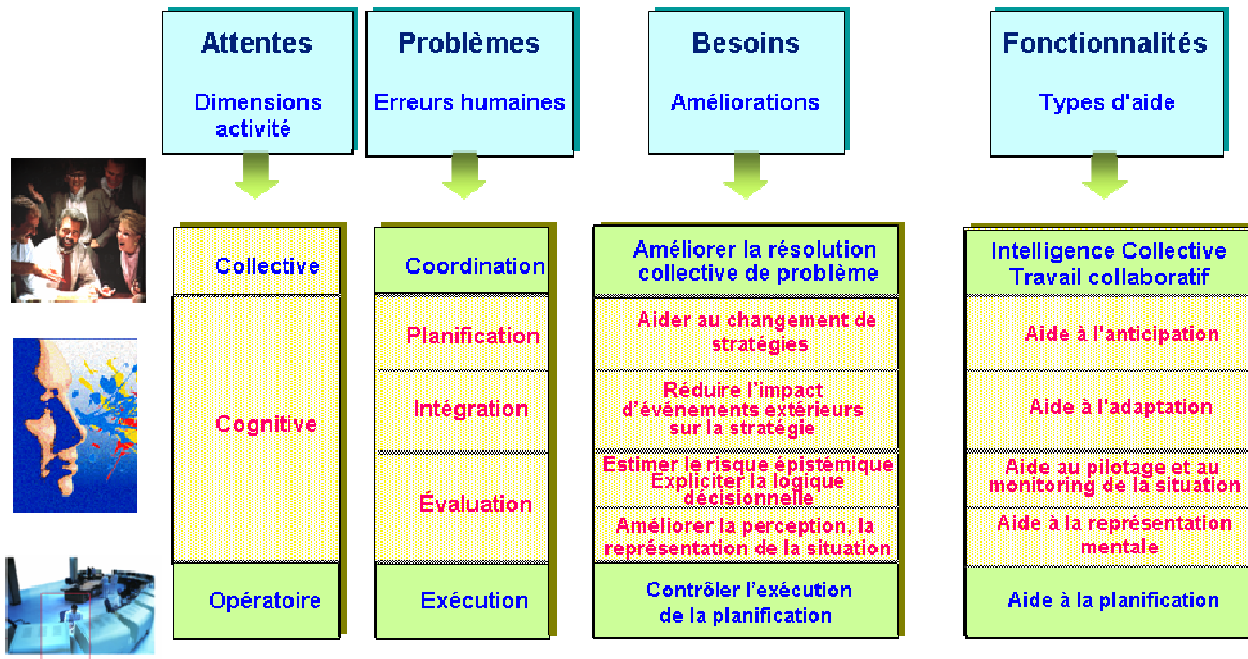


Figure I-13: L'aide à la décision collective sous l'angle de l'erreur humaine

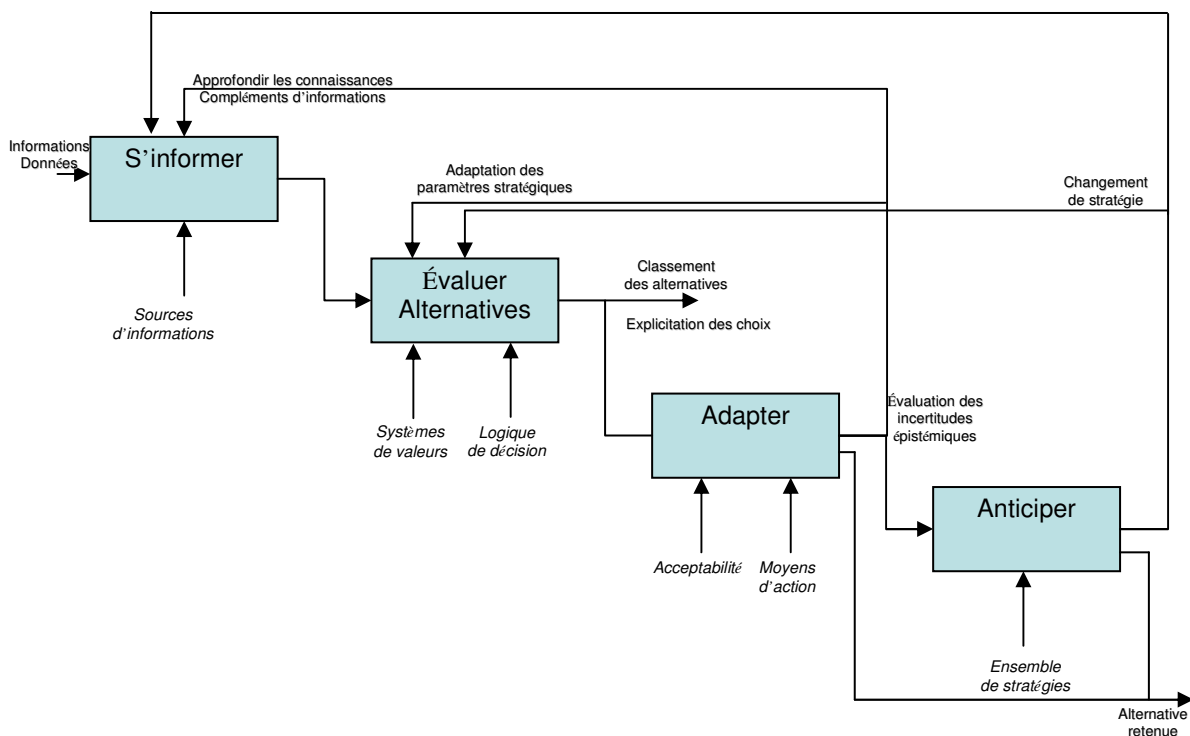


Figure I-14 : Conception du SIADG et fiabilité humaine

Chapitre II.

Le rôle de l'information dans le contrôle d'un processus décisionnel

1. Introduction	46
2. Le modèle S.T.I de H.A. Simon	47
2.1. Rationalité substantive versus rationalité procédurale	47
2.2. Acquisition et traitement de l'information	48
2.3. Synthèse	50
3. L'interprétation cybernétique du modèle S.T.I	51
3.1. Boucles cognitives et dynamique du processus décisionnel	51
3.2. Boucle de rétroaction et contrôle de la dynamique du processus décisionnel	52
3.2.1. La phase d'information	52
3.2.2. La phase de représentation	53
3.2.3. La phase de sélection.....	60
3.2.4. La phase de révision.....	64
3.2.5. Synthèse	66
4. Conclusion.....	68

1. Introduction

Ce chapitre présente un modèle de processus décisionnel multi acteurs et multicritère qui confère à notre SIADG les caractéristiques nécessaires pour répondre aux spécifications requises dans le chapitre précédent.

La décision en organisation est une bonne illustration des processus décisionnels multi acteurs et multicritères. Notre modèle s'inspire à l'origine du modèle S.T.I (Système de Traitement d'Information) de H.A. Simon qui donnait à l'information un rôle central dans le processus décisionnel d'une organisation. Ici, ce modèle des Sciences Humaines et Sociales est interprété dans un formalisme cybernétique où la « vanne de l'information » est vue comme l'organe de commande de la décision. Le contrôle de l'apprentissage—de la capitalisation des connaissances utiles à la décision—induit une dynamique forcée du processus décisionnel. La recherche et le compte-rendu mécanisés des éléments décisionnels manipulés dans cette boucle de régulation reposent sur des techniques de « text-mining » ou fouille de textes. Cette interprétation prétend clairement étendre l'automatisation cognitive de l'activité humaine à des tâches qui par définition sont d'un haut niveau décisionnel.

Nous montrons alors que cette analyse conceptuelle de la décision de groupe conclut à des traitements de l'information attendus du SIADG parfaitement cohérents avec ceux préconisés dans l'approche pragmatique du chapitre précédent : un SIADG doit être tout à la fois un support pour l'acquisition, le partage et le traitement d'informations utiles à l'action, pour le diagnostic de la situation décisionnelle au regard des connaissances disponibles, pour l'évaluation, la comparaison et la légitimation de solutions potentielles, le suivi et le contrôle temporels des incertitudes afférentes à la situation décisionnelle et à l'état des connaissances utiles à l'action.

Nous sommes néanmoins conscients des limites et dangers de cette vision un tant soit peu utopique ou tout du moins très idéaliste de l'informatique décisionnelle poussée à outrance dans l'organisation. Nous ne reviendrons pas sur les querelles que le débat sur l'automatisation cognitive a pu engendrer au milieu des années 80. Mais, afin qu'il n'y ait pas d'ambiguïté quant à notre position sur ce concept, nous précisons que notre démarche s'inscrit dans la perspective suivante : l'automatisation cognitive exacerbée ne doit être qu'un Graal pour l'automaticien et l'informaticien. C'est pour cette raison que le premier chapitre a d'abord proposé une analyse des besoins basée sur le retour d'expérience des SHS et une conception de SIADG guidée par l'analyse de dysfonctionnements sous l'angle de la fiabilité humaine pour donner une lecture rassurante, non pas formelle mais d'usage, de notre approche. Nous ne parlons pas de système de décision, mais bien de système interactif d'aide à la décision : même si l'on essaie dans cette étude de suppléer l'être humain autant qu'il se peut dans le traitement de l'information utile à la décision, son rôle d'arbitre et d'intervenant ultime dans la décision ne saurait être remis en cause.

Le chapitre s'organise donc de la façon suivante. Dans un premier temps, nous redonnons les caractéristiques fondamentales du modèle de Simon pour la décision en organisation. La section suivante synthétise l'interprétation cybernétique qu'Akharraz en a proposée dans sa thèse [Akharraz, 2004; Montmain *et al.*, 2006; Plantié *et al.*, 2002]. Nous étendons ensuite ce travail à un panel de stratégies décisionnelles plus large que celui proposé par Akharraz³. Cette extension mathématique permet la généralisation formelle des concepts de légitimation et de risque associés à une décision et l'extension des fonctionnalités du SIADG. Dans les

³ A.Akharraz a développé son modèle dans le seul cas où la stratégie multicritère est modélisable par une intégrale de Choquet.

chapitres suivants du manuscrit, cette interprétation cybernétique constitue la base de notre modèle, nous pourrions alors nous intéresser plus spécifiquement à la « construction » automatisée d'informations utiles à la décision en nous appuyant sur des techniques de fouille de textes, puis proposer les spécifications d'un actionneur intelligent à intégrer dans le modèle cybernétique de Akharraz : une vanne « intelligente » de l'information qui repose sur une indexation automatisée de l'information dédiée à la décision multicritère. Cette étape fera sauter l'ultime verrou—la sélection automatique de l'information la plus pertinente pour réduire dynamiquement l'incertitude épistémique afférente à une situation décisionnelle— et confèrera au modèle initial de Akharraz un degré d'automatisation supérieur pour l'ensemble des processus cognitifs impliqués dans la décision de groupe.

2. Le modèle S.T.I⁴ de H.A. Simon

2.1. Rationalité substantive versus rationalité procédurale

Pour H.A.Simon, l'émergence de la décision comme domaine d'étude scientifique remonte aux années 1943 et 1948, moment où se créent trois courants de recherche parallèle.

- Les théories mathématiques de la décision centrées autour de la théorie des jeux, des théories normatives avec J.von Neumann, O.Morgensten et Savage. Cette approche est surtout développée en économie mathématique [Neumann *et al.*, 1944; R. Yager, 1979; Savage, 1972] ;
- La cybernétique comme science de la communication et de la commande dans les systèmes naturels et artificiels créée par N.Wiener en 1948 et vite associée aux problématiques de la recherche opérationnelle ;
- H.A.Simon à travers sa thèse publiée en 1947 "A study of the decision making process in administrative organization" [Simon, 1947].

H.A.Simon remet alors en cause le modèle formel et mathématique de la décision encore adopté par la quasi-totalité des économistes classiques, des chercheurs opérationnels et des spécialistes de gestion de production. Cette vision part du postulat selon lequel choisir de façon rationnelle, c'est adopter la démarche déductive et analytique de la logique formelle des propositions. Face à un problème donné, il faut inventorier une liste complète des actions possibles, sélectionner par des procédures mathématiques la meilleure solution, c'est-à-dire celle qui satisfasse au mieux un critère d'utilité et enfin agir en appliquant cet optimum calculé. La référence idéale sera celle de l'abeille qui construit une cellule de cire selon une forme géométrique invariable : un dodécaèdre rhomboédrique dont l'angle de la base est de 136.4 degrés. Les calculs ne montrent-ils pas que cet angle correspond à la forme optimale qui minimise la quantité de cire nécessaire pour stocker un volume donné de miel ! La nature parvient donc à programmer en son sein une procédure totalement rationnelle, ou optimisatrice.

Toute notre cellule cartésienne de la décision est fondée sur ce principe "faire aussi bien que la nature" laquelle serait logique, rationnelle et déductive. La recherche opérationnelle, la théorie des jeux, les calculs statistiques procèdent tous de cette démarche idéale. Cependant dans la réalité quotidienne, on utilise rarement de telles procédures. La réflexion de Simon a consisté à remettre en cause cette démarche en proposant une conception alternative des processus de décision dans les organisations.

D'abord il faut remettre en cause l'idée selon laquelle la décision est une réponse précise à un problème donné, prédéfini. La décision est un processus où problème et réponse se construisent en même temps. Simon ne parle pas de la décision mais de "decision making

⁴ Système de Traitement de l'Information

process". La première phase de la décision consiste à identifier la nature de la question à traiter. Parfois la décision conduira à changer les données du problème.

Le second apport de Simon concerne la conception et l'évaluation des solutions alternatives possibles. L'exemple qu'il cite est celui des joueurs d'échecs. Aux échecs, il est impossible à un cerveau humain, ni même à un ordinateur d'envisager toutes les combinaisons possibles. Les joueurs utilisent donc des heuristiques, stratégies habiles, raisonnées, et non pas une procédure algorithmique qui consisterait à passer en revue la liste complète des solutions.

Pour la plupart des problèmes de la vie quotidienne, nous mettons en œuvre de telles heuristiques, c'est-à-dire des raisonnements plausibles mais non certains (des inférences plutôt que des déductions). Parfois même, il arrive qu'une démarche de la pensée viole les principes de la logique déductive formelle et soit pourtant efficace. L'argumentation de fond de Simon est donc d'abandonner ce concept de stratégie optimale. Premièrement, il est très rare que nous soyons dans une telle situation avec un seul critère. Et à supposer que l'on soit dans une telle situation, on a rarement la capacité cognitive de traiter les milliards d'informations et de solutions possibles.

Simon a longtemps médité sur le concept de rationalité avant de proposer un diagnostic. Il a pendant très longtemps utilisé le concept de "Bound rationality" traduit habituellement par rationalité limitée et qui renvoie à l'idée d'une connaissance imparfaite, ou bornée que le sujet a de son environnement. Dans les années 70, Simon préfère opposer le concept de rationalité substantive, qui est le raisonnement formel, analytique et déductif à la rationalité procédurale qui correspond à la façon dont l'être humain conduit fort correctement sa raison en reliant sans cesse ses intentions et ses perceptions du contexte dans lequel il raisonne.

Il y a une sorte de renversement. Ce n'est pas la rationalité procédurale qui est limitée mais plutôt la rationalité déductive. La rationalité déductive ou substantive ne correspond qu'à une petite partie des formes possibles du raisonnement humain.

Il s'agit d'élargir l'éventail des formes de raisonnement possibles sans se limiter au seul raisonnement déductif. L'étude de la rationalité procédurale ouvre un important champ d'études. Il est expliqué aujourd'hui par les théories de l'argumentation, la nouvelle rhétorique, ou encore certains courants de la psychologie cognitive. On n'est plus dans le cadre de la déduction formelle mais dans celui de la capacité de l'esprit à produire des solutions rusées, malicieuses pour résoudre les problèmes.

Avec l'essor actuel des sciences de la cognition se développent aujourd'hui de nouvelles réflexions épistémologiques des sciences de la complexité, jetant ainsi les bases des "Nouvelles sciences de la décision" [Simon, 1991].

2.2. Acquisition et traitement de l'information

Pour Simon, les modèles mathématiques utilisés en particulier en sciences économiques présente une conception qui se caractérise par quatre traits :

- ignorance des phases d'intelligence et de conception, dont le résultat est supposé faire partie des croyances du décideur dès le départ (ensemble des possibles, forme du problème) ; la phase de sélection est jugée suffisante pour modéliser la décision ;
- ignorance de la phase de bilan, même pour un simple bouclage sur la phase de sélection, puisque celle-ci est seule modélisée, et «linéarité», en conséquence, du raisonnement ;
- recours à une modalité de rationalité forte puisque optimisant le critère (unique) retenu («utilité» pour le consommateur, profit pour le producteur) ;
- réduction de la complexité du système de pilotage par la sorte d'hypostase du décideur.

Une telle conception de la décision, même en se plaçant dans l'hypothèse d'avenir certain, semble aujourd'hui difficilement justifiable sauf dans des cas bien spécifiques. En effet, on

peut faire valoir que ce modèle «traditionnel» de la décision individuelle peut être utilisé à des fins prescriptives dans les domaines purement techniques, et en particulier pour la gestion à court terme d'un parc d'équipements ou d'un chantier. Les phases d'intelligence et de conception sont alors souvent triviales, et les quatre hypothèses implicites relevées plus haut peuvent être admises au moins en première approximation. C'est sur ce créneau spécifique de la gestion de la production et de la gestion financière que fleurissent, surtout depuis les années cinquante, les nombreux modèles de la recherche opérationnelle et du calcul économique dont les apports ne sauraient être discutés.

Il en va tout autrement dans le cadre de la décision organisationnelle où le vocable «système de traitement de l'information» (S.T.I) permet de désigner commodément la lignée des modèles issus de la pensée de H.A.Simon. Cet ensemble de propositions se fonde sur quatre remarques critiques vis-à-vis des applications un peu hâtives du modèle de conception précédent.

Les décisions ne sont pas toutes de même nature ni de même niveau. Certaines sont telles qu'un automate, correctement informé et programmé, est en mesure de les prendre aussi bien—voire de façon plus fine—que l'homme lui-même, dont la capacité purement combinatoire est singulièrement limitée. On les appellera décisions «programmables». Mais, dès lors notamment que des ensembles humains ou des interfaces hommes-machines sont concernés, les décisions sont de l'ordre du «non-programmable». Gommer les phases d'intelligence et de conception du processus de décision est alors une réduction inacceptable.

De façon plus générale encore, une conclusion identique s'impose lorsqu'on remarque que les ensembles de possibles ne sont jamais faciles à explorer. Or le meilleur général n'est sans doute pas celui qui réalise le dosage optimal entre aile droite et aile gauche sur un terrain donné, mais celui qui sait imaginer une possibilité stratégique de plus, permettant d'accroître sa palette de tactiques. Il en va de même du chef d'entreprise, du préfet, ou du couple élevant ses enfants. Les phases d'intelligence et de conception sont véritablement les phases clés du processus de décision.

On pourrait se donner pour programme de recherche de tenter de formaliser la démarche décisionnelle à travers toutes les phases évoquées. Mais on buterait très vite sur des difficultés de calcul inextricables. Une autre voie consiste à essayer plutôt de représenter—et d'améliorer—la façon dont les hommes utilisent leurs capacités de raisonnement et de traitement des informations. Au lieu de chercher à désigner une décision ambitionnant d'être optimale, il est plus modeste dans les objectifs, mais peut-être plus efficace pour le résultat, de chercher à user d'une procédure de traitement de l'information et de raisonnement plus satisfaisante. La rationalité «limitée» ou «procédurale» vient ainsi se substituer à la rationalité optimisante et «substantive».

Enfin, le fait qu'on se place dans le cadre d'une organisation comportant évidemment des services, des départements, etc., non totalement intégrés introduit une contradiction avec la nécessité, pour optimiser de façon significative, de traiter simultanément dans le temps et partout dans l'organisation l'ensemble des informations disponibles. Les problèmes de délégation et de coordination se heurtent, là encore, à d'inextricables difficultés. Aussi, dans le modèle S.T.I., les différentes phases de la décision ne se présentent-elles pas de façon linéaire, mais en boucles. De nombreuses itérations sont nécessaires, au vu de la faible capacité de traitement de l'information des hommes et de la complexité des problèmes de décision, avant qu'un terme puisse être apporté au processus de décision. Davantage encore, chacune des phases engendre des sous-problèmes qui, à leur tour, appellent des phases d'intelligence, de conception, de sélection et de bilan. Les phases sont ainsi des «engrenages d'engrenages».

Par ailleurs, le processus de décision est mis en œuvre, dans une organisation, par un système de décision, lui-même complexe et dont l'hypothèse du décideur peut donner une

représentation mutilante et être une source d'erreurs de décision si elle conduit à ignorer les difficultés de communication et de coordination.

Le modèle S.T.I insiste en définitive sur les aspects cognitifs de la décision, l'acquisition et le traitement de l'information apparaissant comme plus importants pour prendre une «bonne» décision que la recherche fine illusoire d'une décision en apparence «la meilleure». La notion d'organisation est centrale dans la justification de ce modèle.

Il y a encore une dizaine d'années, la mise en œuvre du modèle S.T.I de Simon aurait pu être résumée comme suit. Les techniques de mise en œuvre du modèle S.T.I sont principalement celles de l'informatique, et notamment des Systèmes Interactifs d'Aide à la Décision (SIAD). Ceux-ci constituent une interface entre un ordinateur et une ou plusieurs personnes (SIAD individuel, ou SIADI ; et SIAD de groupe ou SIADG). Sur la «machine», on implante soit un algorithme, correspondant à la phase de sélection, soit un système heuristique pouvant aider à la créativité (aide à la phase de conception) ou, de façon plus ambitieuse, utilisant une base de connaissances (système expert). Dans tous les cas de figure, le système implanté en machine est considéré comme n'«apportant» pas «la» solution, mais comme jouant le rôle d'amplificateur cognitif, mettant l'accent sur une ou plusieurs phases du processus de décision.

Dans ce qui suit nous montrons que les avancées techniques et technologiques des systèmes d'information de ces dernières années, couplées à un véritable accompagnement mathématique et informatique des processus cognitifs collectifs semblent définitivement donner raison à H.A.Simon et nous proposons une mise en œuvre pratique du modèle S.T.I pour la réalisation d'un SIAD collectif basé sur le couplage d'un système de traitement d'informations dynamique et d'un système d'évaluation multicritère.

2.3. Synthèse

Les fondements de la théorie de Simon reposent sur les remarques suivantes :

- Le décideur ne possède pas en fait une connaissance totale de la situation, d'où le terme de « rationalité limitée » cher à H.A. Simon et ses limitations dans la connaissance des faits et hypothèses proviennent principalement des contraintes de l'organisation qui sélectionne ou favorise tel ou tel scénario en fonction de ses intérêts. Dans sa théorie du surcode [Sfez, 1992], le politologue L. Sfez montre que derrière l'image trompeuse d'une décision consciente et unifiée, il y a en fait une multiplicité de rationalités différentes qui s'imbriquent, se superposent, se confrontent. La décision optimale apparaît dénuée de sens dans une évaluation multipoints de vue et multi acteurs ;
- L'incapacité (en tout cas la capacité limitée) de l'homme à traiter l'ensemble des flux d'informations imprécises, incertaines, incomplètes et contradictoires nécessaires à la décision semble montrer que la solution pour une aide à la décision efficace relève, dans ce cas, des systèmes de traitement de l'information. L'acquisition et le traitement de l'information apparaissant comme plus importants pour prendre une «bonne» décision que la recherche fine et illusoire d'une décision en apparence «la meilleure» qui se fondera nécessairement sur des hypothèses et des simplifications de représentations ;
- Selon le modèle de la rationalité limitée, le décideur est naturellement tenté de s'orienter vers une approche monocritère, occultant la prise en compte de la complexité de la réalité. L'approche multicritère de l'aide à la décision permet de pallier cette restriction en augmentant le niveau de réalisme et de lisibilité donné au décideur [Roy *et al.*, 1993] [Slowinski, 1998] ;
- Les différentes phases de la décision ne se présentent pas de façon linéaire, mais en boucles. De nombreuses itérations sont nécessaires, au vu de la capacité cognitive des hommes et de la complexité des problèmes de décision, avant qu'un terme ne puisse être

apporté au processus de décision. Phases d'intelligence (Intelligence), de conception (Design), de sélection (Choice) et de bilan (Review) se succèdent sans logique chronologique préétablie possible (modèle IDCR de H.A.Simon) (voir détails section suivante).

Ces prescriptions sont bien cohérentes avec les résultats du chapitre précédent concernant les apports attendus d'un SIADG. On y retrouve, de façon centrale, le principe de la rationalité limitée (incertitude épistémique) qui recouvre bon nombre des biais cognitifs liés à une mauvaise perception de la situation ou à une mauvaise interprétation de celle-ci, à la complexité d'une évaluation multidimensionnelle de la situation par un groupe d'acteurs n'ayant pas nécessairement les mêmes intérêts et valeurs, à une mauvaise appréhension du temps ou de la causalité des événements, à une vision linéaire du processus de décision, à une mauvaise appréciation de la nature de la situation ou de la réaction de l'environnement, et par suite l'application d'une stratégie d'action inappropriée, etc.

3. L'interprétation cybernétique du modèle S.T.I

3.1. Boucles cognitives et dynamique du processus décisionnel

H.A.Simon [Simon, 1997] distingue quatre phases dans le processus de décision [Lévine *et al.*, 1989], [Turban, 1993] : la recherche d'information, la conception, le choix et la révision.

- L'information ou le renseignement (*intelligence*)

Il s'agit d'identifier les objectifs ou buts du décideur, c'est-à-dire de définir le problème à résoudre. Pour cela, il est nécessaire de rechercher les informations pertinentes en fonction des questions que se pose le décideur. L'acquisition d'informations pertinentes pendant cette phase peut se poser elle-même en termes de décision. En effet, ces informations pertinentes sont à l'origine du processus de décision et leur choix est crucial. Elles influencent fortement les autres phases puisque tous les choix suivants en découlent. L'apport technique de cette thèse est principalement consacré à cette phase ;

- La conception (*design*)

Cette phase comprend la génération, le développement et l'analyse des différentes suites possibles d'actions. Le décideur construit des solutions, imagine des scénarios, ce qui peut l'amener à rechercher de l'information supplémentaire. Pour cela, il va être nécessaire de choisir un ou plusieurs modèles de décision en fonction de la complexité du problème à traiter. Pour le ou les modèles choisi(s), il faut déterminer les variables de décision, la sélection des principes de choix (critère d'évaluation), ainsi que les relations mathématiques ou symboliques ou qualitatives entre ces variables et construire les différentes alternatives ;

- Le choix (*choice*)

Pendant cette phase, le décideur choisit entre les différentes suites d'actions (solutions) qu'il a été capable de construire et d'identifier pendant la phase précédente. Il faut déterminer les critères d'évaluation des différentes solutions envisageables et étudier ou mesurer les conséquences de chaque alternative. L'évaluation des alternatives et le choix final dépendent du type de critères utilisés. Par exemple, trouver la meilleure solution, une solution assez bonne ou satisfaisante, prendre des risques ou non, minimiser des regrets ou un manque à gagner ou maximiser des gains, etc. Cette phase inclut la recherche, l'évaluation et la recommandation d'une solution appropriée au modèle.

- La révision (*review*)

Cette phase est souvent négligée, bien qu'essentielle dans le modèle de Simon, en particulier dans le cas où la décision s'intégrerait dans un processus dynamique. De nouvelles informations pertinentes peuvent influencer tel ou tel choix, voir le modifier complètement.

Ce processus de décision IDCR (*Intelligence-Design-Choice-Review*) est rarement séquentiel, il peut y avoir des « retours en arrière » ou boucles, c'est-à-dire qu'à chaque phase, on peut être amené par exemple à générer une nouvelle alternative ou encore à rechercher de nouvelles informations, puis ensuite modifier le ou les modèles choisis, etc. (Figure II.1). La présence de ces boucles pendant les processus de décision dépend du niveau de *structuration* du problème de décision (contraintes). H.A. Simon parle de causalités enchevêtrées.

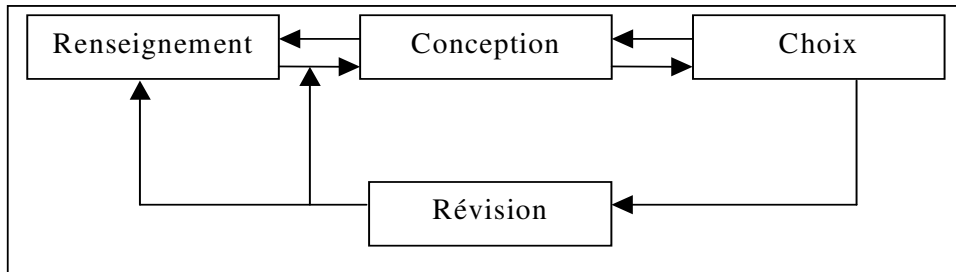


Figure II.1: boucles cognitives dans le modèle IDCR de Simon

3.2. Boucle de rétroaction et contrôle de la dynamique du processus décisionnel

3.2.1. La phase d'information

Supposons donc qu'une organisation se soit dotée d'un intranet qui lui permette de partager différentes bases de connaissances réputées être utiles à la décision que l'organisation souhaite prendre. L'intranet et ses bases de connaissances constituent un système de gestion dynamique des connaissances (SGDC)⁵ qui permet de contrôler l'évolution du corpus de connaissances partagé et produit par le collectif sur le problème à résoudre. Un processus cognitif d'apprentissage collectif s'opère sur cette mémoire partagée en expansion et acquiert alors une dimension dynamique. L'échange de points de vue « métier », de systèmes de valeurs, d'intérêts ou de cultures, trouve ainsi une structure dynamique qui favorise la réactivité, la délibération et l'argumentation et qui, in fine, peut modifier la dynamique du processus d'apprentissage collectif. L'apprentissage d'un savoir et d'une mémoire collectifs est le processus cognitif supporté par le SGDC : cet élément technique est une réponse possible à la phase d'information du modèle de Simon. Dans ce qui suit, il est montré en quoi il constitue également un support naturel pour l'évaluation des alternatives potentielles, la sélection et la justification d'une alternative.

Conformément à la pensée de Simon, le principe est d'améliorer la façon dont les hommes utilisent leurs capacités de raisonnement et de traitement de l'information. On ne cherche pas une solution nécessairement optimale mais satisfaisante au sens d'« argumentable », justifiable ; plutôt qu'un compromis, on cherche un consensus. L'utilisation du SGDC donne une certaine transparence à la décision, une meilleure lisibilité de la logique décisionnelle : en proposant un cadre et une instrumentation explicites pour les processus cognitifs collectifs d'apprentissage, d'argumentation, de délibération, de concertation et de sélection, elle permet (tant dans l'aspect gestion des connaissances que dans la décision elle-même) de favoriser l'échange hiérarchique et transversal de données, d'informations, de connaissances et de décisions à travers une modélisation systémique et mathématique des liens de subordination et de coordination qui lient les acteurs d'un processus de décision organisationnel.

⁵ Nous reviendrons dans le chapitre suivant sur la notion de gestion dynamique des connaissances.

3.2.2. La phase de représentation

Parmi les bases de connaissance gérées par le SGDC, on distingue plus particulièrement les connaissances qui expriment un jugement de valeur de leur auteur relativement à une alternative candidate et une perspective d'analyse. Ces connaissances qui expriment le positionnement d'un acteur sont des éléments de rhétorique directement utiles à la décision ; ils sont l'interprétation d'informations, de données ou de savoirs dans le cadre des objectifs supportant le processus de décision. Ce sont des connaissances actionnables (CAs). Nous reviendrons plus en détail sur cette notion dans le chapitre suivant, mais pour l'instant nous nous contenterons d'assimiler une CA à un score, expression quantitative du jugement de valeur qu'elle exprime.

Par ailleurs, lorsque l'évaluation globale d'un objectif est complexe, il est nécessaire de décomposer l'objectif à atteindre en en structurant l'ensemble des critères d'évaluation. L'approche multicritère de l'aide à la décision augmente le niveau de réalisme et de lisibilité donné au décideur [Pomerol *et al.*, 1993]. Construire un modèle prenant explicitement appui sur plusieurs critères, traduit et formalise, un mode de raisonnement intuitif et naturel face à un problème de décision qui consiste à analyser séparément chaque conséquence [Roy, 1985].

L'espace d'évaluation

Les principes de partage des connaissances et de réactivité aux opinions exprimées par des acteurs de métier, de culture et/ou de fonction très différentes orientent justement notre SIADG vers des outils de décision multicritère. L'aspect cognitif de l'approche implique que l'on doit porter autant d'attention à la sémantique des opérateurs mathématiques d'agrégation ou de fusion multicritère qu'à leurs propriétés mathématiques [Dubois, 1983; Dubois *et al.*, 1985; J-L. Koning, 1990]. Avant de donner quelques éléments mathématiques dans la sous section suivante sur la fusion et l'agrégation multicritère, ce paragraphe est destinée à expliquer comment les CAs sont utilisées dans le processus de décision [Montmain *et al.*, 2002].

Pour évaluer les alternatives que le processus de décision fait émerger, il est nécessaire de définir un espace de représentation : une grille d'évaluation des alternatives potentielles examinées selon un ensemble de critères (Figure II.2). Il est possible de donner une structure hiérarchique à l'ensemble des critères d'évaluation utilisés.

Exemple : La Figure II.2 propose l'évaluation de 12 films examinés selon 4 critères. Un code de couleur a été préféré à une échelle numérique pour des raisons d'interface homme/machine. Le vert foncé exprime une pleine satisfaction, le rouge vif une déception totale. On peut imaginer sur cet exemple que les CAs sont des critiques de cinéphiles recueillies sur le web.

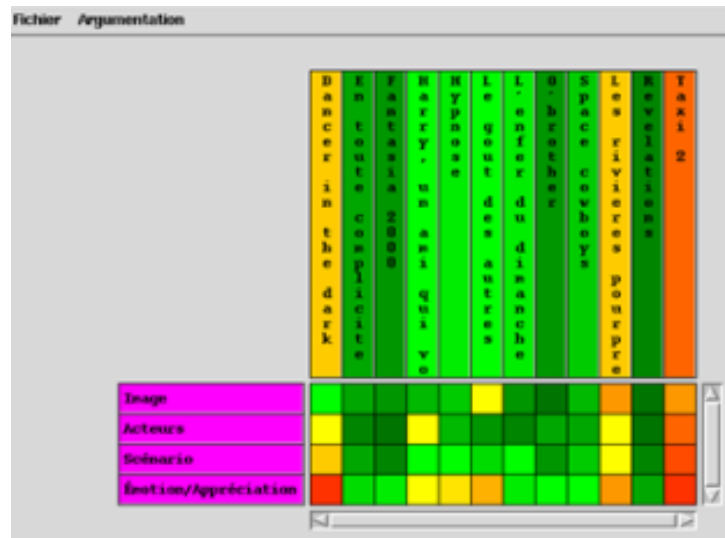


Figure II.2: Un exemple de grilles d'évaluation par les CAs

L'idée de base de cette évaluation multicritère du projet est de se référer à la théorie du choix collectif en identifiant les critères à des *votants* et les concepts envisagés pour le projet à des *candidats à élire*. On cartographie alors les jugements de valeurs ou scores partiels associés aux CAs dans cette grille, ce qui permet de prendre en compte les intensités de préférence des votants : chaque *vote* est une évaluation d'une alternative *i* vis-à-vis d'un critère *j*. L'évaluation globale d'une alternative—son *score*—correspond à l'agrégation des appréciations partielles qu'elle a obtenues selon chaque critère et son obtention est donc assimilée à une procédure d'élection par les critères. Qu'un *électeur* soit un critère et non un individu permet d'éviter l'écueil de la multiplicité et de la variabilité des acteurs du projet. La CA—ou plutôt le score qu'elle porte—est vue comme *la voix ou le bulletin de vote* utilisé par un critère vis-à-vis d'une alternative. La CA cartographiée, participe par le jugement de valeur qui lui est associé à l'évaluation d'une alternative candidate dès lors qu'elle est incorporée dans la base de connaissances du SGDC—*l'urne* (Figure II.3).

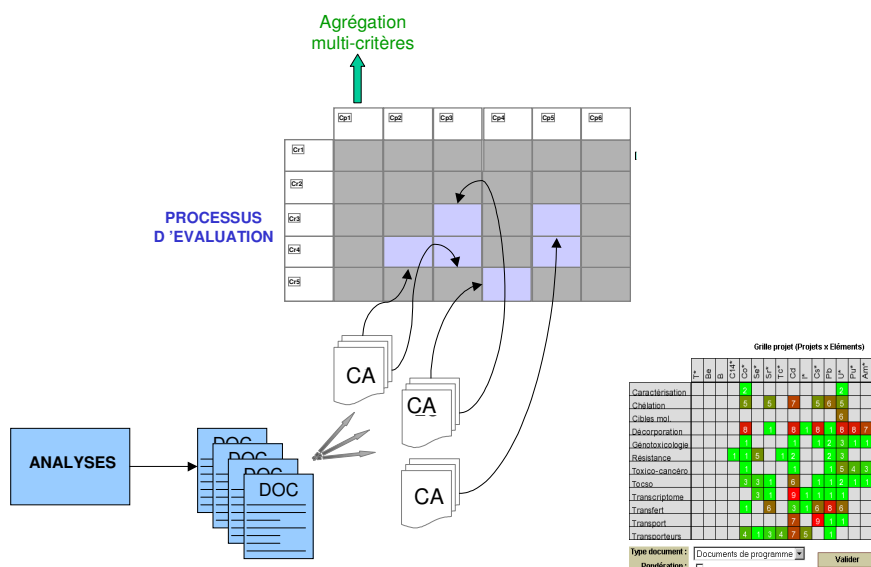


Figure II.3: Le processus de vote

Parce qu'un système de notation numérique dans l'évaluation d'une *alternative i* selon un *critère j* nous est apparu difficile à mettre en œuvre, nous avons opté pour une palette de couleurs plutôt qu'une notation purement quantitative : un rouge vif indique l'incompatibilité de l'alternative *i* avec le *critère j* alors qu'un vert prononcé exprime une adéquation parfaite.

La palette de couleurs dépend bien sûr de la granularité de l'échelle d'appréciation que l'on souhaite ou que l'on est capable d'exprimer.

Une grille de ce type permet donc d'évaluer quantitativement à chaque *instant* chaque *alternative* i selon un *critère* j . L'opération peut bien sûr être répétée selon une fréquence d'échantillonnage adaptée à la dynamique du processus d'apprentissage ou de façon événementielle lorsqu'une CA ou des CAs sont introduites dans la base. On a ainsi la possibilité de suivre dans le temps les évaluations partielles de chaque alternative relativement à chaque critère, mais aussi, de manière plus synthétique, l'évaluation globale d'une alternative, synthèse au sens d'un opérateur d'agrégation des scores partiels obtenus selon chaque critère.

Exemple : La [Figure II.4](#) reprend l'exemple de l'évaluation des 12 films et propose l'évolution de leur évaluation globale dans le temps jusqu'au temps courant « 27 avril 2003 à 04h00 » avec une période d'échantillonnage de 4h.

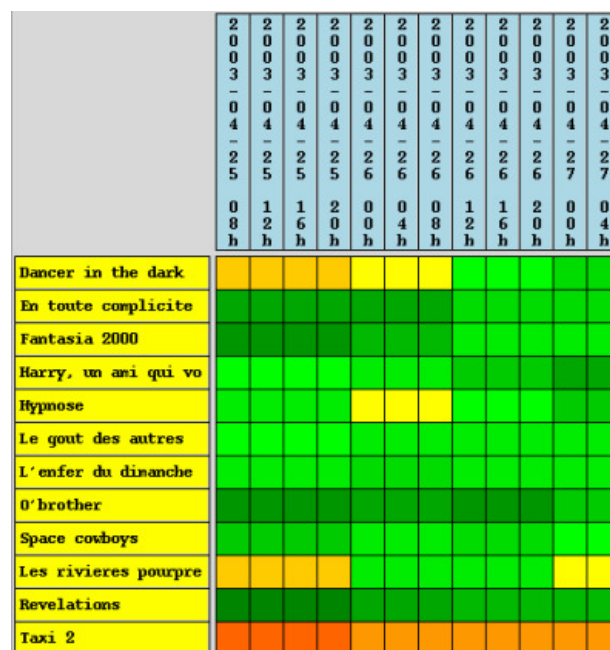


Figure II.4: Suivi de l'évaluation globale des solutions envisagées dans le temps

Agrégation multicritère

Cadre et actants de la procédure de vote définis, il reste à préciser le *processus* ou *mode d'élection*, autrement dit la façon dont sont composés et agrégés les scores des CAs. Cette agrégation porte dans un premier temps sur la combinaison des jugements de valeurs des CAs se rapportant à une seule et même case de la grille d'évaluation (critère i ; alternative j). L'évaluation globale d'une alternative nécessite ensuite l'agrégation de toutes ses appréciations partielles selon chaque critère. Il existe différents opérateurs mathématiques qui permettent de modéliser cette procédure d'agrégation, les différentes propriétés qu'une *procédure de vote* peut satisfaire.

Nous rappelons dans un premier temps la terminologie et les conditions minimales nécessaires dans cette section pour définir un problème d'agrégation multicritère cardinal.

S : ensemble d'alternatives (solutions, projets, actions, options...) parmi lesquelles le décideur doit choisir ;

C : ensemble des critères permettant d'évaluer les alternatives. $X=[0,1]$: l'ensemble des degrés de satisfaction (ou évaluations partielles).

Nous considérons l'ensemble fini des p critères $C = \{c_1, c_2, \dots, c_p\}$, $P(C)$ l'ensemble des parties de C , $x = (x_1, x_2, \dots, x_p)$ représente les valeurs numériques de satisfaction des critères, de sorte que, x_i est le degré de satisfaction du critère c_i .

On considère aussi l'ensemble des alternatives $S = \{s^1, s^2, \dots, s^n\}$. A chaque alternative $s^k \in S$ est associé un profil $s^k = (x_1^k, x_2^k, \dots, x_p^k)$ avec x_i^k l'évaluation partielle de s^k selon le critère c_i .

Pour des raisons de simplification des notations, les valeurs des scores sont exprimées dans l'intervalle $[0,1]$, 1 correspond à la satisfaction complète d'un critère, 0 exprime la non satisfaction complète.

L'un des problèmes d'aide à la décision multicritère est d'aboutir à un classement des alternatives de la meilleure à la moins bonne, au vu de l'ensemble critères, ou tout simplement à trouver la meilleure alternative. Dans l'approche cardinale, l'objectif est de construire une fonction $h: [0,1]^p \rightarrow [0,1]$ telle que, pour chaque alternative s^k : on a $h(s^k) = h(x_1^k, \dots, x_p^k)$, où $h(s^k)$ représente l'évaluation globale de l'alternative s^k relativement à tous les critères, h est l'opérateur d'agrégation à déterminer. Les alternatives s^k peuvent donc ensuite être classées relativement à leur score $h(s^k)$.

Les conditions nécessaires sur l'opérateur h sont les suivantes :

- h est continu ;
- $h(0,0,\dots,0) = 0$ et $h(1,1,\dots,1) = 1$;
- $\forall (u_i, v_i) \in [0,1]^2$, si $u_i \geq v_i$ alors $h(u_1, \dots, u_p) \geq h(v_1, \dots, v_p)$.

Groupes homogènes

Dans ce paragraphe, on considère que tous les critères ont la même importance (participent à même hauteur dans notre procédure de vote).

Pour une opération exprimant la satisfaction simultanée des objectifs un axiome naturel est : $\forall (u_1, \dots, u_p), h(u_1, \dots, u_p) \leq \min(u_1, \dots, u_p)$, c'est-à-dire que l'évaluation globale d'une action ne peut être meilleure que la plus mauvaise des évaluations partielles. On appellera ces opérateurs des conjonctions. Les principales conjonctions associatives sont :

$$\min(u_1, \dots, u_p), \prod_{i=1}^p u_i \text{ et } \max\left(0, \sum_{i=1}^p u_i - p + 1\right) \quad \text{II-1}$$

Un opérateur exprimant la redondance de deux objectifs satisfera un axiome dual du précédent :

$\forall (u_1, \dots, u_p), \max(u_1, \dots, u_p) \leq h(u_1, \dots, u_p)$ c'est-à-dire que c'est la meilleure des évaluations partielles qui conditionne l'évaluation globale. Ce sont les disjonctions dont une sous-classe importante correspond aux disjonctions associatives. Les plus utilisées sont :

$$\max(u_1, \dots, u_p), 1 - \prod_{i=1}^p (1 - u_i), \min\left(1, \sum_{i=1}^p u_i\right) \quad \text{II-2}$$

Une opération exprimant le compromis satisfera l'axiome suivant, complémentaire des deux précédents : $\forall (u_1, \dots, u_p), \min(u_1, \dots, u_p) \leq h(u_1, \dots, u_p) \leq \max(u_1, \dots, u_p)$.

Notre but est alors de procéder à l'identification de l'opérateur h . Cette opération peut être décomposée en deux étapes. On peut cerner le comportement du décideur en lui demandant de fournir une appréciation linguistique du niveau de compatibilité entre conséquence et objectif sur un ensemble réduit d'actions. Il est alors possible de se fixer rapidement sur le type d'agrégation à employer : conjonction, compromis ou disjonction. Une fois déterminée l'une

des trois attitudes fondamentales possibles du décideur, on peut utiliser des familles de fonctions paramétrées pour affiner la détermination de l'opérateur d'agrégation : le choix de h est ramené à un problème d'identification paramétrique.

Les opérateurs conjonctifs sont couverts, par exemple, par l'opération de Yager qui définit une famille paramétrée d'opérateurs :

$$Y_q(u_i) = 1 - \min \left(1, \left(\sum_{i=1}^p (1 - u_i)^q \right)^{1/q} \right) \text{ pour } q > 0 \quad \text{II-3}$$

$$\text{Par exemple : } \lim_{q \rightarrow \infty} Y_q(u_i) = \min(u_i) \text{ et } Y_1(u_i) = \max \left(0, \sum_{i=1}^p u_i - p + 1 \right)$$

La famille des conormes associées permet de couvrir l'ensemble des disjonctions :

$$CY_q(u_i) = \min \left(1, \left(\sum_{i=1}^p u_i^q \right)^{1/q} \right) \text{ pour } q > 0 \quad \text{II-4}$$

$$\text{Les compromis sont couverts par : } Ym_q(u_i) = \left(\frac{\sum_{i=1}^p u_i^q}{p} \right)^{1/q} \quad \text{II-5}$$

Des valeurs remarquables de q donnent la moyenne arithmétique $Ym_1(u_i)$ et la moyenne harmonique $Ym_{-1}(u_i)$; lorsque $q \rightarrow 0$, le compromis est la moyenne géométrique.

Les caractéristiques symétriques de h n'impliquent pas que l'agrégation soit symétrique. En effet, l'agrégation repose sur les ensembles flous décrivant les objectifs partiels, et des dissymétries peuvent facilement être introduites en manipulant les fonctions d'appartenance. Les fonctions d'agrégation ci-dessus peuvent être également généralisées avec des coefficients de pondération ou poids pour introduire les dissymétries de façon naturelle [Dubois, 1983; Dubois *et al.*, 1985].

Des techniques d'optimisation à partir de ces trois familles d'opérateurs permettent de construire des outils d'identification de l'opérateur h lorsque l'on dispose d'une base d'échantillons de tests suffisamment riche (en nombre et en diversité des appréciations partielles). L'identification très précise d'une telle opération est parfois impossible tant l'incertitude sur le paramètre q est grande. Les résultats d'identification doivent alors être accompagnés d'une sensibilité et d'un pouvoir de discrimination afin que l'on soit à même d'accorder un degré de validité au résultat lors d'une utilisation en simulation (recherche d'extrêma « très plats »). Avoir recours à ces tests de sensibilité n'est pas sans implication sur l'interprétation sémantique des coefficients des fonctions d'agrégation. Cela n'invalide pas l'approche pour autant, en effet :

- une identification très précise peut être jugée inutile ; seul un certain comportement de l'opérateur h peut être souhaité. Cela est à rapprocher du caractère arbitraire des valeurs d'appartenance, et de la traduction numérique des jugements émis par le décideur. Il vaut mieux dans ce cas là, se donner un catalogue restreint mais bien dispersé d'opérations simples et identifier h de façon *qualitative* en demandant au décideur de formuler à l'aide d'une échelle linguistique de jugement (totalement compatible, plutôt compatible,..., incompatible) un jugement global sur chaque action.
- si on trouve que plusieurs opérations sont des modèles possibles, on pourra malgré tout tenter de classer les éléments. On pourra introduire une incertitude sur le résultat de l'agrégation. Par exemple si on n'a pas pu (ou voulu) discriminer entre les opérations $h_1, h_2, \dots, h_t$ on peut poser :

$$\forall (u_1, \dots, u_p), h(u_1, \dots, u_p) \in \left[\min_{i=1, \dots, p} h_i(u_1, \dots, u_p), \max_{i=1, \dots, p} h_i(u_1, \dots, u_p) \right] \quad \text{II-6}$$

Groupes non homogènes

Très souvent on peut admettre que les différents critères selon lesquels on évalue les alternatives n'ont pas tous la même importance (tous les critères n'ont pas la même importance relative dans notre procédure de vote). Une façon de modéliser cet aspect est d'affecter un poids w_j à chaque critère. Alors le problème du choix collectif est de trouver des fonctions d'agrégation qui agissent sur l'ensemble des profils d'intensité de préférence pondérés.

La notion d'importance relative des critères (le sens accordé à la notion d'importance restant par ailleurs souvent *vague* dans la littérature) conduit naturellement à introduire des distributions de poids sur les critères dans le cadre de la théorie précédente :

$$\forall s^k \in S, h(s^k) = \sum_{i=1}^p p_i u_i^k \quad \text{II-7}$$

La notion de poids peut être transposée sur de nombreux opérateurs d'agrégation comme évoqués précédemment (en particulier les normes et conormes) dès lors que l'opération présente une structure additive sous-jacente [Dubois *et al.*, 1984]. Ainsi, lorsque h est de la forme $\Psi^{-1}(\Psi(.) + \Psi(.))$, une extension de l'équation précédente peut être proposée :

$$\bar{h}(u_1, \dots, u_n) = \Psi^{-1}\left(p \sum_{i=1}^p p_i \Psi(u_i)\right) \text{ où } \sum_{i=1}^p p_i = 1 \quad \text{II-8}$$

A titre d'exemple, pour le produit, en posant $\Psi = \text{Log}$, on obtient : $\bar{h}(s^k) = \prod_{i=1}^p (u_i^k)^{p p_i}$.

[Cholewa, 1985] décrit une série d'axiomes que les agrégations pondérées doivent vérifier et propose la moyenne arithmétique pondérée comme fonction d'agrégation typique satisfaisant ces axiomes. [Montero, 1989] justifie cette règle comme étant celle qui maximise la décision relative des critères d'évaluation, proportionnellement à l'importance de chaque critère.

Une interprétation du poids peut être la pertinence du critère. Ce niveau de pertinence peut agir en tant que contrainte sur les intensités de préférence qu'un critère peut exprimer. Par exemple, on peut souhaiter qu'un critère ne puisse pas entraîner l'acceptation ou le rejet complet des alternatives. Etant par essence antidémocratique, ce type de poids n'est pas adapté aux règles de type majorité, mais peut être utile pour limiter l'influence d'un critère dictateur ou d'un critère veto, en interdisant des classements extrêmes quand l'influence de certains critères doit être restreinte. Par exemple, si $\varphi(u_1^i \dots u_p^i) = \min_{j=1 \dots p} u_j^i$ alors n'importe quel critère a un droit de veto, mais ce droit devient limité avec l'agrégation pondérée $\varphi(u_1^i \dots u_p^i, w_1 \dots w_p) = \min_{j=1 \dots p} \max(w_j, u_j^i)$ où le critère j ne peut pas exprimer un niveau d'intensité plus petit que w_j , $\min_{j=1 \dots p} (w_j) = 0$.

La règle duale pour la dictature est : $\varphi(u_1^i \dots u_p^i, w_1 \dots w_p) = \max_{j=1 \dots p} \min(w_j, u_j^i)$, où $\max_{j=1 \dots p} (w_j) = 1$.

Les règles de majorité peuvent s'obtenir en utilisant la notion de moyenne ordonnée pondérée (OWA pour Ordered Weighted Average) introduite par [Yager, 1988]. L'idée est d'utiliser un ensemble de poids $w_1, \dots, w_p$ qui ne sont pas a priori affectés à des critères ; par exemple, les plus hauts poids sont affectés aux critères exprimant les plus hautes intensités de préférence pour une alternative donnée.

Considérons σ une permutation de $(1, 2, \dots, p)$ telle que $x_{\sigma(1)}^i \geq x_{\sigma(2)}^i \geq \dots \geq x_{\sigma(p)}^i$.

On considère alors la combinaison convexe : $\varphi(x_1^i, x_2^i, \dots, x_p^i, w_1, w_2, \dots, w_p) = \sum_{j=1}^p w_j x_{\sigma(j)}^i$

La moyenne arithmétique correspond au cas où $w_j = 1/p, \forall j$. $\varphi = \max$ si $w_1 = 1, w_j = 0, j \geq 2$; $\varphi = \min$ quand $w_p = 1, w_j = 0, j \leq p-1$. La règle de majorité "q parmi p" est obtenue avec $w_1 = w_2 = \dots = w_q = 1/q$ et $w_{q+1} = w_{q+2} = \dots = w_p = 0$. Le cas où les poids sont ordonnés tels que $w_1 \geq w_2 \geq \dots \geq w_p$ permet à cette règle de traiter les quantificateurs, puisqu'on suit le principe : plus un classement est élevé plus il est important.

D'autres modèles de majorité floues ont été proposés dans [Kacprzyk, 1987; Zadeh, 1983].

Les règles d'unanimité restreintes diffèrent des règles de majorité quantifiées en ne permettant pas de compensation entre les critères. C'est une altération de la règle du minimum selon laquelle une alternative est collectivement approuvée si tous les critères approuvent séparément cette alternative :

$\varphi(x_1^i, x_2^i, \dots, x_p^i, w_1, w_2, \dots, w_p) = \min_{j=1, p} \max(w_j, x_{\sigma(j)}^i)$, où les poids satisfont la condition $\min_{j=1, p} w_j = 0$. L'approbation exigée de q critères parmi p est obtenue en posant $w_1 = w_2 = \dots = w_q = 0$ et $w_{q+1} = w_{q+2} = \dots = w_p = 1$. Notons que l'approbation partielle ($x_j^i = q/p$) de tous les critères n'est pas considérée identique à l'approbation complète de q critères tandis que ces approbations sont équivalentes pour les opérations OWA. Notons que $\varphi = \max$ pour $w_1 = 0, w_j = 1, \forall j \geq 2$.

Le cas où $0 = w_1 \leq w_2 \leq \dots \leq w_p$, traite de la situation où le seuil q est mal défini comme lorsqu'on exige la satisfaction de "la plupart" des critères où la plupart est ici une proportion décrite par un ensemble flou Q sur $\{1, 2, \dots, n\}$ avec $\mu_Q(i) = w_i$ [J-L. Koning, 1990].

Interaction des critères

Nous avons jusqu'ici évoqué les relations de subordination entre les critères, il peut être nécessaire dans la pratique de considérer également les relations de coordination entre critères d'évaluation. En effet, dans le cas général, les critères sont liés, en plus des liens de subordination, par des relations transversales : ces relations transversales expriment des liens de coordination qui assurent la cohérence de scores d'un même niveau hiérarchique. Pour bénéficier de cette distinction sémantique des relations entre critères, il est nécessaire d'introduire de nouvelles fonctions d'agrégation comme l'intégrale de Choquet [Akharraz, 2004; Grabisch *et al.*, 2000; Montmain *et al.*, 2006]. Cet opérateur est une bonne illustration de l'agrégation de performances liées par des dépendances de subordination et de coordination : elle permet en particulier de modéliser l'importance objective d'un critère et son interaction avec les autres critères d'évaluation.

Plusieurs situations peuvent alors être distinguées :

- les objectifs déclinés en critères d'évaluation sont parfaitement non-interactifs ou indépendants, c'était l'hypothèse des paragraphes précédents ;
- les objectifs sont interactifs ; un type d'interaction concerne la notion de conflit. D'un autre point de vue, ces objectifs sont partiellement *complémentaires*, c'est-à-dire que leur satisfaction simultanée affecte la performance agrégée plus significativement que des satisfactions prises séparément (il y a synergie) ;

- un autre type d'interaction d'objectif est la redondance, c'est-à-dire des objectifs qui sont en quelque sorte interchangeables parce que les performances respectives auxquelles ils se rapportent évoluent dans le même sens.

Il est clair que si l'on est capable de construire une hiérarchie de critères où l'on a pris soin de n'utiliser que des objectifs indépendants et additifs, cette problématique n'a pas lieu d'être. Malheureusement, en pratique, l'indépendance suppose en particulier qu'il n'y ait pas de redondance entre les critères et la classique moyenne pondérée n'est pas capable de rendre compte de cet aspect. Plutôt que d'essayer de construire des structures hiérarchiques indépendantes et additives, on peut modéliser l'interaction d'objectifs comme le permet l'intégrale de Choquet qui, dans le cas 2-additif, prend la forme :

$$CI_g(u_1, u_2, \dots, u_p) = \sum_{\substack{I_{ij} > 0 \\ i > j}} \min(u_i, u_j) I_{ij} + \sum_{\substack{I_{ij} < 0 \\ i > j}} \max(u_i, u_j) |I_{ij}| + \sum_{i=1}^p u_i \left(v_i - \frac{1}{2} \sum_{j \neq i} |I_{ij}| \right) \quad \text{II-9}$$

avec $v_i - \frac{1}{2} \sum_{j \neq i} |I_{ij}| \geq 0$ et où :

- les v_i sont les indices de Shapley et représentent l'importance globale de chaque objectif relativement à tous les autres ($\sum_{i=1}^p v_i = 1$) ;
- les I_{ij} représentent les interactions entre les paires de critères (C_i, C_j) qui prennent leur valeur dans l'intervalle $[-1; 1]$; une valeur de 1 signifie qu'il y a un effet synergique positif entre les deux objectifs, une valeur de -1 qu'il y a une synergie négative, et une valeur nulle que les objectifs sont indépendants.

Ainsi, peut-on décomposer l'intégrale de Choquet 2-additive en trois parties, une conjonctive, une disjonctive et une additive :

- un I_{ij} positif implique que la satisfaction simultanée des critères C_i et C_j a un effet significatif sur la performance agrégée et qu'une satisfaction unilatérale n'a par contre aucun effet ;
- un I_{ij} négatif implique que la satisfaction de C_i ou C_j est suffisante pour avoir un effet significatif sur la performance agrégée ;
- un I_{ij} nul implique qu'il n'existe pas d'interaction entre les deux critères considérés ; ainsi les v_i jouent le rôle de poids des moyennes pondérées classiques.

Les propriétés mathématiques de l'intégrale de Choquet ont été largement traitées dans les travaux de Grabisch [Grabisch *et al.*, 1995; Grabisch *et al.*, 1998; Grabisch *et al.*, 1996, 2000]. Nous ne donnons que quelques exemples remarquables à titre illustratif des propriétés de cette intégrale floue.

L'intégrale de Choquet peut aussi s'écrire sous la forme suivante qui permet de la placer par rapport à la moyenne pondérée associée par les v_i :

$$CI_g(u_1, u_2, \dots, u_p) = \sum_{i=1}^p u_i v_i - \frac{1}{2} \sum_{i > j} |u_i - u_j| I_{ij} \quad \text{II-10}$$

3.2.3. La phase de sélection

Evaluation et classement

Le jugement de valeur ou le score porté par la CA_i concernant l'alternative k selon le point de vue du critère j se note x_j^{k, CA_i} . Le score $X_j^k \in [0, 1]$ obtenu par l'alternative k pour le critère j est le résultat de l'agrégation des scores portés par toutes les CA s se rapportant à la case $(j; k)$

de la grille d'évaluation. Les termes x_j^{k,CA_i} cartographiés dans la grille d'évaluation permettent l'évaluation de l'alternative k selon la dimension j de l'évaluation : une relation d'agrégation du type $X_j^k = g_j(x_j^{k,CA_i})$ est établie, dans laquelle g_j est l'un des opérateurs d'agrégation définis dans la section précédente.

Dans [Akharraz, 2004], g_j est la moyenne arithmétique des x_j^{k,CA_i} . A ce niveau d'agrégation, il nous semble néanmoins plus indiqué d'avoir recours soit à des majorités restreintes soit à des unanimités restreintes si l'on souhaite éviter les compensations entre x_j^{k,CA_i} . Ces familles d'opérateurs s'apparentent davantage au principe de vote collectif (*la quasi-totalité des connaissances x_j^{k,CA_i} permet de conclure que...*) que l'on cherche à modéliser qu'une moyenne arithmétique qui lisse l'ensemble des opinions et en masque complètement la diversité. Par ailleurs, si l'on considère que les x_j^{k,CA_i} sont aussi fonction du temps, la moyenne arithmétique aura l'effet d'un filtre moyenneur et déformera la dynamique d'évolution de X_j^k (la détection de tout revirement d'opinion sera retardée). Enfin, majorités et unanimités restreintes sont des opérateurs à sémantique explicite à ce niveau de l'agrégation (*la plupart des connaissances x_j^{k,CA_i} vont dans ce sens, plus de 75% des connaissances x_j^{k,CA_i} laissent à penser que...*), qui n'exigent pour autant aucun paramétrage : le choix des w_j dans ce type d'opérateurs ne nécessite aucune expertise du domaine (pas d'identification paramétrique).

Enfin, l'évaluation globale de la solution k , S^k , résulte de l'agrégation des X_j^k selon un opérateur d'agrégation h comme défini précédemment. Cependant à ce niveau d'agrégation, il devient raisonnable de penser que h puisse modéliser une stratégie, des priorités, un comportement décisionnel (conjonctif, disjonctif ou de compromis) et on pourra alors procéder à une identification comme nous l'avons évoqué précédemment. A l'issue de cette évaluation multicritère, les solutions sont donc classées selon leur score global $h(s^k) = h(X_1^k, \dots, X_p^k)$.

A ce stade du processus de décision, on dispose donc d'une solution préférée (ou élue) s^i : elle a été retenue suivant une stratégie h . Par définition, on a : $\forall j \neq i, h(s^i) \geq h(s^j)$.

Argumentation de la sélection

Nous avons vu que la légitimation des choix est une dimension essentielle de la décision collective. Intéressons nous donc maintenant à la formalisation du processus de justification de la solution retenue. Il s'agit dans cette étape d'automatiser l'extraction des CAs de la base de connaissance, autrement dit les éléments rhétoriques qui ont conduit à la décision.

Considérons la procédure d'argumentation relative suivante qui permet de justifier le choix de la solution s^i plutôt que la solution s^k . Elle repose sur la notion de contribution d'un critère à l'évaluation globale.

Soit $h(s^1) - h(s^k) = h(x_1^1, \dots, x_p^1) - h(x_1^k, \dots, x_p^k)$

Lorsque h est $n+1$ -différentiable et $\|h^{(n+1)}(x_1, \dots, x_p)\| \leq M, \forall \mathbf{x} = (x_1, \dots, x_p)$

Alors on a :

avec $\mathbf{h} = \mathbf{x}^1 - \mathbf{x}^k$,

$$\left\| h(\mathbf{x}^k + \mathbf{h}) - h(\mathbf{x}^k) - h'(\mathbf{x}_1^1, \dots, \mathbf{x}_p^1) \cdot \mathbf{h} - \dots - \frac{1}{n!} \cdot h^{(n)}(\mathbf{x}_1^1, \dots, \mathbf{x}_p^1) \cdot (\mathbf{h})^n \right\| \leq M \cdot \frac{\|\mathbf{h}\|^{n+1}}{(n+1)!} \quad \text{II-11}$$

Pour $n=1$, la décomposition en contributions marginales des critères au point s^k est donnée par :

$$h(s^i) - h(s^k) = h(\mathbf{x}^k + \mathbf{h}) - h(\mathbf{x}^k) \cong h'(x_1^k, \dots, x_p^k) \cdot \mathbf{h} = {}^t \nabla h(s^k) \cdot \mathbf{h} = \sum_{j=1}^p \frac{\partial h(\mathbf{x}^k)}{\partial x_j} \cdot \mathbf{h}_j \quad \text{II-12}$$

Les contributions marginales $(\frac{\partial h(\mathbf{x}^k)}{\partial x_j} \cdot \mathbf{h}_j)$ peuvent être réécrites de sorte à ce que

$$\frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(j)}} \cdot \mathbf{h}_{\sigma(j)} \geq \frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(j+1)}} \cdot \mathbf{h}_{\sigma(j+1)}.$$

On peut ensuite partitionner les contributions $\frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(j)}} \cdot \mathbf{h}_{\sigma(j)}$ en classes relatives aux ordres de

grandeur du ratio⁶ : $\frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(j)}} \cdot \mathbf{h}_{\sigma(j)} / \frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(1)}} \cdot \mathbf{h}_{\sigma(1)}$. Plus ce ratio est proche de 1, plus la

contribution selon le critère j est forte, plus le critère j représente une dimension essentielle dans la justification de notre choix. On définit ainsi un découpage symbolique sur l'échelle

continue $x / \frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(1)}} \cdot \mathbf{h}_{\sigma(1)}$ dans lequel chacun des termes $\frac{\partial h(\mathbf{x}^k)}{\partial x_{\sigma(j)}} \cdot \mathbf{h}_{\sigma(j)}$ peut être classé (Figure II.5).

On peut alors appliquer un raisonnement aux ordres de grandeur pour expliquer avec le niveau de détail choisi—en un mot, l'essentiel, le principal, le détail et l'anecdotique—en quoi la solution s^i est préférable à la solution s^k ; plus la contribution marginale est importante, plus le critère est une dimension discriminante de la sélection. Les symboles des relations aux ordres de grandeur $=, \cong, \approx, <$ et $<<$ signifient respectivement *égal à*, *voisin de*, *comparable à*, *petit devant* et *négligeable devant* et sont associés aux explications « en un mot », « à l'essentiel », « principalement », « dans le détail » et « de façon anecdotique »

(Figure II.5). Ainsi, à titre d'exemple, si l'on a $\frac{1}{1+e} \leq \frac{\partial h(s^k)}{\partial x_{\sigma(j)}} \cdot \mathbf{h}_{\sigma(j)} / \frac{\partial h(s^k)}{\partial x_{\sigma(1)}} \cdot \mathbf{h}_{\sigma(1)} < 1$, la

contribution du critère j au score global de s^k est voisine de 1 et on dira alors que le score obtenu par s^k selon le critère j est une raison essentielle de son score global.

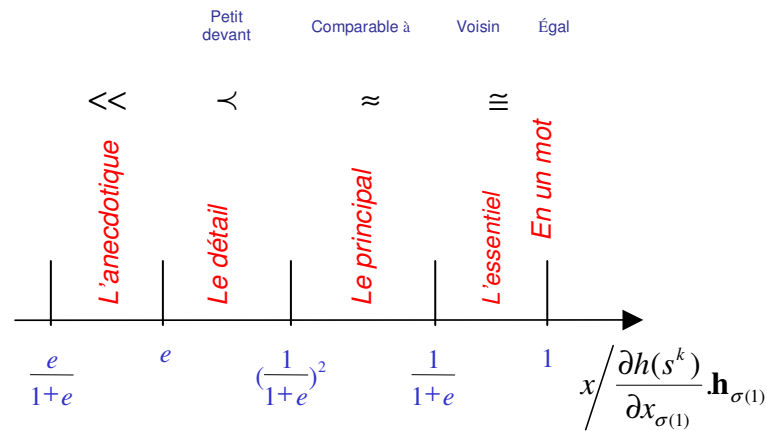


Figure II.5: Découpage en étiquettes symboliques des contributions partielles⁷

⁶ Une relation « $A \text{ r } B$ » est équivalente à « $(A/B) \text{ r } 1$ » et peut être modélisée comme un intervalle flou sur le rapport (A/B) en utilisant un unique paramètre e [Mavrovouniotis *et al.*, 1988][Dubois, 1989].

⁷ Le nombre de classes est un paramètre à fixer qui dépend de la connaissance que l'on souhaite exprimer (le détail et l'anecdotique pourraient par exemple être fondus en une seule catégorie linguistique). Le paramètre e est

On peut également souhaiter disposer d'une procédure de justification dans l'absolu, autrement dit faire valoir quelles sont les dimensions qui expliquent le bon score de s^k . Considérons alors l'expression analytique suivante :

$$h(s^k) - h(\alpha, \dots, \alpha) = \sum_{j=1}^n \frac{\partial h(\alpha, \dots, \alpha)}{\partial x_j} \cdot (x_j^k - \alpha)$$

II-13

Où $h(\alpha, \dots, \alpha) = \alpha$ peut être vu comme une valeur de référence ou un degré minimal de satisfaction attendu. Les contributions marginales associées à cette expression sont alors :

$(\frac{\partial h(\alpha, \dots, \alpha)}{\partial x_j} \cdot (x_j^k - \alpha))$. Elles peuvent être ordonnées de sorte à ce que :

$$\frac{\partial h(\alpha, \dots, \alpha)}{\partial x_{\sigma(j)}} \cdot (x_{\sigma(j)}^k - \alpha) \geq \frac{\partial h(\alpha, \dots, \alpha)}{\partial x_{\sigma(j+1)}} \cdot (x_{\sigma(j+1)}^k - \alpha)$$

puis on peut appliquer un raisonnement aux ordres de grandeur comme précédemment : plus la contribution est importante, plus le critère a été déterminant dans la bonne appréciation de s^k .

Une autre façon d'envisager le niveau de justification est de fixer le pourcentage $\beta\%$ d'explication du score de s^k et de chercher p_0 tel que :

$$h(s^k) - h(\alpha) \triangleq \sum_{j=1}^{p_0 \leq p} \frac{\partial h}{\partial x_j} \cdot (x_j^k - \alpha) = \beta\% \cdot \sum_{j=1}^p \frac{\partial h}{\partial x_j} \cdot (x_j^k - \alpha) \quad \text{II-14}$$

A ce stade de l'argumentation, il est donc possible de sélectionner automatiquement les critères les plus déterminants dans la décision ; autrement dit, les cases (j, k) de la grille d'évaluation pour lesquelles les scores partiels x_j^k associés jouent un rôle décisif dans les choix pris, peuvent être pointées *mécaniquement*. Il faut ensuite pour chacune de ces cases extraire les CAs les plus caractéristiques des arguments que l'on veut faire valoir. En effet, si dans les processus d'agrégation multicritère résumés précédemment, les CAs sont assimilées aux scores partiels qu'elles portent, il ne faut pas oublier qu'une CA est d'abord l'expression en langage naturel de l'interprétation finalisée par son rédacteur d'une information dans le cadre de l'objectif qu'il poursuit, c'est-à-dire un élément d'argumentation ou de rhétorique.

De façon analogue à la sélection des critères les plus pertinents, on peut mettre en place une procédure permettant d'extraire automatiquement les CAs les plus pertinentes pour soutenir la logique décisionnelle que l'on veut faire valoir, autrement dit de générer automatiquement la rhétorique nécessaire pour faire valoir la solution retenue.

Cette fonctionnalité d'aide à l'argumentation souligne l'idée qu'un système d'aide à la décision doit avant tout permettre d'éclairer le décideur au cours des phases du processus décisionnel, et rend ainsi notre système conforme à la définition que propose B.Roy : l'aide à la décision ne relève que de façon très partielle de la recherche de vérité ; les théories, les méthodologies, les concepts, les modèles, les techniques sur lesquels elle s'appuie doivent avoir une ambition différente : raisonner le changement que prépare un processus de décision de façon à accroître la cohérence des objectifs avec le système de valeurs du décideur [Roy, 1985].

Dans cette section 3.2, nous avons repris les principes de l'évaluation multicritère et de l'argumentation tels que Akharraz les a proposés dans sa thèse dans le cas particulier où

à déterminer, une valeur courante de e est 0.1 et correspond à l'idée commune qu'une grandeur devient négligeable devant une autre au-delà d'un rapport 1/10.

l'opérateur h était l'intégrale de Choquet (non différentiable en tout point) et $\forall j, g_j$ une moyenne arithmétique. Nous avons proposé des opérateurs à sémantique plus riche pour les g_j et une écriture nouvelle des expressions analytiques utiles à l'explication lorsque l'opérateur h est différentiable en introduisant les dérivées partielles de cet opérateur ce qui permet de généraliser le concept de contributions marginales d'un critère à tout opérateur différentiable. Nous allons procéder de même pour la phase de révision.

3.2.4. La phase de révision

Dans le schéma de Simon, toute décision raisonnée doit reposer sur une analyse de la situation qui permet d'expliciter et de caractériser les phases de perception, de représentation et de sélection du processus décisionnel, ces trois phases de la décision ne se présentant pas de façon linéaire, mais en boucles (loops). Pour rendre compte de cette boucle d'itération dans le processus de décision, nous allons introduire la notion de contrôle de la dynamique de décision.

A cet effet, nous introduisons la notion de risque décisionnel. Le risque décisionnel est fonction du temps : il est associé à la notion de fiabilité du classement à un instant donné. Pour définir le risque décisionnel, nous proposons de nous baser sur une notion de distance entre l'alternative préférée à un instant donné s^i et les alternatives challengers s^k . La notion de risque est donc ici assez éloignée de la définition probabiliste classique et relève plutôt de la sensibilité des préférences établies à toute perturbation externe. Le risque décisionnel associé à la stabilité du classement est :

$$r = 1 - \min_{\substack{k=1..n \\ i \neq k}} d(s^i, s^k) \quad \text{II-15}$$

Il reste à choisir la distance d . On pourrait prendre la distance qui semble la plus naturelle, c'est-à-dire l'écart entre les scores globaux obtenus par les solutions concurrentes :

$$r = 1 - \min_{\substack{k=1..n \\ i \neq k}} [h(s^i) - h(s^k)] \quad \text{II-16}$$

Néanmoins, dans l'espace d'évaluation multidimensionnel, le challenger qui aurait la meilleure évaluation après s^i n'est pas nécessairement le concurrent le plus *menaçant*. Prenons par exemple le cas où h est une somme pondérée et imaginons que l'on évalue trois élèves A, B et C selon trois matières, mathématiques, français et histoire ([Tableau II-1](#)). Leurs coefficients respectifs sont : 5, 3 et 2. Le tableau suivant récapitule les résultats obtenus par les élèves au premier trimestre :

	Elève A	Elève B	Elève C
Mathématiques	17	20	12
Français	17	20	20
Histoire	17	1	20
Moyenne	17	16.2	16

Tableau II-1 : Exemple de classement et notion d'effort à produire

On a donc $h(A) > h(B) > h(C)$ et avec la définition de d précédente, $r = 1 - 0.8 = 0.2$. Néanmoins, on peut remarquer que B doit gagner 4 points en *Histoire* s'il veut revenir à hauteur de A, toute chose étant égale par ailleurs, alors que C n'a finalement que 2 points à prendre en *Mathématiques*...

Nous choisirons donc une distance qui tienne compte de cette sensibilité à toute amélioration potentielle des solutions concurrentes. Nous parlerons d'« effort » à fournir par les solutions concurrentes pour remettre le classement en question : plus cet effort devra être conséquent, plus le classement sera stable et moins le risque décisionnel sera important.

Ainsi, plus la distance entre l'alternative s^i recommandée et ses principales rivales est faible, plus la sélection de s^i est risquée : il suffirait d'« une faible quantité d'informations nouvelles et pertinentes » pour que le classement s'en trouve modifié. Le classement peut être très sensible à toute nouvelle CA modifiant les scores partiels. On considère que l'on peut prendre une décision lorsque le risque décisionnel est en deçà d'un seuil fixé C_r . Ce seuil définit la notion d'acceptabilité de la décision, la situation est décidable lorsque : $|C_r - r(t)| \leq \varepsilon$ (Figure II.6). ($r(t)$ est le risque à la date t)

Pour que cette contrainte soit vérifiée, il est nécessaire de déterminer les points sur lesquels, des informations complémentaires sont nécessaires pour réduire de manière efficace le risque entre deux dates d'évaluation, autrement dit diminuer l'incertitude épistémique. Néanmoins, l'information a un coût dans tout processus décisionnel : l'information est la matière première du processus de décision, on doit donc en faire un usage optimal. L'algorithme suivant propose donc de déterminer dynamiquement l'information « juste » nécessaire pour rendre une situation décidable, c'est-à-dire pour que $|C_r - r(t)| \leq \varepsilon$ soit vérifiée avec un temps de réponse réduit autant que possible. Il s'agit donc bien d'astreindre la dynamique du processus décisionnel (dynamique forcée) en contrôlant l'apprentissage pour obtenir un temps de réponse inférieur à celui où aucun contrôle n'est effectué quant à la façon dont l'organisation construit le corpus de connaissances qui lui est utile pour arrêter sa décision (dynamique libre) (Figure II.6).

En pratique, il faut donc chercher pour chaque alternative s^k les dimensions de l'évaluation pour lesquelles il est le plus pertinent (au sens de la rentabilité économique si on en revient au coût informationnel) de vérifier qu'il n'existe pas d'informations non encore capitalisées par l'organisation qui augmenteraient significativement le score global de s^k .

Nous proposons d'énoncer l'algorithme général du risque décisionnel sous la forme du problème d'optimisation suivant :

Objectif :

$$d(s^i, s^k) = \min \|f_{x^k}(\delta^k)\|_{L_1} = \min \left(\sum_{j=1}^p f_{j_{x_j^k}}(\delta_j^k) \right) \quad \text{II-17}$$

où $\delta^k = [\delta_1^k, \dots, \delta_p^k]^T$ est le vecteur d'amélioration nécessaire à s^k pour revenir à hauteur de s^i et $f_{j_{x_j^k}}(\delta_j^k)$ est le coût informationnel pour une réévaluation δ_j^k de s^k selon la dimension j depuis x_j^k et $f_{x^k}(\delta^k) \triangleq [f_{1_{x_1^k}}(\delta_1^k) \dots f_{j_{x_j^k}}(\delta_j^k) \dots f_{p_{x_p^k}}(\delta_p^k)]^T$.

Contraintes :

$$\begin{aligned} h(x^k + \delta^k) &= h(x^i) \\ \forall j, \quad 0 \leq \delta_j^k &\leq 1 - x_j^k. \end{aligned} \quad \text{II-18}$$

Soit $\bar{\delta}^k = [\bar{\delta}_1^k, \dots, \bar{\delta}_p^k]^T$ la solution de ce problème.

Notons que si on suppose que les coûts informationnels sont linéaires et que h est linéaire (une somme pondérée par exemple), alors ce problème est un simplex. Dans le cas général, ce problème est un problème d'optimisation non linéaire. Akharraz s'est intéressé à la résolution du problème dans le cas particulier de l'intégrale de Choquet qui est une fonction linéaire par morceaux sur $[0, 1]^n$ [Akharraz, 2004]. Dans une première approximation, on peut s'abstraire du problème de l'identification des fonctions *coût informationnel* et considérer que $\forall i, f_i = Id$ (le coût de l'information est indépendant du score d'une part, de la dimension d'autre part)⁸. Le problème simplifié s'énonce alors :

⁸ Dans l'exemple des élèves, cela revient à dire que progresser de deux points en Mathématiques n'est pas plus difficile quand on a 4 de moyenne que lorsqu'on a 12 d'une part, que le travail ou l'effort fourni pour gagner k points est indépendant de la matière et donc que l'élève ne présente pas d'aptitude particulière d'une matière à l'autre.

$$\text{Objectif :} \quad \min \|\delta^k\| = \min \left(\sum_{j=1}^n \delta_j^k \right) \quad \text{II-19}$$

$$\text{Contraintes :} \quad \begin{aligned} & h(\mathbf{x}^k + \delta^k) = h(\mathbf{x}^i) \\ & \forall j, \quad 0 \leq \delta_j^k \leq 1 - x_j^k \end{aligned}$$

De cet algorithme, on déduit le risque décisionnel : $r = 1 - \min_{\substack{k=1..n \\ k \neq i}} \|\bar{\delta}^k\|_{L_1} / p$ que l'on compare au

seuil d'acceptabilité C (la division par p permet de normaliser la distance maximale à 1). Si le risque est inacceptable ($|C_r - r(t)| > \varepsilon$), il est nécessaire d'acquérir de nouvelles informations afin de pouvoir trancher entre les options. L'algorithme va mettre en lumière les dimensions de l'évaluation qui rendent le choix critique entre s^i et chaque option s^k . C'est-à-dire qu'il va permettre d'identifier les dimensions selon lesquelles, il sera le plus « rentable », le plus pertinent d'acquérir des compléments d'informations. Le principe calculatoire est donné ci-dessous.

L'amélioration minimale nécessaire à s^k pour atteindre le score global de s^i selon la dimension j est notée $\bar{\delta}_j^{(ik)}$. Le gain associé à cette amélioration partielle est donné par :

$$G_j^{k \rightarrow i} = \int_{x_j^k}^{x_j^k + \bar{\delta}_j^{(ik)}} \frac{\partial h}{\partial x_j} dx_j \quad \text{II-20}$$

Le gain global espéré de s^k pour rejoindre s^i est :

$$G^{k \rightarrow i} = h(s^i) - h(s^k) = \sum_{k=1}^p \int_{x_j^k}^{x_j^k + \bar{\delta}_j^{(ik)}} \frac{\partial h}{\partial x_j} dx_j \quad \text{II-21}$$

L'importance relative moyenne (IRM) du critère j le long de la trajectoire résultant de l'algorithme du risque décisionnel, $IRM_j^{k \rightarrow i}$, est définie par :

$$G_j^{k \rightarrow i} = IRM_j^{k \rightarrow i} \cdot \bar{\delta}_j^{(ik)} \Rightarrow IRM_j^{k \rightarrow i} = \frac{1}{\bar{\delta}_j^{(ik)}} \cdot \int_{x_j^k}^{x_j^k + \bar{\delta}_j^{(ik)}} \frac{\partial h}{\partial x_j} dx_j \quad \text{II-22}$$

Plus l' $IRM_j^{k \rightarrow i}$ est grande, plus le critère j est une dimension critique sur laquelle l'apport de nouvelles informations relatives à l'alternative s^k a de « chances » de modifier le classement antérieur des alternatives en concurrence. Par conséquent, les acteurs de la décision auront intérêt à rentrer en priorité de l'information relativement aux dimensions critiques identifiées par l'algorithme du risque. Autrement dit, l'algorithme « délivre » aux acteurs du processus décisionnel, un « signal de commande » qui leur permet de discerner les analyses à approfondir, les points sur lesquels une investigation plus poussée est sans intérêt, etc. C'est en ce sens que l'on parlera de pilotage par le risque décisionnel, la phase de révision permet de calculer les entrées informationnelles pertinentes qui permettront de jouer sur la dynamique du processus décisionnel. Via la phase d'apprentissage ou d'information, on induit une dynamique forcée du processus décisionnel.

Comme pour la phase de sélection, nous sommes partis du modèle d'Akharraz spécifiquement conçu pour l'intégrale de Choquet pour proposer un algorithme du risque décisionnel général pour tout opérateur différentiable avec les conventions dont nous avons convenu pour la phase de justification.

3.2.5. Synthèse

Dans les quatre paragraphes précédents, nous avons ainsi défini une boucle de contrôle dont les éléments caractéristiques sont les suivants (Figure II.6).

Le *processus à contrôler* est le processus d'évaluation assimilé à un système dynamique dont les entrées sont les CAs et la sortie le classement des alternatives possibles.

Dans cette interprétation, le risque décisionnel peut être vu comme la *variable régulée*. Compte tenu des remarques précédentes, remarquons que contrôler le risque décisionnel,

revient à contrôler l'entropie au sens large du terme de la base de connaissance utilisée pour la décision.

Le système d'information ou SGDC est l'actionneur de ce schéma de contrôle : il est l'organe d'action qui « ouvre ou ferme la vanne d'informations utiles à la décision ». Le calcul du risque et des critères sensibles du classement constitue le régulateur du système.

L'acceptabilité de la décision prend une interprétation bien particulière dans ce cadre : c'est la consigne fixée sur le risque décisionnel. Si l'on est en dessous de ce seuil, la situation est décidable, aucune étude supplémentaire n'est requise pour proposer un choix qui ne pourra être « stratégiquement » désavoué (arrêt ou ralentissement des études devenues inutiles compte tenu de l'état suffisant des connaissances). Par contre, si l'on ne parvient pas à ramener le risque décisionnel en dessous de la consigne fixée, indépendamment de nouveaux flux d'informations, la stratégie d'agrégation doit être révisée (les paramètres la déterminant peuvent être réajustés, les critères d'évaluation réexaminés) : aucun consensus ne pourra être trouvé avec cette stratégie (Figure II.6).

Toute information que l'on rentre dans le SGDC sans que la dimension à laquelle elle se réfère au calcul du risque décisionnel est une perturbation : ainsi, l'inertie psychologique ou la controverse sont typiquement deux comportements organisationnels identifiables à des perturbations de la dynamique d'évolution du processus de décision (ce sont des perturbations sur l'actionneur ; l'incertitude sur la stratégie appliquée pourrait constituer une perturbation en sortie) (Figure II.6).

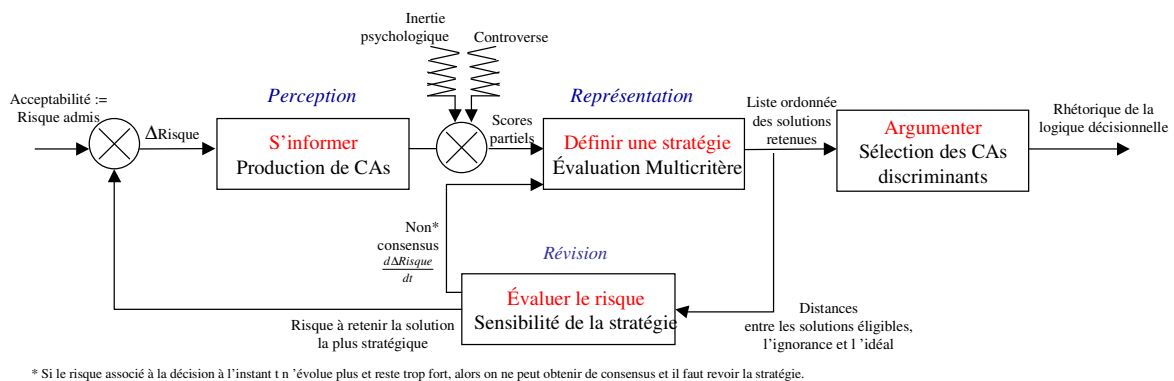


Figure II.6: Contrôle de l'information décisionnelle ou pilotage par le risque

Revenons au modèle S.T.I de H.A.Simon. Les causalités enchevêtrées mises en exergue dans le modèle de l'économiste trouvent ici une interprétation en termes de boucle de contrôle. La phase d'information ou de perception est à associer à l'actionneur de notre schéma, la phase de représentation (le processus d'évaluation multicritère ici) est le système à contrôler, la phase de sélection (argumentation de la logique de choix) est la matrice d'observation du système. En effet, la sortie du processus d'évaluation est le classement des alternatives sur la base des scores qui leur ont été attribués au cours de l'évaluation. Ce classement n'est qu'un artifice mathématique pour rendre compte de la combinaison des intensités de préférence des alternatives en fonction des critères à un instant donné. Les scores sont assimilables à des variables d'état dont le décideur n'a « que faire » pour justifier ses actions. Comme nous l'avons expliqué dans la section précédente, une dimension essentielle de notre système d'aide à la décision est sa capacité à justifier ses évaluations et par conséquent un classement résultant des CAs du SGDC. La fonctionnalité de légitimation joue le rôle d'une matrice d'observation : elle transforme le « vecteur d'état », les scores, le classement, en un « observable », la rhétorique de la logique décisionnelle appliquée pour obtenir ce classement. Cette rhétorique de logique décisionnelle, qui prendra la forme de rapports d'argumentation remis par le décideur à ses administrés, aux pouvoirs publics, etc., constitue l'observable du processus d'évaluation par les CAs.

4. Conclusion

Nous avons proposé dans ce chapitre une synthèse rapide du modèle S.T.I de Simon pour un processus décisionnel collectif. Akharraz a donné une interprétation cybernétique de ce modèle dans le cas particulier où la stratégie d'agrégation est modélisée par une intégrale de Choquet. [Akharraz, 2004; Montmain *et al.*, 2006] est consacrée à la mise en œuvre de l'intégrale de Choquet dans les calculs liés à l'explication et au risque décisionnel. À l'inverse, notre démarche n'était pas de choisir un opérateur d'agrégation spécifique et d'en utiliser toutes les propriétés pour optimiser les calculs. Nous nous sommes fixés comme objectif de donner un modèle formel général que l'on peut instancier par tout opérateur d'agrégation du moment qu'il est différentiable. Nous sommes donc restés au niveau formel et n'avons pas proposé de méthodes de résolution particulière en ce qui concerne l'algorithme du risque décisionnel car les propriétés spécifiques des opérateurs font qu'il est difficile de prôner une méthode de résolution générale. On peut juste dire de manière très abstraite que ce problème d'optimisation non linéaire contraint peut être résolu par un algorithme de programmation quadratique séquentielle (sequential quadratic programming (SQP)). Un certain nombre de packages comme NPSOL, NLPQL, OPSYC, OPTIMA, MATLAB, et SQP, permettent de résoudre ce problème. Dans sa version la plus basique, l'algorithme de SQP remplace la fonction *objectif* par une approximation quadratique et les contraintes par des approximations linéaires. Mais encore une fois, notre objectif était un modèle formel, pas une technique de résolution dédiée à un opérateur ou une famille d'opérateurs précis.

Dans la boucle de régulation que nous proposons, nous avons montré comment le traitement de l'information pouvait être automatisé pour justifier une logique décisionnelle, pour estimer la sensibilité ou la fiabilité d'un classement, pour déterminer quels apports informationnels nouveaux seraient les plus pertinents pour conclure au plus vite à une situation décidable. Le « signal de commande » que fournit l'algorithme du risque décisionnel consiste à indiquer aux acteurs de la décision, les dimensions selon lesquelles l'évaluation n'est pas suffisamment fiable, c'est-à-dire les dimensions j^* où il n'existe pas d'éléments décisionnels (CAs) suffisamment significatifs et pertinents pour justifier que la solution s^j est préférable à s^k . Si l'on revient sur la structure de la base des CAs sous la forme d'une grille d'évaluation, cela signifie que l'algorithme identifie les cases (j^*, k) à renseigner prioritairement, sur lesquelles s'informer plus précisément. Autrement dit, cette boucle permet d'identifier les informations dont l'acquisition serait la plus à même de réduire les incertitudes épistémiques dans un processus d'évaluation itératif.

Dans [Akharraz, 2004], le « signal de commande » (les couples (j^*, k)) est en fait une simple recommandation aux acteurs de la décision, ce sont ces acteurs—des êtres humains—qui sont supposés se procurer cette information. Le rôle de « l'actionneur » est joué par un être humain auquel on donne simplement des conseils quant aux informations qu'il a intérêt à rentrer prioritairement dans ses bases.

Les chapitres techniques à suivre de ce mémoire proposent une solution pour mécaniser le traitement de l'information au niveau de l'actionneur. Le principe est de s'appuyer sur des techniques de classification automatique pour aller chercher dans des bases externes, sur le web, etc., les informations requises par le régulateur (relatives aux couples (j^*, k)). On va donc s'attacher par la suite à mécaniser l'actionneur de notre boucle de régulation pour tendre vers un système d'apprentissage et de décision sans intervention humaine. L'actionneur intelligent devra donc être capable de distinguer dans une base une information relative à un couple (j^*, k) , d'estimer s'il s'agit d'une information favorable ou non à la solution k au regard du critère j^* et enfin, de lui attribuer un jugement de valeur quantitatif, un score, afin qu'elle puisse automatiquement être utilisée dans l'évaluation au pas de temps suivant.

Il s'agit là d'un schéma très théorique de la décision avec une automatisation cognitive poussée aux extrêmes. Nous verrons qu'outre les problèmes déontologiques qu'une telle approche peut soulever, il subsiste des limites formelles et techniques à notre approche qui préservent l'avenir de nos décideurs...

Néanmoins, nous avons vu dans ce chapitre que l'interprétation cybernétique du modèle de Simon permettait de retrouver et d'instrumenter les grandes fonctionnalités du SIADG du chapitre précédent (suivi dans le temps, évaluation, compréhension, justification, conseil d'action, etc.). La seule différence réside dans le fait que l'être humain est exclu de la boucle de notre modèle formel alors que dans l'approche fiabiliste que nous avons présentée, on mettait en avant l'idée de support, d'assistance pour pallier la faiblesse ou l'erreur humaine.

Chapitre III. Gestion des connaissances et décision

Gestion des connaissances et décision

1. Introduction	72
2. La course à l'information à l'ère du WWW	72
2.1. La croissance exponentielle de l'information	72
2.2. Un besoin avéré ou suscité ?	73
2.3. La surcharge cognitive	74
2.4. Indexer pour mieux gérer	74
3. La notion de connaissance.....	75
3.1. Nature de la connaissance	75
3.2. Une typologie des connaissances pour une organisation	77
4. La gestion des connaissances	78
5. Le processus d'indexation : de l'information à la connaissance	79
5.1. Indexation manuelle	80
5.2. Indexation automatique	81
5.3. Comparaison des deux approches	81
6. Connaissance actionnable	81
7. Instanciation de ces notions à notre problématique	82

1. Introduction

Les deux chapitres précédents ont mis en avant le rôle de l'information dans la décision collective. Le modèle que nous avons proposé au chapitre 2 est dans la lignée du modèle Information Processing System (I.P.S) de H.A. Simon, qui, le premier, dès 1947, a souligné l'importance de la phase d'information dans une décision d'organisation [Newell *et al.*, 1972; Simon, 1983]. Les fonctionnalités d'évaluation, de comparaison, de justification des choix que nous avons spécifiées puis formalisées reposent sur le concept élémentaire et fondamental de connaissance actionnable (CA), savoir élémentaire utile à l'action. La CA est un jugement, une appréciation qui correspond à l'interprétation par un acteur du processus décisionnel, de la pertinence et de l'utilité d'une information sous la perspective d'analyse d'une dimension de la décision. C'est de l'acquisition, de la disponibilité des CAs que dépend donc la pertinence de notre SIADG. Aussi, nous a-t-il semblé légitime de revenir dans ce chapitre 3 sur la notion de connaissance d'une part, sur la gestion de la connaissance d'autre part.

Ainsi, ce chapitre présente quelques éléments bibliographiques relatifs à la notion d'apprentissage pour la décision qui nous permettent de resituer nos travaux dans la problématique de la gestion des connaissances. Le chapitre est composé de trois sections. La première s'attache à cerner le concept complexe de connaissance. Dans un second temps, nous abordons la problématique de la gestion des connaissances. Enfin, nous recadrons ces éléments dans notre projet : apprendre, générer et capitaliser des CAs pour décider.

Ce chapitre n'a pas la prétention de faire une synthèse de tout ce qui a pu être écrit et formalisé autour du concept de connaissance tant dans les Sciences Humaines et Sociales qu'en Intelligence Artificielle (pour ne citer qu'elles), il a seulement pour but de légitimer notre interprétation de la notion de connaissance et plus spécifiquement de connaissance actionnable dans le contexte de la décision en organisation.

2. La course à l'information à l'ère du WWW

2.1. La croissance exponentielle de l'information

De nos jours l'explosion des technologies de l'information a fait de l'Internet un outil, incontournable tant au niveau personnel que professionnel occupant une place prépondérante dans notre société et notre vie. Le WWW nous offre un monde de l'information prodigieux sans limite, ce qui pourrait apparaître comme quelque chose de véritablement bénéfique, mais d'un autre côté, c'est l'ouverture à un monde inexplorable et incontrôlable (en terme de quantité et de fiabilité des informations délivrées), où l'ensemble des choix possibles est quasi infini...

Le problème pour le décideur en quête de l'« information pertinente » n'est plus l'accès à cette information qui tend à être uniformément distribuée, mais de savoir la trouver dans d'imposants flux informationnels extrêmement « bruités ». Le journal « Le Monde », du 14 mars 2001, annonçait qu'une étude de l'Université de Berkeley avait établi que l'homme créerait plus d'information dans les deux à trois ans à venir qu'au cours des 40000 dernières années... Il n'est plus possible, pour l'individu et même à l'échelle d'une entreprise, d'une organisation, d'analyser dans sa globalité de telles quantités d'informations pour y trouver une réponse à un problème précis. Utiliser les nouvelles technologies pour acquérir la bonne information au bon moment, signifie qu'il faut savoir traiter des flux d'informations de façon « intelligente ». L'accès à la pertinence requiert un filtrage intelligent.

Cette croissance exponentielle de l'information est encouragée par le comportement boulimique des cyberconsommateurs qui accumulent les données, les documents, ainsi que par le besoin des entreprises d'avoir un accès rapide, voire quasi immédiat dans certains secteurs comme l'économie, à cette information. L'obtention rapide de données fiables et qualifiées devient un enjeu majeur, parfois critique, d'indépendance et de compétitivité.

Dans certains domaines, se pose un problème crucial d'utilisation et de stockage intelligent des données. Que ce soit sur Internet (sur des banques de données internationales) ou en interne dans les entreprises, la simple recherche d'informations prend désormais une part très importante du temps de travail.

Nous sommes donc face à une problématique nouvelle sur laquelle repose de plus en plus les différents secteurs d'activités de l'organisation. La constitution d'une mémoire d'entreprise, la gestion de sa connaissance interne et aussi externe est stratégique. Il s'agit non seulement de mémoriser correctement l'information (mise en forme, unification) mais surtout de la retrouver facilement. L'information gérée par l'entreprise se transforme donc en connaissances primordiales pour la survie et l'innovation de l'entreprise.

2.2. Un besoin avéré ou suscité ?

On peut s'interroger sur la relation de causalité qui lie l'information et le problème à résoudre...

Prenons tout d'abord un exemple de la vie quotidienne, chacun d'entre nous semble de plus en plus dépendant de l'information : un cyberconsommateur souhaitant acheter n'importe quel bien ou service sur le net ne saurait s'affranchir de consulter l'ensemble des offres commerciales via un comparateur de prix (ex : Kelkoo), de s'enquérir des biens ou produits réputés équivalents (service proposé sur la quasi-totalité des e-retailers), de s'assurer de la fiabilité du site sur lequel il envisage de faire son achat via un e-recommander (ex : Ciao.com)... En retour, les « marques » mettent à la disposition du consommateur de plus en plus d'informations relatives à leurs produits ou services sur la toile. Sur cet exemple banal, on voit à quel point la société de l'information a emprise sur l'individu.

Si la quantité d'informations disponible est phénoménale, elle n'est pas un gage de qualité et de fiabilité. Quelle crédibilité accorder à l'évaluation d'un produit lorsqu'elle est proposée par le site marchand lui-même (ex : Fnac, Amazon) ? Le cyberconsommateur se trouve vite démuni dans cet océan d'informations. Sa connaissance souvent très superficielle du site, du produit ou du service recherché, les informations souvent subjectives, parfois mêmes contradictoires qu'il consulte ne lui permettent pas d'évaluer et de comparer objectivement, rationnellement et exhaustivement une pléthore de sites tous susceptibles de satisfaire à son besoin [McNee *et al.*, 2003; Terveen *et al.*, 2001]. Ainsi confronté à cet inépuisable espace des possibles, les clients sans a priori, ont tendance à se tourner naturellement vers les opinions et les expériences d'autres cyberconsommateurs [Montmain *et al.*, 2005]. Ce comportement est à l'origine du concept de e-recommandation : les sites de e-recommandation sont dédiés à soutenir, gérer et automatiser les partages d'opinions et de recommandations de communautés de cyber-consommateurs (ex : Ciao.com, Leguide.com). On en vient à consulter de l'information sur l'information, à recourir aux statistiques pour estimer la fiabilité de ce qu'on peut lire [Denguir-Rekik *et al.*, 2006a] !...

Dans le monde de l'entreprise, les applications intranet se multiplient et deviennent le véritable poumon de la vie professionnelle. Chaque employé a accès aux informations de l'entreprise que ce soit des informations générales concernant le fonctionnement administratif mais également les données techniques nécessaires au travail quotidien. Les ERP sont souvent

de plus en plus accessibles via intranet et permettent de programmer et réguler l'activité quotidienne, en proposant par exemple la liste des tâches à effectuer dans la journée.

2.3. La surcharge cognitive

Aujourd'hui chaque utilisateur peut donc accéder à pratiquement toute l'information qu'il souhaite. Cependant l'abondance d'information ne résout pas les problèmes. Elle produit un phénomène bien connu qui est la surcharge cognitive. [Rouet *et al.*, 1995] donne la définition suivante : « la surcharge cognitive désigne un excès de traitement à réaliser ». [Clavien *et al.*, 2003] donne une définition légèrement différente : « Le flux continu d'information est ininterrompu, et trop rapide pour laisser le temps à l'utilisateur de construire la connaissance. On parle alors de surcharge cognitive ».

Pour faire référence à notre chapitre 1, la dite surcharge cognitive était déjà diagnostiquée dans les salles de conduite d'installations fortement informatisées par les gens de la fiabilité humaine il y a de cela vingt cinq ans ! Le débat portait sur l'intégration d'indicateurs toujours plus sophistiqués et plus nombreux dans les postes de conduite informatisés, l'incapacité de l'opérateur à traiter une avalanche d'alarmes dans le temps et l'espace. L'informatisation des contrôles et de certaines actions peut être poussée jusqu'à l'automatisation cognitive, c'est-à-dire le remplacement de l'opérateur pour certaines fonctions de décision [Després, 1991]. Si on reconnaît aujourd'hui de plus en plus clairement que, contrairement à ce qui était prévu, l'automatisation (au sens de la commande) a pu parfois rendre les tâches des opérateurs plus difficiles, la même chose peut être imaginée pour l'automatisation cognitive [Bainbridge, 1991].

Dans le cadre de notre étude où il s'agit d'alimenter le SIADG avec les informations pertinentes et utiles pour produire une décision, la surcharge cognitive provient de ce que l'utilisateur ne contrôle ni la quantité, ni la complétude, ni la qualité, ni la fiabilité des flux informationnels qu'il n'a de toutes façons pas la capacité (temps, compétences, etc.) d'analyser avec les précautions nécessaires que cela exigerait. Les questions qui se posent alors sont les suivantes : quelle est l'information utile ? Comment la repérer ? Comment la (re)trouver ?

Chaque utilisateur tente par divers moyens de se constituer une liste de références (que ce soit des sites, des revues, ou autres) lui permettant d'aller chercher la bonne information au bon moment. Mais cette attitude reflète davantage un besoin qu'elle n'apporte une solution efficace à terme.

2.4. Indexer pour mieux gérer

Il devient donc nécessaire de répertorier l'information contenue dans de grandes masses de documents. La problématique est alors de savoir comment annoter, c'est-à-dire indexer cette information pour la retrouver facilement.

Cette indexation doit-elle dépendre de l'usage que l'on fera de l'information ? Quelles sont les qualités d'une indexation ? Comment mettre en regard diverses techniques et les comparer suivant des critères objectifs : temps d'indexation, expression de la requête, pertinence du résultat obtenu, etc. ? L'indexation est-elle une activité automatisable ? Existe-t-il des méthodes hybrides qui permettent un gain significatif par rapport aux techniques existantes ?

Avant de tenter de répondre plus précisément à ces questions dans le chapitre suivant, il faut replacer l'indexation dans la problématique plus générale de la gestion des connaissances.

3. La notion de connaissance

L'analyse de la littérature existante dans ce domaine ne permet pas de définir une notion unique de « connaissance ». Cependant on peut déjà noter deux angles d'analyse de la notion de connaissance. D'abord, il est intéressant de définir la connaissance par opposition à deux autres notions qui du point de vue du langage commun sont à peu près employées indifféremment, la donnée et l'information. Par ailleurs, la nature et les modes de transmission de la connaissance ouvrent une autre perspective d'analyse qui donne lieu à une typologie des savoirs dans l'organisation.

3.1. Nature de la connaissance

Trois concepts principaux sont classiquement distingués : donnée, information et connaissance. Cette distinction nous est utile pour montrer le degré de progression dans les niveaux conceptuels associés à chaque notion. Cependant tous les auteurs ne sont pas d'accord sur cette distinction comme l'évoque [Vacher, 2000].

3.1.1. La donnée

Pour Ermine et Tsuchiya, les données sont des faits de base, qui apparaissent au cours de la réalisation d'une tâche [Ermine J.L. *et al.*, 96; Tsuchiya, 1995]. Elles peuvent être transcrites sous forme de chiffres, de mots ou de symboles, de figures... Pour [Brooking, 1998], les données sont des faits, des images, des nombres présentés sans aucun contexte. La notion de donnée est donc perçue comme la couche de plus bas niveau dans la hiérarchie conceptuelle du savoir. C'est sur elle que s'élaborent les notions d'information et de connaissance, elles sont la matière première de l'apprentissage d'un savoir. Dans le concept de donnée, il n'y a aucune notion d'interprétation ni de contexte, elle a une valeur quantitative ou qualitative, se présente sous la forme d'un nombre, d'un extrait électronique de texte ou d'une image, etc.

3.1.2. L'information

[Ermine J.L. *et al.*, 96] s'étonne du débat que soulève la définition du concept d'information, puisqu'il existe une définition précise, mathématique de celui-ci. Elle a été donnée par Shannon en 1949, et constitue ce qu'on appelle la théorie de l'information. Les premiers chercheurs qui ont essayé de définir mathématiquement la notion d'information contenue dans un message ont dressé une liste des propriétés que devrait vérifier la quantité d'information à partir d'une définition plausible. Ils ont ainsi procédé sans s'en rendre compte comme les thermodynamiciens du siècle précédent. C'est ainsi qu'« au sens strict de la théorie de l'information, l'information est une quantité, mesurée à l'aide d'une formule qui est sensiblement la même, (mais avec un signe inversé) que celle utilisée par le physicien Ludwig Boltzmann à la fin du XIXe siècle pour mesurer l'entropie des gaz » (Philippe Breton). La théorie de Shannon fournit donc l'aspect structurel de l'information. L'aspect fonctionnel de l'information concerne le traitement de l'information, parfois assimilée à la notion d'informatique.

Le dictionnaire Larousse donne cette définition : « Élément de connaissance susceptible d'être codé ou représenté à l'aide de conventions pour être conservé, traité ou communiqué ». Pour [Baizet, 2004] citant [Ermine J.L. *et al.*, 96; Tsuchiya, 1995], les informations sont pour les uns des données triées, sélectionnées et organisées par un individu dans un but précis, pour les autres des données auxquelles sont associées des significations par la description de méthodes et procédures d'utilisation. Ces informations sont à la base d'une possible communication et sont donc formulées, explicitées. Pour [Tsuchiya, 1993], quand on donne un sens à une donnée à travers un cadre interprétatif, elle devient information. Dans [Brooking, 1998] les

informations sont des données organisées présentées en contexte comme par exemple : les statistiques concernant les maisons d'habitation en Australie, etc.

De cette rapide synthèse, nous retiendrons les éléments suivants. La notion d'information présente donc un degré conceptuel plus élevé que la notion de donnée dans la valeur et la signification qu'elle occupe dans l'application où elle est utilisée. La notion d'information est toujours associée à la possibilité de traitement informatique c'est-à-dire son stockage sous forme exploitable pour les applications. Pour ces raisons, une information doit donc être associée à une représentation formelle permettant sa traduction informatique et conceptuelle.

3.1.3. Définition de la notion de connaissance

Dans [Ermine J.L. *et al.*, 96; Tsuchiya, 1995], les connaissances sont définies soit comme des informations affinées, synthétisées, systématisées, soit comme des informations associées à un contexte d'utilisation. Pour [Brooking, 1998], la connaissance est l'information en contexte, associée à une compréhension de son mode d'utilisation, comme par exemple : la connaissance au sujet du drainage de l'eau dans une rue, déduite de l'observation d'un schéma et la compréhension des influences de l'emplacement des maisons d'habitation sur ces drainages. Ce qu'il résume par l'équation symbolique : connaissance = information en contexte + compréhension. Pour [Bonjour *et al.*, 2002], la connaissance est une information affinée, synthétisée, se rapportant à un contexte spécifié et assurant une finalité bien précise.

[Tsuchiya, 1993] formule sa conception de la notion de connaissance d'une autre manière : « L'information ne devient connaissance que lorsqu'elle est comprise par le schéma d'interprétation du receveur qui lui donne un sens (sense-read). Toute information inconsistante avec ce schéma d'interprétation n'est pas perçue dans la plupart des cas. Ainsi la « commensurabilité » des schémas d'interprétation des membres de l'organisation est indispensable pour que les connaissances individuelles soient partagées. »

Pour [Penalva *et al.*, 2002] La connaissance, à l'inverse de l'information, repose sur un engagement, des systèmes de valeurs et de croyances, sur l'intention. La connaissance est bâtie à partir de l'information pour faire quelque chose, pour agir...

On peut considérer cette citation extraite de [Tsuchiya, 1993] comme une synthèse équitable des précédents points de vue : « Although terms 'datum', 'information', and 'knowledge' are often used interchangeably, there exists a clear distinction among them. When datum is sense-given through interpretative framework, it becomes information, and when information is sense-read through interpretative framework, it becomes knowledge. »

[Mayère, 1997] affirme encore que concernant la distinction tentée entre information et connaissance, il est très délicat en pratique de tenter une distinction claire entre ces termes car ils dépendent du point de vue de l'observateur. On remarquera simplement que l'utilisation du mot connaissance renvoie plus explicitement à une dynamique de création et d'échanges, dans lesquels les individus ont une position privilégiée.

Sous l'éclairage de ces définitions, la connaissance correspond au niveau conceptuel le plus élevé de l'apprentissage d'un savoir. La donnée est un élément brut qui n'est pas exploitable en tant que tel. Cette donnée devient information dans le contexte d'un cadre d'interprétation et de fait son traitement informatique est possible. La connaissance est un degré supplémentaire dans l'échelle conceptuelle et nécessite un contexte et surtout cette notion de « commensurabilité », c'est-à-dire cette exigence de compréhension mutuelle des acteurs dans l'échange de ces connaissances.

Une autre perspective d'analyse de la connaissance correspond à l'idée de classer les connaissances de par leur nature et leur apprentissage. C'est ce que nous allons résumer dans le paragraphe suivant.

3.2. Une typologie des connaissances pour une organisation

Dans *The Knowledge Creating Company*, Nonaka et Takeuchi analysent la performance des entreprises japonaises en matière d'innovation par leur capacité à combiner deux types de connaissances [Nonaka et al., 1996] :

- 1/ *les connaissances tacites* qui regroupent deux formes de savoir. Ces connaissances sont souvent essentielles pour l'entreprise, mais elles sont malheureusement peu ou pas toujours explicitables :
 - *La connaissance pratique* (ou savoir-faire) qui correspond aux routines : elle s'acquiert principalement par expérimentation et imitation ;
 - Les compétences comprennent également la *connaissance du contexte*, plus souvent appelée culture d'entreprise. Il s'agit de toutes les connaissances partagées sans forcément être explicitées, comme par exemple reconnaître un chercheur à sa tenue décontractée et un commercial à sa chemise bleue. Cela peut être de rester tard au bureau dans telle entreprise ou d'arriver tôt dans telle autre : rien n'est dit mais faire différemment serait jugé de façon très négative par les collègues.
- 2/ *Les connaissances explicites* qui sont facilement formalisables et se transmettent avec peu de perte d'intégrité. Ce sont des informations mises en forme pour une situation donnée, c'est-à-dire interprétées et mises en contexte (en ce sens elles deviennent connaissances). Elles peuvent être inscrites dans des supports : documents (narratifs, techniques ou descriptifs), plans, livres, dossiers informatiques, CD-ROM, etc. Ces supports mobilisent généralement des langages codifiés : normes, procédures, bons de commandes, factures, rapports, comptes rendus, publications, brevets, etc. Les valeurs chiffrées et les schémas ont une place importante. Ces documents sont de plus en plus électroniques avec l'avantage d'en faciliter la transmission.

Les auteurs présentent quatre processus dont il faut organiser la combinaison pour favoriser la création dynamique de connaissances (et donc l'innovation), c'est-à-dire le passage d'une connaissance à l'autre :

Venant de	Vers	Connaissance tacite	Connaissance explicite
Connaissance tacite		Socialisation	Extériorisation
Connaissance explicite		Internalisation	combinaison

Figure III.1 : Les processus de transformation des connaissances de l'entreprise

- 1/ La *socialisation* représente le transfert de connaissances tacites. Elle a lieu au cours d'échanges d'expériences, de rencontres où des histoires sont racontées (sur l'entreprise, sur la vie de chacun, etc.), de toutes sortes de contacts que l'organisation peut favoriser,
- 2/ L'*extériorisation* permet à des connaissances tacites de devenir explicites lorsque cela est possible. On mobilise le plus souvent les métaphores et les analogies,
- 3/ L'*internalisation* est le processus d'appropriation de connaissances explicites par les individus. Par exemple, on applique une norme sans y faire référence comme si elle était devenue évidente, on sait faire quelque chose qui auparavant nécessitait d'être décrit, etc.

4/ La *combinaison* est l'échange de connaissances explicites pour en créer de nouvelles : elles peuvent être reformulées pour être mobilisées dans un nouveau contexte, elles peuvent être mises sur des supports différents, elles se renvoient les unes aux autres, etc.

[Vinck, 1997] insiste sur l'inévitable coexistence des connaissances explicites et tacites (ou implicites) dans l'entreprise, du point de vue de l'individu. Il illustre ces deux concepts par la métaphore de l'iceberg de la connaissance, qui fait apparaître une partie immergée (tacite) et émergée (explicite) de la connaissance.

Ces deux notions sont essentielles pour comprendre les modes de fonctionnement des outils de gestion des connaissances. Un bon outil de gestion des connaissances sera celui qui permettra de stocker intelligemment les connaissances explicites et permettra aux utilisateurs de formaliser de façon plus aisée les connaissances tacites.

4. La gestion des connaissances

Nous allons donner un aperçu très limité des points de vue que l'on peut trouver dans la littérature sur la problématique de la gestion des connaissances. Nous examinons quelques approches représentatives dans le but d'en extraire les concepts importants et utiles pour définir et situer nos travaux par rapport aux approches existantes.

La notion de gestion des connaissances est déclinée sous différentes expressions qui comportent des ressemblances mais ne sont pas identiques :

- Gestion des connaissances ;
- Capitalisation des connaissances ;
- Knowledge management.

[Ballay, 1997] présente la gestion des connaissances comme un moyen permettant autant que possible, de valoriser les capacités et l'expérience de chacun à la place qui lui convient le mieux, de faire circuler l'information utile et d'aider à trouver au bon moment celle dont on a réellement besoin dans l'action.

De son côté, Prax, qui préfère utiliser l'expression "Knowledge Management", propose une définition en trois niveaux [Prax, 2000] :

- Le Knowledge Management est une approche qui tente de manager des items aussi divers que pensées, idées, intuitions, pratiques, expériences, émis par des gens dans l'exercice de leur profession ;
- Le Knowledge Management est un processus de création, d'enrichissement, de capitalisation et de diffusion des savoirs qui implique tous les acteurs et l'organisation, en tant que consommateurs et producteurs ;
- Le Knowledge Management suppose que la connaissance soit capturée là où elle est créée, partagée par les hommes et finalement appliquée à un processus de l'entreprise.

Une autre approche de la gestion des connaissances est celle développée par [Grundstein, 2000]. Celle-ci consiste à associer la fonction de management à la capitalisation des connaissances : « il faut insister sur le fait que la capitalisation des connaissances est une problématique permanente, omniprésente dans les activités de chacun, qui devrait de plus en plus imprégner la fonction de management ». Il précise encore : « capitaliser les connaissances, c'est considérer certaines connaissances utilisées et produites par l'entreprise

comme un ensemble de richesses et en tirer des intérêts contribuant à augmenter la valeur de ce capital » [Grundstein, 1995].

Nonaka et Takeuchi [Nonaka *et al.*, 1996] ont donné les fondements de la gestion des connaissances : « By organizational knowledge creation we mean the capability of a company as a whole to create new knowledge, disseminate it throughout the organization, and embody it in products, services and systems. ». Dans le contexte de concurrence mondiale actuel, la connaissance doit être un support de l'action, mais, plus exactement et pragmatiquement, un vecteur de créativité et de compétitivité. La gestion des connaissances selon Nonaka a donc pour but de s'appuyer sur un vécu pour stimuler l'apparition de nouvelles connaissances.

Pour [Ermine, 1996; Penalva *et al.*, 2002] La gestion des connaissances a été définie comme la mise en place d'un système de gestion de flux cognitifs qui permet à tous les composants de l'organisation à la fois d'utiliser et d'enrichir le patrimoine de connaissances de cette dernière. La gestion des connaissances se propose ainsi de repérer, formaliser, partager et valoriser les connaissances de l'organisation et en particulier celles qui revêtent un caractère stratégique et décisionnel.

Parmi toutes les démarches envisageables en gestion des connaissances, on constate deux tendances [Penalva *et al.*, 2002] :

- La *capitalisation des savoirs et du savoir-faire*, dont l'objectif principal est de consigner les connaissances stratégiques, et qui porte donc l'effort sur la sélection et la structuration des connaissances ;
- Le *partage dynamique des connaissances* qui ne préjuge pas de leur utilisation future, et n'élimine pas des connaissances dont l'intérêt pourrait être révélé plus tard.

C'est l'Intelligence Artificielle qui a développé la première voie avec une ingénierie des connaissances fondée sur la constitution de structures de données exploitables et « validables » dans les domaines techniques où une expertise stable peut être dégagée. La mise en forme de cette expertise est conditionnée par la finalité : l'implémentation en machine à traiter l'information. L'utilisation massive des réseaux d'information et de communication a relancé récemment l'intérêt de constituer des corpus d'éléments de connaissances dynamiques partageables entre acteurs humains. Les systèmes de gestion dynamique des connaissances (SGDC) explorent en priorité cette seconde voie. Un corpus de connaissances y est vu comme recouvrant un domaine de connaissances qui est aussi un domaine d'action et de décision.

Constituer et gérer via l'outil informatique une mémoire commune de l'organisation pour assister celle-ci dans ses prises de décisions, ses actions semble donc être un objectif communément admis. Le consensus entre les approches est plus difficile à établir quant à la constitution de cette mémoire. Il s'agit non seulement de mémoriser correctement l'information (mise en forme, unification) mais surtout de la retrouver facilement. Afin de mettre en œuvre cette étape il faut avoir recours à des méthodes d'indexation.

5. Le processus d'indexation : de l'information à la connaissance

Pour gérer une mémoire commune comme le propose la gestion des connaissances, intéressons nous d'abord aux moyens de constituer la base de connaissance. Plusieurs approches sont possibles. Nous devons pour cela répertorier l'information contenue dans une grande masse de documents : c'est le processus d'indexation, qui permet d'associer à chaque document textuel un certains nombres de repères qui permettent ensuite à des outils informatiques de retrouver l'information ainsi indexée. Nous n'aborderons l'indexation dans ce

chapitre que relativement à l'étape qu'elle constitue en gestion des connaissances puisque nous débattons des aspects techniques dans les deux chapitre suivants.

Selon la définition du Petit Larousse, indexer un document consiste à construire la "liste alphabétique des mots, des sujets, des noms apparaissant dans un ouvrage, une collection, etc. avec les références permettant de les retrouver".

Il s'agit donc d'un repérage lexical interne à un document ou un ensemble de documents [Crampes *et al.*, 2002]. Appliquée à un centre de documentation, l'indexation a pour but d'effectuer un classement approprié de documents afin de pouvoir par la suite les retrouver avec plus de facilité et les consulter. Quelques mots-clefs judicieusement choisis suffisent souvent à la tâche. La recherche actuelle de documents sur Internet déroge peu à la règle même si la quantité des documents et des index est d'une autre ampleur. Internet introduit cependant une différence majeure de nature. En effet, l'indexation traditionnelle est le fait d'un(e) documentaliste, à charge pour le lecteur potentiel de retrouver manuellement l'ouvrage recherché par l'utilisation de l'index. A l'inverse, la documentation en ligne est souvent indexée électroniquement. La recherche est alors déléguée par l'utilisateur à des robots, les moteurs de recherches. Il ne s'agit plus de retrouver un point précis sur un rayonnement, mais d'identifier des objets dont on ignore jusqu'à l'existence et dont l'emplacement importe peu.

Deux approches différentes sont représentatives des différents courants de l'indexation.

- L'indexation automatique ou « plein texte » qui est fondée sur les techniques de « fouille de textes » ou de « text-mining », dont la principale caractéristique est de tenter d'automatiser le processus d'indexation ;
- L'indexation manuelle par la formalisation ou discrétisation de l'information, dont la principale caractéristique est de tenter de conceptualiser l'information au moyen d'un processus humain.

Nous allons synthétiser ces deux approches.

5.1. Indexation manuelle

L'indexation manuelle consiste à demander à un opérateur spécialisé ou non, de représenter un document textuel ou un extrait de texte, dans une forme exploitable informatiquement. Un exemple d'indexation manuelle est proposée dans [Crampes *et al.*, 2000] qui propose de fonder celle-ci sur une ontologie du domaine d'étude.

Dans de nombreux domaines (TALN, Intelligence Artificielle, Productique, Multimédia, Biologie), on a vu apparaître, ces dernières années, le terme d'ontologie. Si la définition de ce terme semble parvenir à un consensus [Chabert-Ranwez, 2000; Iksal, 2002], son utilisation réelle est beaucoup plus opaque. De l'avis général, elle sert à uniformiser un langage ou tout du moins à désambiguïser un vocabulaire, et doit favoriser ainsi la communication entre les acteurs d'un projet. En fait, même si ce point est souvent passé sous silence, elle est surtout un formidable outil pour forcer les acteurs d'un projet à structurer leurs informations, même si son rôle s'arrête souvent là. Dans la communauté Web sémantique et Intelligence Artificielle, elle est aussi fréquemment utilisée pour indexer des documents dans la mesure où elle est directement inspirée des thésaurus dont les documentalistes usent abondamment dans leur classification.

L'ontologie représente deux notions importantes relativement à un domaine : les concepts du domaine et les relations entre ces concepts. On se référera à [Chabert-Ranwez, 2000] pour une explication détaillée de l'ontologie. [Crampes *et al.*, 2000; Crampes *et al.*, 2002] proposent à l'utilisateur de se limiter à l'utilisation des concepts et des relations contenues dans

l'ontologie pour représenter un texte. Il s'agit alors de parcourir le texte manuellement et d'associer au fur et à mesure une portion de texte à un concept de l'ontologie ou à une relation de cette ontologie entre plusieurs concepts. Le résultat obtenu est un modèle formel du texte de départ, exploité ensuite par des outils informatiques de rapprochement d'informations similaires en se basant sur l'ontologie du domaine [Planté, 2000]. Le texte lui-même n'est plus disponible en tant que tel.

5.2. Indexation automatique

Elle a pour but de représenter la connaissance tout en conservant toutes les nuances et toutes les subtilités que peut exprimer le langage naturel sans perte d'information et de connaissance.

L'indexation automatique comporte deux étapes :

- Définir un modèle de représentation de documents. Différents modèles existent que nous détaillerons dans le chapitre suivant. La plupart des modèles qui nous intéressent sont des modèles vectoriels. En général ces modèles sont fondés sur des mots clés ;
- Représentation du document par son modèle.

L'indexation automatique doit permettre à l'opérateur de s'affranchir d'une tâche manuelle d'indexation.

5.3. Comparaison des deux approches

Indexation manuelle et automatique sont deux voies opposées. L'indexation automatique tente de palier au principal défaut de l'indexation manuelle : le processus manuel long et fastidieux d'indexation. Cette raison est amplement justifiable pour favoriser l'indexation automatique quand il s'agit d'intégrer dans une base de connaissance un grand nombre de documents dans lesquels peu de structure repérable existe. Un autre argument également en faveur de l'indexation automatique, est qu'il est difficile de demander à des utilisateurs d'une base de connaissance de procéder à une indexation manuelle chaque fois qu'il souhaite intégrer un document dans la base de connaissance.

En revanche, l'indexation automatique a comme inconvénient principal d'être plus imprécise qu'une indexation manuelle. Dans l'indexation automatique c'est la machine qui décide des mécanismes d'indexation fondés souvent sur des règles donnant des résultats nécessairement moins précis que ce que ferait un opérateur. Si on a pour objet de construire un référentiel métier, alors cela peut être un investissement rentable au niveau de l'entreprise de miser sur une indexation manuelle reposant sur un modèle de type ontologie parce que la granularité de description de la base de connaissance sera plus fine.

6. Connaissance actionnable

La valorisation de l'information dans le cadre de la décision n'est pas un concept nouveau [Simon, 1980]. Nous avons décrit dans ce qui précède ce processus de valorisation de la donnée à la connaissance. La question suivante est : quand est-ce qu'une connaissance devient utile à l'action ? Nous allons introduire la notion de connaissance actionnable (CA).

Ce néologisme, comme le rappelle J.L. Le Moigne [Le_Moigne, 1998], a été introduit dans la littérature organisationnelle par D. Schön en 1983 (actionable knowledge), afin de dépasser le distinguo habituel entre savoir et savoir-faire, c'est-à-dire la séparation entre la composante épistémique (la connaissance) et la composante pragmatique (l'action). C. Argyris, qui dès 1974 a proposé avec D. Schön des modèles organisationnels pour l'action, parle en 1995 de "savoir pour agir" (knowledge for action) [Argyris *et al.*, 1978]. Il ressort de ces travaux que

c'est par la réflexion sur ses actions, sur ses savoirs d'expérience, que le praticien (le sujet connaissant engagé dans l'action) peut mieux prendre conscience des stratégies d'actions qu'il a élaborées, et donc pourra les améliorer. La traduction de savoirs tacites en savoirs d'action constitue le cœur même du processus d'apprentissage.

En développant le concept dual d'action intelligente, H.A. Simon avait attiré l'attention sur la légitimité d'une rationalité procédurale, qui ne se réduit pas à la logique et au calcul fondé sur des faits ayant valeur de vérité, et qui tient pour tout aussi satisfaisants les modes cognitifs que sont la délibération et l'argumentation. Il s'agit alors de concevoir et construire des connaissances qui soient représentations d'expériences, d'actions, de réflexions, de délibérations [Le_Moigne, 1998].

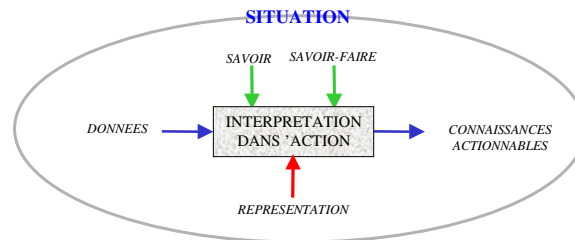


Figure III.2: Connaissance Actionnable (CA)

Dès lors, la gestion des connaissances actionnables n'est pas une encapsulation du savoir, mais un suivi dynamique d'un corpus de connaissances en expansion. Elle est définie comme processus de partage dynamique de connaissances [tacites] utiles à l'action [collective] [Penalva, 2000].

7. Instanciation de ces notions à notre problématique

Replaçons ces considérations générales dans notre processus de décision organisationnelle multicritère. Les notions abordées dans ce chapitre se rapportent bien sûr à la phase d'information du modèle de Simon, c'est-à-dire, pour éviter de retomber dans l'équivoque, à la phase d'apprentissage des connaissances nécessaires et utiles à la décision.

Force est donc de constater que l'abondance de l'information dans notre société a vraisemblablement changé le paradigme de la décision : on pouvait autrefois énoncer « Le pouvoir, c'est le savoir ». Seul un petit nombre d'élus qui « gouvernaient » savaient et pouvaient donc décider. Aujourd'hui, « tout le monde » est supposé avoir accès à l'information, le paradigme tombe... Le véritable problème, comme nous l'avons exposé dans ce chapitre, devient plutôt la capacité à gérer et traiter ce flux informationnel pour en tirer les connaissances pertinentes et utiles au bon moment. Il faut construire un *filtre de pertinence* pour alimenter l'évaluation et la sélection d'alternatives du processus de décision.

La décision est un choix ponctué d'une action. Il paraît ainsi légitime de postuler qu'un élément décisionnel ne saurait être une donnée, une information, ni même une connaissance, mais une connaissance actionnable.

La gestion des connaissances dans notre projet devient donc le partage et le contrôle de CAs qui soutiennent la décision. Il nous faut nous doter d'un outil informatique, un SGDC, qui assure cette tâche. La mémoire dynamique des CAs gérée par le SGDC doit être perçue comme le support de la logique de décision, la traçabilité des choix retenus. Le SGDC ayant pour ultime objet la décision, sa(ses) base(s) de connaissances doit(vent) faire l'objet d'une indexation spécifique sur laquelle nous allons revenir dans ce paragraphe. C'est cette indexation qui permettra de calibrer le *filtre de pertinence* recherché. C'est la condition sine qua non pour éviter toute surcharge cognitive.

Sur le plan technique, le SGDC sera assimilé, par la suite, à un outil intranet : serveur WEB s'interfaçant à des bases de données SQL et documentaires. La technologie du WEB permet de consulter facilement les documents ou informations, supportant ainsi les processus de partage et de délibération sur lesquels nous avons insisté auparavant. La collecte de documents ou d'informations, largement répartis dans l'organisation, est généralement une opération lourde et difficile. Cette solution technique s'affranchit de cet obstacle, en remplaçant la collecte par un dépôt à l'initiative de chaque « acteur - relais ».

Revenons donc plus précisément sur la notion de CA dans notre projet. Nous considérons que la CA est un savoir élémentaire qui se distingue d'une simple information par le fait qu'elle comporte une part d'interprétation liée à la personne qui l'énonce d'une part, et que cette interprétation tend à la rendre utile à l'action d'autre part. En faisant référence aux différentes considérations que nous avons évoquées dans ce chapitre, on écrira donc qu'une CA est définie tout à la fois comme :

- Une donnée informative jugée utile et qui prend du sens dans un contexte ;
- Un savoir élémentaire interprété par la personne qui l'énonce ;
- Une trace des raisonnements menés par les acteurs du processus décisionnel;
- Une entité minimale intelligible, de sens partageable et réutilisable dans le contexte.

Une CA est l'interprétation en termes de finalité par son auteur d'une information dans le cadre de son projet d'action : elle sera un élément d'argumentation et de rhétorique pour la justification des choix. Le SGDC assure le partage et le suivi des CAs dans l'organisation. Dans notre cadre multicritère, la formalisation de la CA s'explicite donc par les caractéristiques suivantes :

- Une *synthèse* (σ) qui est la valeur informative de la CA. C'est la partie descriptive qui peut préciser, par exemple, la source d'information, le sujet abordé, les références, hypothèses, le contexte, etc. Elle est en langage naturel ;
- Le *commentaire* (χ) qui est l'expression d'un jugement de valeur. Il correspond à une appréciation de l'auteur de la CA relativement au système de valeurs véhiculé par le processus décisionnel. Il est en langage naturel et est un élément de rhétorique potentiel ;
- Le système de valeurs de la décision se déclinant en critères d'évaluation dans notre évaluation multicritère, une CA est donc un jugement de valeur sur une alternative ϕ sous l'angle d'analyse d'un critère α . Les *coordonnées* (α, ϕ) de la CA dans la grille d'évaluation (Figure 3 du chapitre II, p critères en ligne, n alternatives en colonne) permettent de cartographier la base de connaissance du SGDC ;
- le *score partiel* (ρ) qu'elle porte. Ce score partiel doit en toute logique être cohérent avec le commentaire de la CA : il doit être la traduction numérique de l'appréciation en langage naturel portée par le commentaire de la CA ;
- Une *date* (τ) qui permet de repérer dans le temps la CA. Une interprétation peut évoluer dans le temps, un historique géré par le SGDC en permet le suivi.

Les trois chapitres qui suivent sont consacrés à la génération automatisée de CAs depuis des sources d'informations numériques. Il s'agit de construire des CAs à partir de l'analyse informatisée de bases de documents numériques. L'enjeu est la construction du *filtre* qui permettra cette extraction/génération. Il repose sur l'indexation automatique de fragments de textes numériques conformément à la description précédente de la notion de CA. Le

traitement du contenu de documents, textes, articles, etc., doit permettre d'identifier les éléments nécessaires à la construction de CAs, dont seront extraits objets et commentaires des CAs. L'automatisation de l'attribution d'un score au commentaire et de la cartographie de la CA dans la grille d'évaluation (*critère, alternative*) à une date donnée repose sur des techniques de classification ; c'est de ces techniques que les deux chapitres qui suivent nous entretiennent. Ainsi, cela revient à considérer que le contenu des textes d'une base documentaire numérique est indexé par des quintuplets $(\sigma, \chi, (\alpha, \varphi), \rho, \tau)$, qui définissent des CAs et donnent donc un accès direct aux éléments de décision disséminés dans le corpus documentaire.

Automatiser la phase d'apprentissage confère à notre système un fort niveau d'automatisation cognitive comme nous en avons débattu dans le chapitre II. Les deux chapitres qui suivent vont donc s'attacher à décrire l'indexation automatique de sources électroniques pour générer des CAs qui viendront alimenter les phase d'évaluation, de sélection et de révision du processus de décision global. Les techniques en jeu relèvent de la classification automatique.

Chapitre IV.

Représentations de corpus documentaires et techniques de classification

1. La représentation de corpus documentaires	86
1.1. Le mot comme unité linguistique.....	87
1.2. La phrase comme unité linguistique.....	87
1.3. La technique des « n-grammes ».....	87
1.4. Traitement préliminaire.....	88
1.5. Différents types de représentations textuelles	89
1.6. Conclusion.....	92
2. Les techniques de classification	92
2.1. Le représentant moyen d'une classe.....	94
2.2. Les classifieurs	94
2.3. Classifieur de Bayes	95
2.4. Autres classifieurs en traitement automatiques du langage	97
2.5. Extension de la sémantique distributionnelle.....	97
2.6. Techniques de segmentation automatique de texte	97
2.7. Validation	98
2.8. Mesure de performance	100
3. Conclusion.....	102

Automatiser la phase d'information du processus de décision consiste à alimenter la phase d'évaluation multicritère en CAs. Les scores des CAs, se référant à des critères, seront utilisés pour évaluer les alternatives ; les commentaires des CAs constitueront l'argumentaire sur lequel s'appuie la logique de décision. Nous cherchons donc à informatiser l'extraction de CAs depuis une base de textes jusqu'à la cartographie et au « scoring » des CAs pour l'évaluation multicritère. Très schématiquement, il s'agit d'abord d'extraire, de documents électroniques, des fragments de texte se rapportant à une alternative φ et un critère α , on associe alors un fragment de texte à une classe—un thème, défini par le couple (α, φ) . Dans un second temps, il faut procéder à l'attribution d'un score à cet extrait : là encore, la question se pose comme un problème de classification, où il faut mettre en correspondance le commentaire contenu dans la CA avec une valeur numérique dans une échelle d'intervalle de granularité plus ou moins fine sur $[0, 1]$. Ces classifications nécessitent plusieurs étapes dont le détail technique fera l'objet du chapitre suivant et permettent d'indexer la base de documents numériques en vue de l'évaluation multicritère comme nous en avons convenu au chapitre précédent.

Avant de proposer une solution technique dédiée à notre problématique, nous proposons une revue synthétique des outils de classification les plus usités dans le domaine de la représentation de corpus de connaissances.

1. La représentation de corpus documentaires

Nous présentons dans ce paragraphe plusieurs représentations de corpus documentaires correspondant à différentes méthodes qui sont néanmoins toutes basées sur des modélisations vectorielles.

Les modèles qui suivent sont d'usage très général, c'est-à-dire qu'ils n'ont pas été développés en vue d'une transcription des documents numériques manipulés en entités manipulables pour une finalité précise (cartographie, traduction, résumé automatique...) ; ils ont le plus souvent été conçus pour modéliser un corpus de connaissances. Cette démarche est à l'opposé de la nôtre où c'est l'évaluation multicritère, dont les critères sont déterminés a priori, qui définit complètement l'indexation du corpus de textes et rend par là même la représentation caduque pour tout autre usage : on retrouve bien ici le distinguo connaissance/connaissance actionnable. Ce n'est qu'a posteriori que les auteurs de ces représentations « polyvalentes », les ont employées pour des analyses de textes spécifiques.

Considérons le problème générique qui est d'affecter un texte à une catégorie—un thème par exemple. L'objectif est donc de mettre en correspondance textes et catégories. Un être humain, de part ses capacités cognitives, dégage rapidement le sens du texte et l'utilise pour rapprocher le texte d'un thème. Par conséquent, lorsque nous délégons cette tâche à un ordinateur nous devons représenter chacun des textes dans un formalisme permettant d'appréhender de manière plus ou moins fine sa sémantique si l'on veut avoir une classification pertinente.

Il est souvent admis que le sens d'un document peut être porté par un ensemble d'unités linguistiques particulières, caractéristiques plus ou moins élaborées issues de l'analyse du corpus documentaire [Besançon, 2001]. Chaque unité linguistique définit une « idée », et par conséquent, un ensemble de mêmes unités linguistiques est censé définir un sens identique. Les premières unités linguistiques à s'être imposées comme représentatives du sens sont le radical et le lemme des mots. La reconnaissance de ces unités linguistiques nécessite d'effectuer un prétraitement linguistique des mots du texte. Certaines unités linguistiques, plus

rudimentaires, ne nécessitent aucun traitement a priori, comme le « sac de mots » ou la phrase.

1.1. Le mot comme unité linguistique

Dans cette approche l'unité linguistique choisie est le mot tel qu'il apparaît dans le document. Cette approche ne nécessite aucun traitement préliminaire, chaque mot est extrait du texte en considérant des séparateurs entre chaque mot, comme l'espace, la virgule, la tabulation, le point et la ponctuation en général. L'approche par *sac de mots* est généralement associée à une procédure de filtrage. En effet, le nombre de mots caractérisant un corpus de documents peut rapidement atteindre une centaine de milliers. Il devient alors nécessaire, afin d'envisager un traitement analytique des textes, de ne conserver qu'un sous-ensemble de ces mots. Classiquement, le filtrage repose à la base sur les fréquences d'occurrence des mots dans le corpus. Nous verrons plus loin dans cette section les différentes techniques de filtrage que propose la littérature.

1.2. La phrase comme unité linguistique

Certaines approches utilisent non pas des mots, mais des groupes de mots, voire des phrases, comme éléments représentatifs du sens. L'intérêt d'utiliser une phrase ou un groupe de mots est double. Tout d'abord ce type d'unités linguistique possède une unité de sens plus complète qu'un simple mot de par la « contextualisation » des mots [Lewis, 1992]. En effet, l'association d'un ensemble de mots, même dans le désordre, offre davantage d'information sur le champ sémantique dans lequel on se trouve. De plus, ce type d'unités linguistiques définit une relation d'ordre entre les mots : un des avantages de cette représentation est qu'elle ouvre la porte à la gestion des cooccurrences de mots. En pratique, les expérimentations avec ce type de représentation ne sont, pour l'instant, pas très concluantes. En effet, si la phrase est une unité linguistique à forte valeur sémantique ajoutée, la fréquence d'apparition de groupes de mots ne permet pas d'offrir des statistiques fiables. Le grand nombre de combinaisons entre les mots engendre des fréquences trop faibles pour être exploitables.

1.3. La technique des « n-grammes »

Les n-grammes ont été introduits par Shannon en 1948 [Shannon, 1948]. En général le n-gramme se définit comme étant une séquence de n caractères consécutifs. Cependant il existe quelques variantes dans la littérature. En effet dans [Caropreso *et al.*, 2001], l'auteur définit le n-gramme comme étant une séquence de n "mots désuffixés" et dans [Cavnar *et al.*, 1994] l'auteur définit le n-gramme comme étant une séquence de n caractères non obligatoirement consécutifs. Nous nous contenterons ici de la notion de n caractères consécutifs. Ainsi, le mot : « porte » est composé des n-grammes suivants :

- bi-grammes : « _p », « po », « or », « rt », « te », « e_ »
- tri-grammes : « __p », « _po », « por », « ort », « rte », « te_ », etc.

De manière générale pour une chaîne de k caractères entourée de blancs, on génère $k + 1$ bi-grammes, $k + 1$ tri-grammes ou plus généralement $k + 1$ n-grammes [Cavnar *et al.*, 1994]. Pour un document quelconque l'ensemble des n-grammes est obtenu en déplaçant une fenêtre de n caractères sur le texte. Avant chaque déplacement de 1 caractère de la fenêtre, le n-gramme contenu dans la fenêtre est défini comme étant un des n-grammes du document. Une fois extraits tous les n-grammes d'un document, on définit par profil d'un document, la liste des n-grammes triés par ordre décroissant de leur fréquence d'apparition. Les avantages liés à l'utilisation des n-grammes sont multiples. Tout d'abord, les méthodes basées sur les n-grammes sont indépendantes de la langue. Elles ne nécessitent ni la segmentation des

documents en unités linguistiques, ni la mise en place de prétraitements tels que le filtrage, la désuffixation ou la lemmatisation. Les n-grammes sont aussi tolérants aux bruits (principalement les fautes de frappe quelles qu'elles soient) [Manning *et al.*, 1999]. Pour un mot mal orthographié dans un texte, avec l'utilisation des n-grammes, le mot n'est pas perdu dans sa totalité. Par exemple le mot « gamme » mal orthographié peut donner « game ». Les bi-grammes de ces deux mots donnent :

- « gamme » : « _g », « ga », « am », « mm », « me », « e_ »
- « game » : « _g », « ga », « am », « me », « e_ »

On remarque que malgré la faute de frappe, il n'y a qu'un bi-gramme qui diffère entre le mot correct et le mot mal orthographié.

1.4. Traitement préliminaire

1.4.1. Stemmatisation (radicalisation ou « stemming »)

L'utilisation de mots comme unité linguistique est possible, mais pose toutefois un certain nombre de problèmes. En effet, il existe plusieurs centaines de milliers de mots dans un corpus documentaire et associer un sens à chacun de ces mots n'est pas nécessairement des plus pertinents. En effet, beaucoup de mots ont des racines communes, des sens communs. Attribuer un sens différent à « chantaient » et « chantent » par exemple, relèverait d'une redondance sans pertinence sémantique.

La stemmatisation qui se nomme également radicalisation ou en anglais « stemming » est un prétraitement qui consiste à trouver la racine de chaque mot. C'est un traitement qui procède à une analyse morphologique du texte. Une implémentation très connue est celle de Porter [Porter, 1980]. Cet algorithme consiste principalement à effectuer une « désuffixation ». C'est un traitement fondé sur l'étude de la morphologie des mots. Il se base sur un dictionnaire de suffixes qui permet d'extraire le radical du mot.

La stemmatisation de : « Ces vases sont disproportionnés », grâce à l'algorithme de Porter donne : « Ce vas sont disproportion ». Une vraie radicalisation aurait donné : « Ce vase sont proportion ».

1.4.2. Lemmatisation

La lemmatisation nécessite une analyse plus poussée que la stemmatisation. La lemmatisation se fonde sur un lexique. Un lexique est un ensemble de lemmes, que l'on peut assimiler globalement aux entrées d'un dictionnaire tels que ceux que l'on trouve dans le commerce. L'objectif de la lemmatisation est d'associer à chaque mot une entrée dans le lexique. Or, l'analyse morphologique est insuffisante pour extraire les lemmes d'un texte car de nombreux mots de même graphie peuvent provenir de différents lemmes. Par exemple « offense » peut être considéré comme le lemme définissant la parole ou la faute qui blesse, ou comme le verbe "offenser" conjugué. Cette ambiguïté se résout en analysant la catégorie grammaticale du mot en question dans la phrase. La lemmatisation nécessite donc de réaliser une analyse supplémentaire : l'analyse syntaxique. La lemmatisation recouvre donc deux analyses regroupées sous le terme d'analyse morphosyntaxique. Depuis la fin des années 80, les lemmatiseurs sont capables d'associer à chaque mot d'un texte, son lemme, grâce à un étiqueteur morphosyntaxique (nom, verbe, adjectif, etc.) dont les taux de réussite avoisinent les 90% [Schmid, 1994]. La construction d'un lemmatiseur nécessite néanmoins un étiquetage de quelques milliers de mots.

1.4.3. La « stop-list »

Avant d'effectuer un des prétraitements précédents, il est usuel d'utiliser une « stop-list ». L'objectif de la « stop-list » est d'éliminer tous les mots ne participant pas activement au sens du document. La « stop-list » est une liste répertoriant tous les mots outils (pronoms, articles, etc.) et les mots trop fréquents pour être discriminants. D'un point de vue linguistique les mots outils sont par définition des mots "vides" de sens. Et, d'un point de vue statistique, les mots trop fréquents (et de distribution uniforme) ne sont d'aucune aide à un processus de catégorisation puisque non discriminants. De par sa nature ce prétraitement s'effectue donc en amont des autres prétraitements linguistiques et son objectif est d'éliminer toutes les unités linguistiques non discriminantes. L'inconvénient de la « stop-list » est qu'elle reste dépendante d'une langue donnée. La « stop-list » pose un autre problème plus fondamental. En effet, qui peut dire a priori que tel type syntaxique ou mot « outil » est vide de sens ? On risque dans certains cas d'omettre certaines unités linguistiques qui auraient permis d'apporter une aide précieuse à la tâche de classification. En particulier, il est peut être imprudent d'établir cette liste avant d'avoir fait le choix de la technique de classification qui peut parfois nécessiter toutes les informations disponibles.

1.5. Différents types de représentations textuelles

1.5.1. Le modèle vectoriel

Le rôle de la représentation textuelle est de projeter les documents au sein d'une représentation mathématique qui en permette le traitement analytique, tout en conservant au maximum la sémantique de ceux-ci. Pour réaliser cette projection la représentation mathématique généralement utilisée est l'utilisation d'un espace vectoriel comme espace de représentation cible.

La caractéristique principale de la représentation vectorielle est que chaque « brique de base du sens » ou unité linguistique (segment textuel) est associée à une dimension propre au sein de l'espace vectoriel. Deux textes utilisant les mêmes segments textuels seront donc projetés sur des vecteurs identiques. Le formalisme le plus utilisé pour représenter les textes est le formalisme vectoriel issu de [Salton, 1971] et [Salton, 1983]. Dans ce qui suit, nous allons décrire quatre représentations vectorielles différentes et insister sur leurs similarités mais aussi sur leurs différences qui confèrent à chacune des avantages spécifiques.

La différence principale des représentations que nous allons décrire est la détermination des dimensions de l'espace vectoriel. Chaque dimension est associée à des notions différentes en fonction de l'approche choisie. Dans certains cas il s'agit tout simplement de mots, dans d'autres il s'agit plutôt de concepts. Cette détermination des dimensions relève d'un processus d'indexation.

1.5.2. Le modèle vectoriel SMART

Dans [Salton, 1971] et [Salton, 1983], les documents sont représentés selon le modèle vectoriel dont l'implémentation la plus connue est SMART. Dans ce formalisme, chaque dimension de l'espace vectoriel correspond à un élément textuel, nommé terme d'indexation, préalablement extrait par calcul selon plusieurs méthodes préconisées par Salton. Cette action de sélection fait appel à un premier processus d'indexation. Ce processus répertorie toutes les unités linguistiques de tous les documents du corpus de connaissance.

Le processus d'indexation nécessite lui-même l'étape préliminaire consistant à regrouper des mots par familles syntaxiques grâce au processus de lemmatisation que nous avons décrit précédemment.

Salton a proposé plusieurs versions du processus de traitement préliminaire des unités linguistiques. Dans un premier temps ce traitement était une stématisation. Puis avec l'évolution des techniques informatiques, de nombreuses solutions sont apparues pour l'analyse morphosyntaxique des textes. Ainsi, dans une implémentation plus évoluée, les éléments retenus pour les traitements ultérieurs sont les mots racines ou lemmes ou encore « morphèmes lexicaux » du texte. Le traitement est alors une lemmatisation. Le processus d'indexation lui-même consiste alors à effectuer un simple inventaire complet de tous les lemmes du corpus en cours d'étude.

Ensuite, vient le processus de sélection des lemmes qui vont constituer ce que l'on peut nommer les « unités linguistiques » du domaine ou les dimensions de l'espace vectoriel de représentation du corpus documentaire. Dans ce formalisme, chaque dimension de l'espace vectoriel correspond à un segment textuel—l'unité linguistique, préalablement extraite du corpus documentaire. L'unité linguistique que nous avons introduite précédemment coïncide ici avec la notion de dimension de la représentation vectorielle. Une dimension peut donc, selon le cas, être un mot, un paragraphe, une lettre ou un assemblage de lettres. Les unités linguistiques sont choisies dans l'ensemble des lemmes en fonction de leur pouvoir de discrimination, elles constituent les termes d'indexation.

La sélection des termes d'indexation correspond donc à une réduction de la dimension de l'espace effectuée en fonction de critères de discrimination. D'une manière générale, les techniques de sélection des termes d'indexation font partie d'un domaine de recherche propre (« feature selection ») qui trouve des applications en classification et en recherche documentaire.

Le critère de sélection des termes d'indexation le plus utilisé est la fréquence en documents IDF. Il consiste à calculer le nombre de documents dans lesquels apparaît un lemme puis de prendre l'inverse de ce nombre. Le résultat obtenu est nommé IDF (Inverse Document Frequency). La fréquence en documents est, pour une unité linguistique et une collection de documents donnée, le nombre de documents la contenant. [Salton *et al.*, 1975] ont montré que, pour une collection de documents D , la sélection des unités linguistiques qui ont une fréquence en documents entre $|D|/100$ et $|D|/10$ génère le plus souvent un ensemble d'unités linguistiques ayant un pouvoir de discrimination satisfaisant pour la recherche documentaire. Ce critère de sélection est souvent utilisé car il est simple et rapide à mettre en œuvre et il ne nécessite pas de connaissances particulières autres que la donnée brute du corpus.

A ce stade, lorsqu'un lemme l_i est sélectionné, il constitue une dimension de l'espace vectoriel de représentation du corpus de textes. Par suite, un texte ou document D est représenté par un vecteur $D = (ft_1, \dots, ft_i, \dots, ft_{|V|})$ dans lequel V représente le *vocabulaire* ou l'ensemble des unités linguistiques sélectionnées ($|V|$, son cardinal) et ft_i , le nombre d'occurrences du lemme l_i dans le document D .

Cette technique, encore très utilisée aujourd'hui, est une vision purement statistique de l'extraction d'information qui a cependant prouvé son efficacité. L'unique prise en compte de la sémantique de l'information tient dans le processus de lemmatisation et dans le fait que les unités linguistiques retenues symbolisent les tendances macroscopiques du sens du corpus documentaire.

1.5.3. Le modèle LSI

Le modèle LSI (« Latent Semantic Indexing ») est une variante du modèle vectoriel standard qui tente de prendre en compte, pour les représentations des documents, la structure

sémantique des unités linguistiques, potentiellement implicite (i.e. latent), représentée par leurs dépendances cachées [Deerwester *et al.*, 1990].

Pratiquement, les techniques LSI utilisent la matrice (*unités linguistiques j x documents i*) du modèle vectoriel standard, dans laquelle chaque élément w_{ij} est le nombre d'occurrences de l'unité linguistique u_j dans le document D_i . Une décomposition en valeurs singulières (SVD) de cette matrice est effectuée et seuls les k premiers vecteurs propres sont pris en compte (k prend typiquement une valeur entre 100 et 300).

Mathématiquement, si M est la matrice « *unités linguistiques j x documents i* », la décomposition en valeurs singulières s'écrit $M = U \cdot \Sigma \cdot V^t$ avec U , V orthogonales ($U \cdot U^t = 1$, $V \cdot V^t = 1$) et Σ est diagonale. On approxime la matrice Σ par la matrice $\hat{\Sigma}$ réduite aux k premières dimensions (nulle sur les autres) : $\hat{M} = \hat{U} \cdot \hat{\Sigma} \cdot \hat{V}^t \approx U \cdot \Sigma \cdot V^t = M$.

Ce modèle permet donc de représenter les documents dans un espace réduit de dimension k . Il est toutefois important de noter que chacune des dimensions de l'espace de représentation final correspond à une « combinaison linéaire des unités linguistiques ». L'espace de représentation n'a donc pas pour support un ensemble de termes d'indexation, ce qui rend les dimensions relativement difficiles à interpréter directement.

Ce modèle permet également, de façon symétrique, de représenter les termes par des vecteurs qui sont une indication du profil d'occurrence du terme dans les documents. Cette propriété peut donc être utilisée pour établir une notion de similarité entre termes, ou encore représenter un document comme la moyenne des vecteurs représentant les termes qu'il contient.

1.5.4. La sémantique distributionnelle

Dans le modèle DSIR de [Besançon *et al.*, 2000], les documents sont représentés sous forme de vecteurs obtenus par un calcul de cooccurrence entre les termes d'indexation d'un corpus documentaire. Une matrice de cooccurrence M est calculée en prenant en compte l'apparition conjointe de 2 termes présents dans la même phrase. Les termes d'indexation constituant la base vectorielle sont sélectionnés de la même manière que pour Salton. Les vecteurs directeurs de Besançon sont la somme pondérée de 2 vecteurs : $\alpha V + (1-\alpha) M \cdot V$, où V est un vecteur mot-clés comme proposé par Salton, et $M \cdot V$ est le produit du vecteur V avec la matrice de cooccurrences M construite sur le corpus. Le facteur α est calculé expérimentalement.

Le résultat est une représentation des documents sous forme de vecteurs qui portent à la fois la représentation statistique des mots clés et la contribution des autres termes du corpus pour chaque mot clé. La représentation de chaque texte exprime à la fois des fréquences d'apparition de mots clés mais également les contributions des liens de cooccurrence du corpus documentaire dont le texte fait partie.

Dans ce modèle, le choix des termes d'indexation pour la tâche de classification diffère très peu de celui de Salton. Comme pour les unités linguistiques du modèle vectoriel de Salton, les termes d'indexation peuvent être choisis dans l'ensemble de toutes les unités linguistiques en fonction de leur fréquence en documents. Cette technique a l'avantage d'introduire le concept de « sème » directement relié à la cooccurrence des unités linguistiques et l'indexation prend ainsi une dimension sémantique plus prononcée.

1.5.5. Le modèle des vecteurs conceptuels

La représentation selon [Chauché, 1990; Jaillet, 2005], bien que se basant aussi sur un modèle vectoriel pour représenter les documents, reste fondamentalement différente de la

représentation de Salton. Les dimensions de l'espace vectoriel ne sont pas associées ici à des lemmes mais à des concepts [Chauché, 1990; Jaillet, 2005]. L'espace de concepts utilisé est celui du thésaurus Larousse [Larousse, 1992] composé de 873 concepts hiérarchisés en 4 niveaux. Par exemple, le mot "mélodie", défini par les concepts 741, 781 et 784 (respectivement, phrase, musique et chant) du thésaurus, sera représenté par un vecteur de dimension 873 dont toutes les composantes seront nulles sauf celles associées aux concepts 741, 781 et 784 qui seront toutes trois égales à $1/3$. La valeur accordée à chaque composante non nulle est identique car l'on ne fait aucune hypothèse sur la prédominance d'un concept particulier par rapport aux autres. Le vecteur est ensuite normalisé. Une phrase sera représentée comme une somme normalisée des vecteurs des mots qui la compose. Un texte sera par suite représenté par un vecteur calculé comme la somme normalisée des vecteurs de ses phrases, etc. Cette modélisation originale est à la fois statistique et sémantique : statistique car les coordonnées des vecteurs représentent des fréquences pondérées de concepts ; sémantique car l'espace vectoriel est un espace de concepts à forte valeur cognitive.

1.6. Conclusion

A travers les différentes représentations que nous avons présentées, nous voyons que chacune d'elles tente d'exprimer à sa façon le sens contenu dans les documents. Nous avons choisi de les présenter dans cet ordre parce que cela correspond à une évolution vers des représentations de plus en plus conceptuelles. En effet, la représentation de Salton est fondée sur la statistique de mots significatifs autrement dit des mots-clés lemmatisés ; la représentation LSA introduit l'abstraction des dimensions essentielles d'un corpus à la façon d'une analyse par composantes principales (ACP) via les vecteurs propres du domaine relatif au corpus ; la représentation de Besançon via les cooccurrences introduit la notion de sème dans l'indexation du corpus et enfin la représentation de Chauché se réfère aux concepts de la langue. Sans doute, cette dernière représentation est-elle plus proche que les autres de ce qu'un être humain entend classiquement par *sens*. En effet, le Thésaurus des concepts de la langue introduit les idées du langage et les structures qui les lient ; l'objet de cette entreprise ne doit rien au traitement automatique du langage mais correspond au besoin des hommes de constituer une structure sémantique de la langue.

2. Les techniques de classification

Nous n'avons pas la prétention dans ce paragraphe de passer en revue l'ensemble des techniques de classification mais plutôt de mentionner celles qui sont les plus utilisées en traitement automatique du langage.

Bien connue en apprentissage, la classification, appelée également *induction supervisée*, consiste à analyser de nouveaux candidats et à les affecter, en fonction de leurs caractéristiques ou attributs, à telle ou telle classe prédéfinie.

Dans [Denis *et al.*, 2000], [Gilleron *et al.*, 2000], les méthodes de classification ont pour but d'identifier les classes auxquelles appartiennent des objets à partir de certains traits descriptifs. Elles s'appliquent à un grand nombre d'activités humaines et conviennent en particulier au problème de l'aide à la prise de décision. Il s'agira, par exemple, d'établir un diagnostic médical à partir de la description clinique d'un patient, de donner une réponse à la demande de prêt bancaire de la part d'un client sur la base de sa situation personnelle, de déclencher un processus d'alerte en fonction de signaux reçus par des capteurs. Une première approche possible pour résoudre ce type de problème est l'approche " systèmes experts ". Dans ce cadre, la connaissance d'un expert (ou d'un groupe d'experts) est décrite sous forme de règles. Cet ensemble de règles forme un système expert qui est utilisé pour classer de

nouveaux cas. Cette approche, largement utilisée dans les années 80, dépend fortement de la capacité à extraire et à formaliser les connaissances de l'expert.

Nous allons considérer ici une autre approche pour laquelle la procédure de classification sera générée automatiquement à partir d'un ensemble d'exemples. Un *exemple* consiste en la description d'un cas avec la classification correspondante. Par exemple, on dispose d'un historique des prêts accordés avec, pour chaque prêt, la situation personnelle du demandeur et le résultat du prêt (problèmes de recouvrement ou non). Un système d'apprentissage doit alors, à partir de cet ensemble d'exemples, extraire une procédure de classification qui, au vu de la situation personnelle d'un client, devra décider de l'attribution du prêt. Il s'agit donc d'induire une procédure de classification générale à partir d'exemples. Le problème est donc un problème inductif, il s'agit d'extraire une règle générale à partir de données observées. La procédure générée devra classer correctement les exemples de l'échantillon mais surtout avoir un bon pouvoir prédictif pour classer correctement de nouvelles descriptions.

Les méthodes utilisées par les systèmes d'apprentissage sont très nombreuses et sont issues de domaines scientifiques variés. Les méthodes statistiques supposent que les descriptions des objets d'une même classe se répartissent en respectant une structure spécifique à la classe. On fait des hypothèses sur les distributions des descriptions à l'intérieur des classes et les procédures de classification sont construites à l'aide d'hypothèses probabilistes. La variété des méthodes vient de la diversité des hypothèses possibles. Ces méthodes sont appelées semi paramétriques. Des méthodes non paramétriques (sans hypothèse a priori sur les distributions) ont été également proposées en statistiques. Les méthodes issues de l'intelligence artificielle sont des méthodes non paramétriques. On distingue les méthodes symboliques (la procédure de classification produite peut être écrite sous forme de règles), des méthodes non symboliques ou adaptatives (la procédure de classification produite est de type "boîte noire"). Parmi les méthodes symboliques, les plus utilisées sont basées sur les arbres de décision ou la logique floue. Pour les méthodes adaptatives, on distingue deux grandes classes : les réseaux de neurones et les algorithmes génétiques. L'apprentissage automatique, dans une définition très générale, consiste en l'élaboration de programmes qui s'améliorent avec l'expérience.

Les méthodes d'apprentissage à partir d'exemples sont très utilisées dans la recherche d'informations dans de grands ensembles de données. En effet, l'évolution de l'informatique permet de nos jours de manipuler des ensembles de données de très grande taille (entrepôt de données). La problématique consiste à constituer un corpus représentatif du domaine de connaissances dans lequel on opère, et à trouver les règles qui le régissent, ou bien à constituer un modèle opérationnel de ce corpus. Un modèle permet d'une part pour un observateur humain de comprendre l'image du corpus, et d'autre part pour la machine de se fonder sur ce modèle pour prédire le comportement à adopter quand un nouveau *candidat* à la classification arrive.

Pour [Witten *et al.*, 2005], l'apprentissage est un ensemble de techniques pour découvrir et décrire des structures et modèles à partir de données afin d'aider à la compréhension de ces données et effectuer des prédictions fondées sur les modèles identifiés.

Les techniques d'apprentissage liées au problème de classification sont tout d'abord liées à l'existence a priori de classes. C'est ce que l'on nomme l'apprentissage supervisé, car comme le dit [Witten *et al.*, 2005] « en un sens, la méthode opère par supervision du fait que l'on fournit a priori la solution réelle pour chacun des échantillons de l'ensemble d'apprentissage ». Le but est alors de fournir un modèle de chaque classe de l'ensemble d'apprentissage. Nous citerons plusieurs façons d'y parvenir :

- Calculer un représentant moyen de chaque classe, et utiliser une distance à ce représentant ;
- Le classifieur de Bayes ;
- Le classifieur de Bayes Multinomial ;
- La technique des Arbres de décision ;
- Les SVP (Support Vector Machine).

2.1. Le représentant moyen d'une classe

Une première méthode de classification utilise la technique du vecteur moyen. Dans cette technique, il s'agit de trouver un représentant de la classe. Plusieurs méthodes existent. En particulier, signalons celle qui consiste à calculer le barycentre des individus de la classe et à choisir comme représentant de la classe, l'individu le plus proche de ce barycentre. C'est en particulier le cas du modèle LSI.

Les modèles de [Besançon *et al.*, 2000; Chauché, 1990; Jaillet, 2005] représentent également chaque classe par son vecteur moyen. Celui-ci est calculé en effectuant la somme normalisée des vecteurs représentant les textes d'une même classe du corpus d'apprentissage. On obtient ainsi un vecteur unique qui ne correspond à aucun des documents du corpus, mais qui est le prototype d'une classe.

2.1.1. Distance au représentant moyen

Pour associer tout nouveau document à une des classes prédéfinies, la notion de distance est introduite. Il s'agit de trouver quelle est la classe la plus proche du nouveau document. La distance la plus utilisée et la plus connue est « cosinus » qui est le cosinus de l'angle formé entre deux vecteurs. Les principales distances utilisées en catégorisation sont le cosinus et la distance euclidienne :

$$\text{Distance du cosinus (« cosinus ») : } \cos(v_1, v_2) = \frac{v_1 \cdot v_2}{\|v_1\| \cdot \|v_2\|} = \frac{\sum_{k=1}^n (v_{1k} \cdot v_{2k})}{\sqrt{\sum_{k=1}^n v_{1k}^2 \sum_{k=1}^n v_{2k}^2}} \quad \text{IV-1}$$

$$\text{Distance euclidienne (L2) : } L2(v_1, v_2) = \|v_1 - v_2\|_2 = \sqrt{\sum_{k=1}^n (v_{1k} - v_{2k})^2} \quad \text{IV-2}$$

[Besançon *et al.*, 2000; Chauché, 1990] comme [Besançon *et al.*, 2000; Chauché, 1990] ont recours à la mesure du cosinus. Cette mesure normalisée ne fait entrer en compte que l'angle que forment les deux vecteurs v_1 et v_2 .

2.1.2. Recherche de l'appartenance à une classe

Dans la méthode du vecteur moyen, l'opération qui permet d'associer un texte (candidat) à une classe est donc la suivante : les classes sont représentées par leur vecteur moyen respectif et pour une distance donnée, il suffit de déterminer le vecteur caractéristique de l'individu à classer, puis de calculer les distances de cet individu au vecteur moyen de chaque classe. La classe d'appartenance de l'individu est alors celle dont la distance est la plus faible.

2.2. Les classifieurs

La seconde façon d'effectuer une classification est de trouver les caractéristiques de chaque classe et d'y associer une fonction d'appartenance. Parmi les méthodes utilisant ce procédé

nous pouvons citer les arbres de décision [Crawford *et al.*, 1991; Quinlan, 1993], les classifieurs de Bayes [Wang *et al.*, 2003], la méthode des SVM (Support Vector Machine) [Joachims, 1998], etc. La classification n'est plus alors fondée directement sur la notion de distance, mais sur un processus plus élaboré que nous allons voir dans les paragraphes suivants. A ce sujet nous avons présenté dans [Planté *et al.*, 2005b] une approche de classification fondée sur les arbres de décision destinée au traitement de notre problématique.

2.3. Classifieur de Bayes

2.3.1. Le classifieur de type Naïve Bayes (CNB)

Le classifieur de type Naïve Bayes est un catégoriseur de type probabiliste fondé sur le théorème de Bayes. Si l'on considère $v_j = (v_{j1}, \dots, v_{jk}, \dots, v_{jd})$ un vecteur de variables aléatoires représentant un document d_j , la probabilité que ce dernier appartienne à la catégorie

$$c_i \in C \text{ (selon le théorème de Bayes) est définie par : } P(c_i | v_j) = \frac{P(c_i)P(v_j | c_i)}{P(v_j)} \quad \text{IV-3}$$

Chaque variable aléatoire v_{jk} du vecteur v_j représente l'occurrence de l'unité linguistique k retenue pour la classification dans le document d_j . v_{jk} prend donc les valeurs 0 ou 1.

La procédure de classification par le CNB est d'affecter d_j à la catégorie $c_{\max}$ dont la probabilité est la plus élevée, c'est-à-dire : $c_{\max} = \arg \max_{c_i \in C} P(c_i | v_j)$.

L'approximation de la probabilité conditionnelle $P(v_j | c_i)$ s'effectue alors sur l'ensemble d'apprentissage en faisant l'hypothèse que les v_{jk} sont indépendantes ($P(v_j | c_i) = \prod_k P(v_{jk} | c_i)$). L'approximation de $P(c_i)$ quant à elle est définie de la façon suivante :

$$\hat{P}(c_i) = \frac{\text{nombre de documents} \in c_i}{\text{nombre total de documents}} \quad \text{IV-4}$$

2.3.2. Le modèle multinomial

Le modèle multinomial de Bayes est un classifieur de type Naïve Bayes auquel est appliquée une loi de probabilité multinomiale.

Dans ce modèle chaque document est représenté par une séquence aléatoire de termes d'indexation. De ce fait, la taille du document et le nombre d'occurrences des unités linguistiques qu'il contient rentrent en compte. La probabilité d'avoir le mot clé m_k présent dans un document appartenant à la classe c_i est définie par $P(m_k | c_i)$. La probabilité d'avoir n fois le mot clé m_k dans un document appartenant à la classe c_i est définie par $P(m_k | c_i)^n$. L'approximation de la probabilité d'apparition d'un texte de longueur L ayant pour vecteur v_j et appartenant à la classe c_i est donc :

$$\hat{P}(v_j | c_i, L) = \prod_k P(m_k | c_i)^{ft(m_k | v_j)} \quad \text{IV-5}$$

où $ft(m_k | v_j)$ représente le nombre de fois où le mot clé m_k est présent dans le document de vecteur v_j . Cette approximation est la base du modèle multinomial et c'est elle qui permettra d'évaluer la probabilité que l'on recherche $\hat{P}(c_i | v_j)$.

2.3.3. Le classifieur de type Naïve Bayes multinomial (CNBM)

Si l'on applique le théorème de Bayes sur la probabilité $P(c_i | v_j, L)$, on a :

$$P(c_i | v_j, L) = \frac{P(v_j | c_i, L)P(c_i | L)}{\sum_{c'_i \in C} (P(v_j | c'_i, L)P(c'_i | L))} \quad \text{IV-6}$$

où $P(v_j | c_i, L)$ est la probabilité d'avoir v_j élément de c_i connaissant sa longueur L ; $P(c_i | L)$ correspond à la probabilité a priori qu'un document de longueur L appartienne à c_i . Par la suite nous considérerons que la catégorie d'un document ne dépend pas de sa longueur et, par conséquent : $P(c_i | L) = P(c_i)$. L'estimation de $P(v_j | c_i, L)$ est plus complexe. A moins que la base d'apprentissage soit infinie, il est impossible d'estimer cette probabilité sans faire d'hypothèse simplificatrice. En effet à cause du nombre considérable de documents—et donc de vecteurs différents qu'il est possible d'obtenir à partir des termes d'indexation—une évaluation empirique ne peut être envisagée pour $P(v_j | c_i, L)$. Néanmoins, si l'on considère que les termes d'indexation apparaissent de manière indépendante et que la probabilité d'apparition d'un terme d'indexation ne dépend que de la catégorie associée au document et pas de la taille de celui-ci alors l'équation IV-5 est une approximation valide de $P(v_j | c_i, L)$.

La dernière étape consiste à approximer $P(m_k | c_i)$. Dans le modèle multinomial, on utilise généralement l'approximation de Laplace (estimateur de Laplace) [Vapnik, 1982] :

$$\hat{P}(m_k | c_i) = \frac{1 + ft(m_k | c_i)}{|V| + \sum_{m_k \in V} ft(m_k | c_i)} \quad \text{IV-7}$$

Avec :

- $|V|$ le nombre de termes d'indexation du modèle ;
- $ft(m_k | c_i)$ le nombre de fois où m_k apparaît dans des documents de la catégorie c_i .

En combinant IV-5 et IV-6 et connaissant IV-7, on obtient une approximation de $\hat{c}_{\max}$:

$$\hat{c}_{\max} = \arg \max_{c_i \in C} \frac{\hat{P}(c_i) \prod_k \hat{P}(m_k | c_i)^{ft(m_k | v_j)}}{\sum_{c'_i \in C} \hat{P}(c'_i) \prod_k \hat{P}(m_k | c'_i)^{ft(m_k | v_j)}} \quad \text{IV-8}$$

Il est aussi possible d'utiliser le logarithme afin de simplifier les produits par des sommes. Cette fonction étant monotone et strictement croissante, elle n'influe pas sur le résultat de $\hat{c}_{\max}$. L'équation IV-8 devient donc :

$$\hat{c}_{\max} = \arg \max_{c_i \in C} \left(\frac{\log \hat{P}(c_i) + \sum_k ft(m_k | v_j) \log \hat{P}(m_k | c_i)}{\log \left(\sum_{c'_i \in C} (\hat{P}(c'_i) + \sum_k ft(m_k | v_j) \log \hat{P}(m_k | c'_i)) \right)} \right) \quad \text{IV-9}$$

La phase d'apprentissage sur le corpus de documents permet de faire les comptages nécessaires pour estimer les probabilités qu'exige l'application de IV-9. La classification consiste alors à calculer le vecteur v_j , les $ft(m_k | v_j)$, puis à appliquer IV-9 pour attribuer le document à la classe correspondante à $\hat{c}_{\max}$.

2.4. Autres classifieurs en traitement automatiques du langage

Parmi les classifieurs que l'on retrouve le plus souvent en traitement automatique du langage, on peut citer les arbres de décision. Un arbre de décision est une représentation graphique d'une procédure de classification. Les nœuds internes de l'arbre sont des tests sur les champs ou attributs, les feuilles sont les classes. Lorsque les tests sont binaires, le fils gauche correspond à une réponse positive au test et le fils droit à une réponse négative. Pour classer un enregistrement, il suffit de descendre dans l'arbre selon les réponses aux différents tests pour l'enregistrement considéré.

On peut citer également la méthode SVM (Support Vector Machines) [Burges *et al.*, 1995; Christianini *et al.*, 2000] qui est un algorithme d'apprentissage supervisé pour les problèmes de classification à 2 classes.

Les principes de ces deux techniques sont donnés en annexe.

2.5. Extension de la sémantique distributionnelle

Qu'advient-il du processus de classification lorsque le document n'est plus modélisé par un vecteur de fréquences ? C'est le cas de la représentation par sémantique distributionnelle du paragraphe 1.5.4.

Une extension de cette technique est caractérisée par le fait que chaque texte y est représenté par sa matrice de cooccurrence. Une classe peut aussi être représentée par une matrice de cooccurrence : elle est construite à partir des matrices des documents de la classe sur la base d'apprentissage. Chaque matrice est ensuite normalisée. Enfin, la classification consiste à mesurer la distance d'un texte à une classe, autrement dit la distance entre deux matrices. Plusieurs mesures de distances existent. La solution la plus courante est de calculer la différence entre les deux matrices, puis de calculer la plus grande valeur propre de cette matrice différence. La classe d'appartenance du document est alors celle pour laquelle la distance matricielle est la plus faible.

2.6. Techniques de segmentation automatique de texte

Un problème de classification particulier et récurrent en traitement automatique du langage est le découpage d'un document en catégories ou thématiques. Dans un document, il s'agit d'extraire des fragments de textes relativement à des thématiques définies au préalable.

Le principe consiste à trouver les ruptures thématiques dans un texte, c'est-à-dire le changement présumé de sujet dans le texte, puis à associer chacune de ces parties aux thématiques. Généralement ce type de tâche est effectué manuellement, mais aujourd'hui des techniques informatiques ont commencé à voir le jour pour résoudre de façon automatique cette tâche.

Nous citerons trois techniques abordant ce problème :

- la méthode « texttiling » de Marty Hearst. Le principe consiste à repérer les unités d'information dans un texte que l'on peut assimiler à des segments thématiquement homogènes. La similarité entre deux phrases adjacentes tient compte des mots communs entre les phrases, du nombre de nouveaux mots apparus dans la seconde, des occurrences de termes sémantiquement proches. La notion est étendue à deux segments adjacents. Lorsqu'il y a rupture de cette mesure de similarité entre deux segments, alors il y a changement thématique. La méthode du « Text-tiling » propose un algorithme pour résoudre cette question [Hearst, 1994] [Hearst, 1997].

- la méthode du « rhétoric parsing ». Dans cette méthode, on cherche à représenter la rhétorique du texte afin d'en extraire les idées les plus marquantes. L'approche exploite les théories du discours pour trouver les relations entre les propositions des phrases (noyau et nucléide). La méthode identifie des unités textuelles, puis les relations rhétoriques qui les lient [Marcu, 1997].
- la méthode d'Exploration Contextuelle. La méthode d'Exploration Contextuelle [Minel *et al.*, 2001] vise à identifier les connaissances linguistiques dans le texte en les restituant dans leurs contextes et en les organisant en tâches spécialisées. L'approche développée est fondée sur la construction manuelle d'une base de données de marqueurs linguistiques et une expression de règles d'exploration contextuelle. Ces règles appliquées aux phrases du texte source vont filtrer les informations sémantiques indépendantes du domaine avec les étiquettes sémantiques hiérarchisées comme : énoncés structurants, définition, causalité, etc. La stratégie de sélection des unités saillantes est fonction des besoins des utilisateurs (filtrage d'informations). Le résultat est un ensemble de phrases représentant les différentes informations saillantes du texte.

2.7. Validation

La validation est une phase indispensable à tout processus d'apprentissage. Elle consiste à vérifier que le modèle construit sur la base d'apprentissage est un classifieur performant. Dans ce cas, *performant*, signifie qu'il permet de classer tout individu avec le minimum d'erreurs possible.

Les méthodes de validation vont dépendre de la nature de la tâche et du problème considéré. Nous distinguerons deux modes de validation : statistique et par expertise. Pour certains domaines d'application (le diagnostic médical, par exemple), il est essentiel que le modèle produit soit compréhensible. Il y a donc une première validation du modèle produit par l'expert, celle-ci peut être complétée par une validation statistique sur des bases de cas existantes.

Pour la validation statistique, la première tâche à réaliser consiste à utiliser les méthodes élémentaires de statistique descriptive. L'objectif est d'obtenir des informations qui permettront de juger le résultat obtenu, ou d'estimer la qualité ou les biais des données d'apprentissage.

Pour la classification supervisée, la deuxième tâche consiste à décomposer les données en plusieurs ensembles disjoints. L'objectif est de garder des données pour estimer les erreurs des modèles ou de les comparer. Il est souvent recommandé de constituer :

- Un ensemble d'apprentissage ;
- Un ensemble de test ;
- Un ensemble de validation.

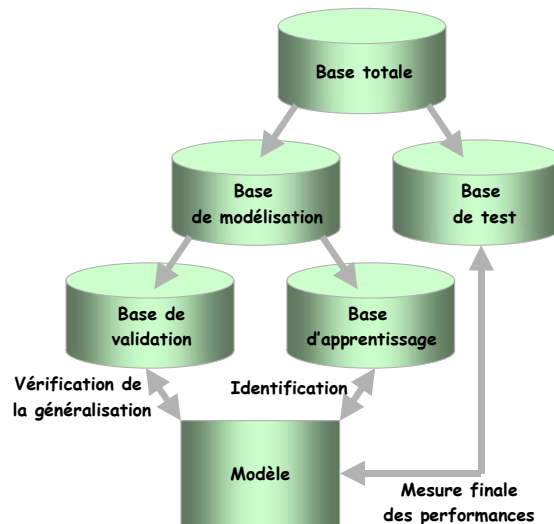


Figure IV-1 : Processus de validation par le test

L'ensemble d'apprentissage permet de générer le modèle, l'ensemble test permet d'évaluer l'erreur réelle du modèle sur un ensemble indépendant évitant ainsi un biais d'apprentissage. Lorsqu'il s'agit de tester plusieurs modèles et de les comparer, on peut sélectionner le meilleur modèle selon ses performances sur l'ensemble de validation et ensuite évaluer son erreur réelle sur l'ensemble de test (Figure IV-1).

Lorsqu'on ne dispose pas de suffisamment d'exemples, on peut se permettre d'apprendre et d'estimer les erreurs avec un même ensemble par la technique de validation croisée ou du *Bootstrap*.

2.7.1. Validation croisée

La validation croisée ne construit pas de modèle utilisable mais ne sert qu'à estimer l'erreur réelle d'un modèle selon l'algorithme suivant (Figure IV-2) :

Validation croisée ($S;x$) :

// S est un ensemble, x est un entier

Découper S en x parties égales $S_1, \dots, S_x$

Pour i de 1 à x

Construire un modèle M avec l'ensemble $S - S_i$

Evaluer une mesure d'erreur e_i de M avec S_i

Fin Pour

Retourner l'espérance mathématique des mesures des erreurs

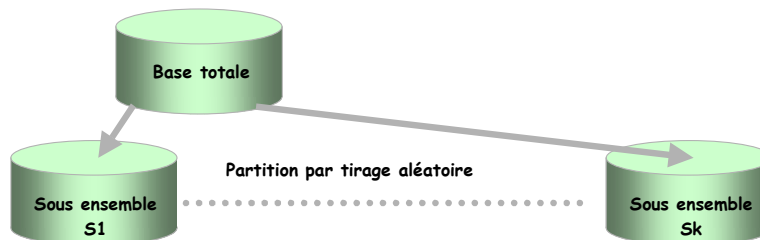


Figure IV-2 : Processus de validation croisée

Si la taille des S_i est de un individu, on parle alors de validation par « leave one out ».

En général le nombre x de parties est fixé à 10.

2.7.2. Bootstrap

Etant donné un échantillon S de taille n , on tire avec remise un ensemble d'apprentissage de taille n (un élément de S peut ne pas appartenir à l'ensemble d'apprentissage, ou y figurer plusieurs fois), l'ensemble test est S (Figure IV-3).

L'estimation de l'erreur réelle est alors l'espérance mathématique des erreurs obtenues pour un certain nombre d'itérations de l'algorithme d'apprentissage.

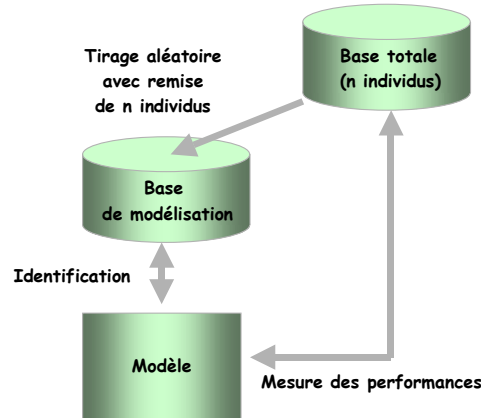


Figure IV-3 : Processus de Bootstrap

Ces deux méthodes fournissent de bons estimateurs de l'erreur réelle mais sont très coûteuses en temps de calcul. Elles sont très utiles pour les petits échantillons.

2.8. Mesure de performance

Afin de valider correctement la procédure de classification, nous utilisons des mesures de performances sur les résultats de la classification.

L'efficacité peut se définir selon plusieurs critères. Les deux critères généralement utilisés pour évaluer un processus de catégorisation sont : la précision et le rappel.

Nous allons donner une définition formelle de ces deux mesures. Pour cela nous allons définir les quatre notions suivantes pour une classe i :

- VP_i est l'ensemble des textes de la classe i bien classés ;
- FP_i est l'ensemble des textes assignés par erreur à la classe i ;
- FN_i est l'ensemble des textes de la classe i non classés i par le classifieur ;
- VN_i est l'ensemble des textes n'appartenant pas à la classe i et identifiés comme tels.

On peut visualiser ces notions sur le tableau suivant :

Classe i		Classement de l'Expert	
		VRAI	FAUX
Classement du Système	POSITIF	VP_i	FP_i
	NEGATIF	FN_i	VN_i

Tableau IV-1 : les quatre possibilités d'un classifieur

2.8.1. Précision

La précision est, pour une classe, parmi tous les documents que le système a attribués à cette classe ceux que l'expert a confirmé appartenir à cette classe. Dit autrement cette mesure indique la capacité du classifieur à classer correctement les documents.

Formellement la précision s'exprime de la façon suivante : $P_i = \frac{VP_i}{VP_i + FP_i}$ IV-10

Ce ratio permet de savoir en particulier si le classifieur, quand il classifie des documents, n'affecte pas trop de documents à une classe par erreur.

2.8.2. Rappel

Le rappel est, pour une classe, le rapport entre le nombre de documents attribués à la classe par le système sur le nombre de documents que l'expert a attribués à la classe. Dit autrement cette mesure indique la capacité du classifieur à classer correctement l'intégralité des documents.

Formellement le rappel s'exprime de la façon suivante : $R_i = \frac{VP_i}{VP_i + FN_i}$ IV-11

Le rappel permet de savoir si le classifieur est performant dans sa capacité à extraire de l'ensemble des documents ceux qui sont attribués à la classe en cours d'analyse tout en ayant peu d'oublis.

2.8.3. Notion de pertinence

La pertinence est la capacité du classifieur à bien classer les éléments qui lui sont soumis. C'est la somme du nombre de documents attribués à chaque classe par le système sur le nombre de documents que l'expert a attribués à chaque classe.

Formellement la pertinence s'exprime de la façon suivante :

$$Pt_i = \frac{VP_i + VN_i}{VP_i + FN_i + FP_i + VN_i} \quad \text{IV-12}$$

2.8.4. F-Mesure

Une fois définies les deux notions de rappel et de précision, plusieurs indicateurs de synthèse ont été imaginés. Le plus couramment employé est la F-Mesure ou F-beta de [Rijsbergen, 1979] :

$$F_\beta^i = \frac{(\beta^2 + 1)P_i R_i}{\beta^2 P_i + R_i} \quad \text{IV-13}$$

Cette mesure fusionne précision et rappel et donne une évaluation de synthèse de la classification. Le coefficient β indique le poids que l'on souhaite donner à la précision par rapport au rappel.

2.8.5. F score

Le Fscore est la mesure F-beta vue au paragraphe précédent en attribuant 1 au paramètre β , soit la mesure suivante [Rijsbergen, 1979] :

$$F_1^i = \frac{2P_i R_i}{P_i + R_i} \quad \text{IV-14}$$

Lorsque β vaut 1, le F-score est la moyenne harmonique de la précision et du rappel. Pour que le F-score soit important les deux composantes doivent être importantes. Cette mesure est fréquemment utilisée quand il est important d'avoir un équilibre entre précision et rappel.

Dans notre problématique, nous aurons souvent des classifieurs à deux classes comme nous le verrons dans le chapitre suivant. Ainsi, un candidat sera affecté à l'une ou l'autre classe. De ce fait, une faible précision sur l'une des classes donnera un faible rappel sur l'autre classe et réciproquement. Il est donc important dans notre cas d'avoir une mesure qui donne une égale importance à la précision et au rappel. Nous retiendrons donc la valeur $\beta = 1$ pour évaluer la performance de nos classifieurs.

Des aménagements à la mesure F-score peuvent être envisagés pour affiner et augmenter l'exigence de résultat de la classification [Nakache *et al.*, 2005]. Le F-score est une mesure largement employée qui permet de comparer les approches sur des benchmarks.

3. Conclusion

L'objectif que nous nous sommes fixés à la fin du chapitre précédent consiste à extraire d'un corpus documentaire des connaissances actionnables. Ces CAs sont structurées conformément à une grille d'évaluation multicritère : une CA est donc indexée par l'alternative ϕ et le critère d'évaluation α auxquels elle se rapporte. Par ailleurs, la CA est porteuse d'un score ρ qui n'est autre que la transcription numérique du commentaire de la CA, lui-même exprimant un jugement de valeur quant à la satisfaction du critère α par l'alternative ϕ : il s'agit donc de mettre en correspondance un jugement de valeur en langage naturel avec une échelle numérique (discrète). Ces deux points peuvent être interprétés comme des problèmes de classification dans le domaine du traitement automatique du langage naturel (TALN). Nous aurons donc recours aux outils de ce chapitre pour résoudre les problèmes de TALN que soulève notre problématique. Finalement, les classifieurs permettront d'indexer le corpus des connaissances utiles à la décision par ces caractéristiques (ϕ , α , ρ) et il sera alors possible d'avoir un accès direct aux CAs pour raisonner un choix, justifier d'une préférence.

Dans le chapitre qui suit nous revenons plus en détail sur les choix retenus tant pour la représentation des connaissances que pour la classification. La représentation vectorielle des documents dans un espace de mots-clés est le modèle de base que nous avons retenu. Comme nous l'avons signalé, les espaces engendrés par les mots-clés ne sont généralement pas viables de par leur dimension excessive. La spécificité de notre problématique nous a permis de construire les filtres adéquats pour diminuer significativement l'espace de représentation du corpus documentaire.

Si l'on prend l'exemple de la classification de dépêches selon des rubriques comme le cinéma, le sport ou l'économie, les catégories thématiques sont « très éloignées », les méthodes SMART, LSA, les vecteurs conceptuels se prêtent bien à la résolution d'un tel problème [Jaillet *et al.*, 2004]. Dans notre problématique, les catégories ne sont pas aussi démarquées les unes des autres. Les critères d'évaluation (si l'on reprend l'exemple du cinéma : scénario, réalisation, etc.) sont des concepts qui peuvent être sémantiquement beaucoup plus proches que ne le sont « cinéma » et « économie » ; ce qui demande une granularité de description plus fine que le thésaurus des vecteurs conceptuels. Aucune des méthodes que nous avons exposées dans ce chapitre n'a donné de résultats concluants sur notre problématique. Dans le chapitre qui suit, nous nous efforçons donc de montrer quels ont été les compromis retenus pour borner l'espace de représentation tout en proposant une granularité de description suffisamment fine pour modéliser des notions thématiques voisines (nos critères).

Chapitre V.

Extraction automatisée de CAs dans un processus d'évaluation multicritère

*« Sans la liberté de blâmer,
il n'est point d'éloge flatteur ».*
Le Mariage de Figaro (V, 3), 1784
Beaumarchais

1. Description générale de la chaîne de traitement.....	104
2. Corpus documentaire et base d'apprentissage	106
3. Représentation vectorielle du corpus de documents	107
3.1. Index et représentation vectorielle	108
3.2. Constitution de l'ensemble d'apprentissage de chaque classifieur	109
3.3. Réduction de l'index dédié à un classifieur	111
3.4. Ajout des synonymes	112
4. La classification des textes.....	113
4.1. le classifieur filtrage des textes évaluatifs et narratifs.....	113
4.2. Le classifieur « classification critère ».....	114
4.3. Attribution d'un score au commentaire d'une CA	115
4.4. Score d'une CA	117
5. Quelques justifications	117
6. Conclusion.....	119

1. Description générale de la chaîne de traitement

La chaîne de traitement que nous présentons ci-après détaille le processus d'extraction de CAs à partir d'une base de documents utilisés dans le cadre d'un processus d'évaluation multicritère comme nous l'avons présenté dans [Plantié *et al.*, 2005b]. Les Connaissances utiles à l'Action (la décision multicritère ici) sont dans notre cas des jugements de valeur à l'égard de n alternatives sous l'angle d'analyse de p critères d'évaluation.

L'extraction de CAs consiste donc :

- à extraire d'une base documentaire tout fragment de texte faisant référence à la grille (critères X alternatives) ;
- à déterminer dans ce fragment de texte ce qui relève d'un simple résumé informatif (*synthèse* σ —information objective) et ce qui constitue un jugement de valeur ou une appréciation (*commentaire* χ —interprétation subjective pour le référentiel d'évaluation) ;
- à cartographier ce fragment de texte dans la grille d'évaluation (autrement dit, lui attribuer des coordonnées (critère α , alternative φ)) ;
- à mettre en correspondance cette appréciation exprimée en langage naturel avec un score p relatif à une échelle d'évaluation numérique (on supposera qu'il s'agit d'une échelle de ratio ou d'intervalle) ;
- à enregistrer et dater (τ).

Suite à ce traitement, la CA est alors repérée dans le référentiel du SGDC par le quintuplet $(\sigma, \chi, (\alpha, \varphi), p, \tau)$. La base de documents est alors indexée en CAs pour le processus de décision relativement au référentiel des n alternatives et p critères d'évaluation. On peut alors interroger la base sous la forme de requête du type « CAs / CA(p) $\geq$ *score* & Alternative = φ & Critère = α » : on aura ainsi toutes les évaluations supérieures ou égales à *score* de l'alternative φ au regard du critère α ; CA(χ) constituera un élément rhétorique du score attribué et pourra être utilisé ultérieurement pour justifier de la logique de décision globale ; CA(σ) resitue le contexte de l'évaluation.

Cette indexation d'une base documentaire en CAs pour un processus d'évaluation multicritère nécessite plusieurs traitements qui reposent essentiellement sur des techniques de classification comme décrites au chapitre précédent. Les documents sont soit automatiquement téléchargés depuis le net ou rentrés par des utilisateurs dans une base à une date donnée.

Prenons un exemple simple. Supposons qu'un cyberconsommateur souhaite acheter un appareil photo sur le net. Comme nous l'avons vu au chapitre III, les sites de recommandation sont dédiés à soutenir, gérer et automatiser les partages d'opinions et de recommandations de communautés de cyberconsommateurs : notre acheteur pourra donc consulter pléthore de retours d'expérience sous forme de courts textes critiques sur toute une gamme d'appareils photos et ce selon différents critères d'évaluation avant de se décider à faire son achat. Nous disposons donc d'une « base documentaire » constituée des critiques d'autres cyberconsommateurs ou bien d'articles d'évaluation rédigés par des experts appartenant soit à des sites marchands soit à des organismes de consommateurs. Ces avis n'ont pas de forme « normalisée », ce sont de simples textes exprimant l'opinion d'un individu sur un article

selon une perspective d'analyse donnée—un critère d'évaluation. Ces appréciations ou critiques en langue naturelle sont généralement constituées de quelques phrases lorsqu'il s'agit d'un retour d'expérience, de quelques paragraphes lorsqu'il s'agit d'une critique de spécialiste. Un appareil photo est jugé suivant un ensemble de critères comme par exemple : qualité de la prise de vue, facilité de prise en main, mise au point, etc.

Le diagramme qui suit donne le schéma général du traitement nécessaire à l'indexation de la base documentaire par les CAs en insistant sur les deux activités principales :

- L'extraction des CAs à partir des textes relativement à un référentiel de critères d'évaluation ;
- l'évaluation d'intention de ces CAs, c'est-à-dire la détermination des scores correspondant aux jugements de valeur qu'elles expriment.

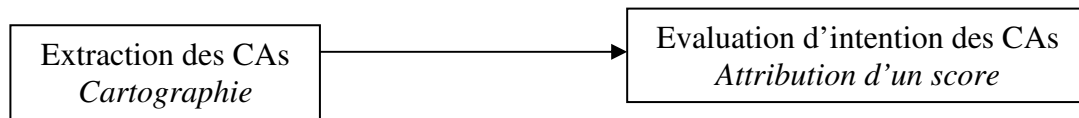


Figure V-1: Processus général

Le diagramme qui suit donne un schéma plus détaillé du traitement relatif à ces deux activités selon le formalisme UML (Figure V-2). Ce chapitre a pour objet de décrire ce diagramme.

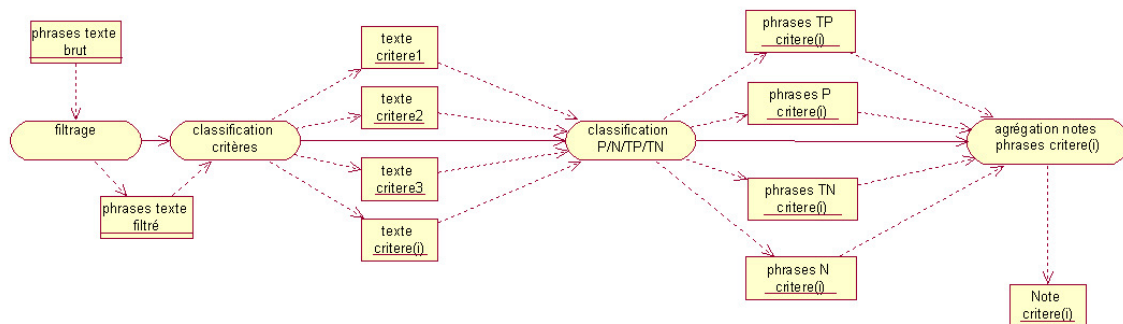


Figure V-2 : Détail du processus général

Les *symboles rectangulaires* représentent les objets du système d'information et les *rectangles à coins arrondis* représentent les activités (ou fonctions, ou tâches).

La tâche d'extraction mentionnée dans la Figure V-1 est ici décomposée en deux activités : « filtrage » et « classification critère ». Le « filtrage » sépare, dans un texte, les passages narratifs des jugements (intentions, appréciations, etc.). Il repose sur un premier classifieur. La « classification critère » associe ensuite un critère à chaque fragment de texte issu du filtrage. Ces activités sont présentées en détail aux paragraphes 4.1 et 4.2.

Les deux activités suivantes « classification P/N/TP/TN » et « agrégation notes » correspondent à la décomposition de « Evaluation d'intention ». La première établit la correspondance entre le jugement de valeur d'une CA et un score (Positif (P), Négatif (N), Très Positif (TP) et Très Négatif (TN)). La seconde agrège les scores de toutes les CAs portant sur un même couple (critère α , alternative φ) pour donner l'évaluation globale de φ sous l'angle de α . Cette activité est présentée en détail aux paragraphes 4.3 et 4.4.

Avant cela, nous décrivons l'élaboration de l'espace des critères d'évaluation et de la base de connaissance nécessaire aux classifieurs. Puis, nous explicitons le processus de vectorisation

et de diminution de l'espace vectoriel qui est utilisé dans la construction de tous les classifieurs. Et enfin, nous présentons la génération des classifieurs.

2. Corpus documentaire et base d'apprentissage

La première étape de l'apprentissage nécessaire à la réalisation de notre système d'aide à la décision est la détermination des critères de jugement auxquels se réfèrent les connaissances actionnables du système. Ces critères sont choisis par les utilisateurs du système d'aide à la décision en amont de tout apprentissage. Les critères s'expriment sous forme de courts aphorismes comme par exemple : qualité de la prise de vue (pour un appareil photo).

Le choix des critères est suivi par l'étape de construction de la base d'apprentissage pour laquelle, à chaque fragment de texte de la base, doit être associé l'un des critères. Compte tenu du travail d'indexation manuelle sur la base d'apprentissage, le nombre de critères doit rester limité. De plus, la performance des classifieurs automatiques qui seront édifiés sur cette base diminue avec le nombre de critères.

Pour initier l'apprentissage, il faut donc analyser chaque texte du corpus documentaire, le fragmenter si nécessaire lorsque celui-ci est composé d'éléments faisant référence à plusieurs critères et associer chacun de ces fragments de texte à un critère unique. La construction de la base est la tâche préliminaire à l'apprentissage de tous les classifieurs nécessaires à l'extraction des CAs.

L'origine de la base de documents dépend du type d'application. Les textes sont recueillis :

- soit sur Internet selon une procédure manuelle ou à l'aide d'un moteur d'extraction. Il s'agit d'extraire automatiquement à partir de pages Internet des retours d'expériences, des avis de consommateurs, des critiques d'art, etc. sur des sites identifiés a priori par l'utilisateur ;
- soit sur l'intranet d'une organisation ou d'une entreprise qui se prête à des pratiques de travail collaboratif et utilise son intranet pour partager connaissances et opinions émis par les membres ou les collaborateurs sur le projet, la réforme, etc.

Si le recueil des documents textuels constituant la base de connaissance est une tâche relativement aisée, la caractérisation des extraits textuels associés à chaque critère l'est beaucoup moins. Cette tâche fastidieuse demande une attention particulière dont la fiabilité d'apprentissage des classifieurs dépend complètement.

Comme nous le voyons sur le diagramme de classe UML qui suit, la base de connaissance est constituée de jugements de valeurs en langage naturel, les critiques. Chaque critique contient une appréciation sur l'une des alternatives. Si nécessaire, chaque critique est ensuite décomposée en fragments de texte afin qu'un fragment de texte se réfère à un et un seul critère.

Pour la partie de la base de connaissance réservée à l'apprentissage, cette opération est manuelle. L'apprentissage supervisé s'opère ensuite sur cette partie indexée de la base.

L'indexation manuelle de la base nécessite donc de distinguer dans une critique rentrée dans la base à une date donnée (Figure V-3) :

- Une partie dite narrative, dans laquelle l'auteur recadre dans son contexte et décrit brièvement l'alternative sur laquelle il va donner son opinion ;
- Une partie dite évaluative qui constitue la prise de position de l'auteur relativement à l'alternative ;

- La partie évaluative est ensuite décomposée en fragments de texte, chacun d'entre eux correspondant à un critère d'appréciation.

C'est sur ces couples (fragment de texte, critère) que sont élaborées les CAs. En effet, à chaque couple est attribué un score qui doit être cohérent avec le jugement de valeur de l'auteur. On fait l'hypothèse que ce score est exprimé dans une échelle de ratio ou d'intervalle lors de cette indexation manuelle. En pratique, chaque phrase des fragments de texte identifiés est associée à un score et le score du fragment est obtenu par agrégation des notes de chacune des phrases le constituant.

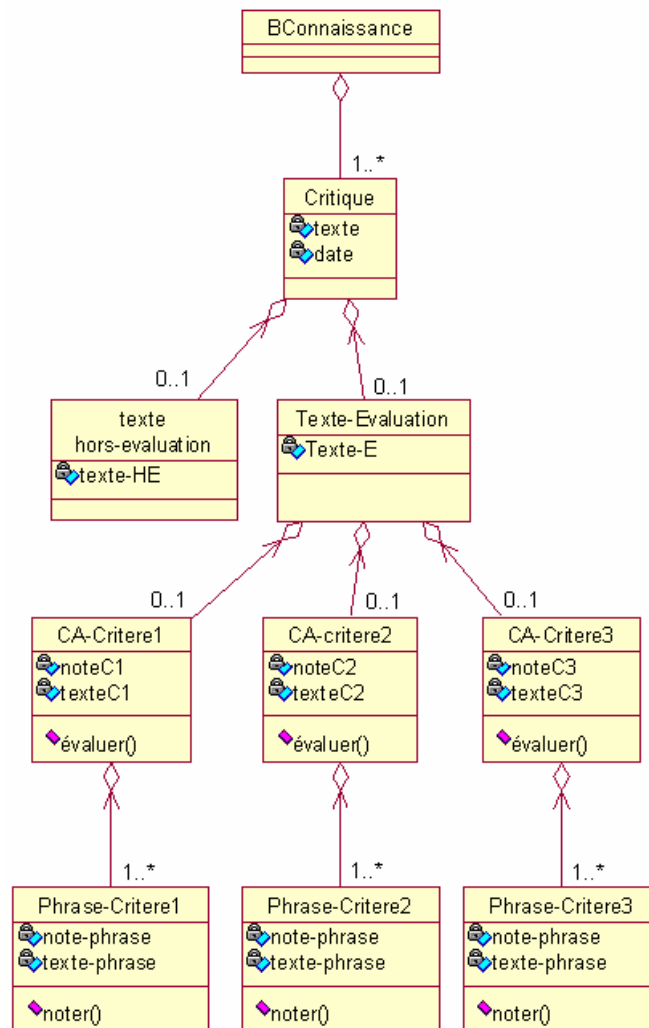


Figure V-3 : Structure de la base de connaissance pour trois critères

3. Représentation vectorielle du corpus de documents

Le diagramme d'activité de la Figure V-4 montre les différentes étapes qui mènent de la base d'apprentissage à la définition de l'espace de représentation utilisé par un classifieur. Comme nous l'avons vu précédemment, nous devons avoir recours à trois classifieurs : 1) Filtrage pour la reconnaissance des parties narratives/évaluatives des textes ; 2) Affectation d'un fragment de texte à un critère et 3) Attribution d'un score à un fragment de texte évaluatif. Chacun de ces classifieurs repose sur une représentation vectorielle spécifique des textes de sa base d'apprentissage : l'index du classifieur (paragraphe 3.3).

La représentation vectorielle d'un corpus de documents (voir paragraphe 1 du chapitre IV) nécessite en premier lieu un prétraitement linguistique pour éliminer tous les mots ne participant pas activement au sens du document (la stop-list). Sur les textes « épurés » est ensuite construit un premier index des textes de la base d'apprentissage. Chaque texte est alors représenté par ses coordonnées (occurrences de mots) dans cet index. La dimension de l'espace de représentation utile au classifieur est enfin réduite en introduisant les classes de celui-ci pour filtrer l'index initial (paragraphe 3.3). C'est ce que nous exposons ci-dessous.

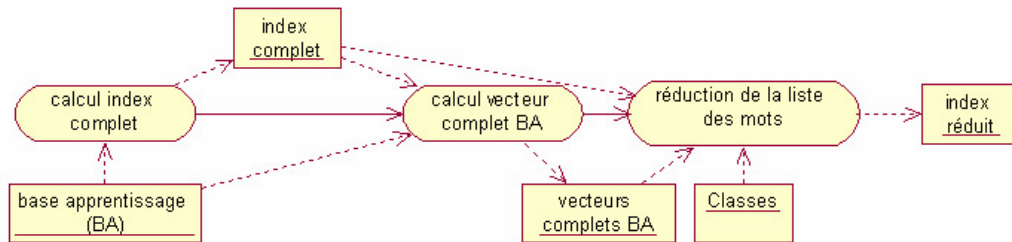


Figure V-4 : Processus de calcul des index

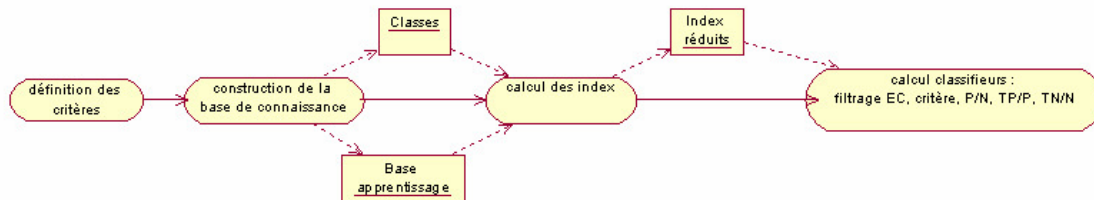


Figure V-5 : Traitements liés à l'apprentissage

3.1. Index et représentation vectorielle

Le traitement linguistique consiste à effectuer une lemmatisation de chaque texte comme indiqué au chapitre précédent. Nous utilisons pour cela l'outil « Synapse » [Synapse, 2001]. Une fois la lemmatisation de tous les textes réalisée, nous établissons la liste de tous les mots présents dans le corpus de textes de la base d'apprentissage. Un filtre linguistique est appliqué afin d'éliminer quelques types de mots considérés inutiles pour la catégorisation : les seuls éléments grammaticaux que nous éliminons sont les articles définis et indéfinis, la ponctuation faible et les prépositions. Nous faisons l'hypothèse que les articles et la ponctuation faible ont peu d'influence sur le sens des textes ou plus exactement sur leur impact dans le processus de catégorisation. A la sortie de cette étape, nous faisons l'hypothèse que nous avons conservé l'ensemble des mots et expressions utiles aux classifieurs. Nous pensons en particulier aux adverbes et autres types grammaticaux de nature similaire susceptibles d'apporter une contribution plus ou moins forte, de nuancer un jugement de valeur.

Comme nous l'avons fait remarquer dans le chapitre précédent, les tâches de catégorisation que nous cherchons à effectuer consistent à différencier des catégories « relativement » proches. Différencier deux critères d'évaluation dans un même domaine n'est pas, catégoriser des news selon leur thème, sport, politique, cinéma, etc. La vectorisation du corpus de documents ne saurait se suffire des 873 concepts du thésaurus Larousse comme proposé dans la méthode des vecteurs conceptuels (paragraphe 1.5.5 du chapitre IV). La granularité de l'espace de représentation doit être plus fine. Toutes les subtilités de la langue contenues dans les textes à analyser sont a priori nécessaires pour garantir une bonne performance de la

catégorisation. Nous conservons donc le maximum de mots pour la suite des opérations et ne procédons à aucune autre simplification grammaticale ou linguistique, seule une réduction de l'espace de représentation liée à la spécificité des classifieurs sera autorisée.

A la sortie de cette étape, la phrase « la fleur est belle », par exemple, est transformée en « fleur être beau » qui est la forme lemmatisée et filtrée de la phrase originale. Ce traitement linguistique donne accès à une représentation formelle.

Après avoir traité tous les textes par lemmatisation, l'étape suivante consiste à énumérer tous les mots lemmatisés ou lemmes qui constitueront les dimensions de l'espace vectoriel de représentation des textes. Cette énumération se fait sur le corpus d'apprentissage. Nous supposons que ce corpus est représentatif du domaine à traiter et que l'énumération des mots lemmatisés constitue le vocabulaire V du domaine. Nous obtenons une liste de mots m_i , $i..|V|$ que nous désignons par « index complet ».

L'index complet établi, il est maintenant possible de construire une représentation vectorielle du corpus de textes de la base d'apprentissage. Dans ce modèle, chaque lemme m_i constitue une dimension de l'espace vectoriel de représentation des textes. Ainsi un texte D est représenté par un vecteur $D = (ft_1, ..., ft_i, ..., ft_{|V|})$ où ft_i représente le nombre d'occurrences du lemme m_i dans le document D , $i..|V|$.

3.2. Constitution de l'ensemble d'apprentissage de chaque classifieur

A ce stade de la procédure, tout document de la base d'apprentissage (BA) est représenté par un vecteur dont les dimensions correspondent à l'index complet et les composantes sont les fréquences d'occurrence des termes de l'index dans le document. Les caractéristiques des classifieurs ne sont jusqu'ici pas intervenues dans les calculs.

Les textes sont alors vus comme des ensembles de phrases. Chacune des phrases d'un texte est étiquetée en vue de la construction des classifieurs. Ainsi, chaque phrase sera étiquetée de la façon suivante :

- Une première étiquette désigne la phrase comme étant un fragment de texte élémentaire soit narratif (elle appartient à l'ensemble des phrases Hors Contexte Evaluation, noté HCE) soit évaluatif donc se rapportant nécessairement à l'un des critères (Cr) d'évaluation (elle appartient à l'Ensemble des phrases utiles à l'Evaluation, noté Cr) ;
- Chaque phrase est ensuite étiquetée par le critère Cr_i auquel elle se rapporte dans l'ensemble Cr ;
- La troisième étiquette est le score associé à une phrase se rapportant à un critère Cr_i . Nous nous limitons à une échelle de valeur à quatre niveaux. Cette étiquette peut prendre les valeurs $\{TN, N, P, TP\}$.

Sur la base d'apprentissage, nous constituons donc conformément aux diagrammes des Figure V-3 et Figure V-5, les ensembles de phrases ou plutôt de vecteurs associés à ces phrases qui supporteront la construction de tous les classifieurs :

- Un ensemble de vecteurs hors contexte d'évaluation (HCE) et de phrases utiles à l'évaluation parce que se rapportant aux critères (Cr) ;
- L'ensemble des vecteurs de Cr est divisé en autant de sous-ensembles qu'il y a de critères d'évaluation ;

- Un ensemble de vecteurs positifs **P** (donc P ou TP) et un ensemble de vecteurs négatifs **N** (donc N ou TN) ;
- Un ensemble de vecteurs positifs (P) (resp. négatifs N) et un ensemble de vecteurs TP (resp. très négatifs TN).

Notons que l'introduction des ensembles **P** et **N**, puis P/TP et N/TN, permet de procéder à une classification hiérarchique : d'abord, on sépare les jugements de valeur favorables des jugements de valeur défavorables, puis pour chacun de ces deux sous-ensembles, les jugements de valeur sont gradués, quantifiés pour affiner l'échelle d'évaluation. Nous justifierons cette manière de procéder en temps voulu.

Nous calculons alors les vecteurs de chaque ensemble comme précisé au paragraphe précédent en rajoutant les coordonnées « classe ».

Chaque vecteur (phrase) peut donc être noté de la façon suivante :

$V_{D_i P_j} = (ft_{D_i P_j 1}, \dots, ft_{D_i P_j k}, \dots, ft_{D_i P_j |V|}, C_{C, D_i P_j}, C_{E, D_i P_j})$ est le vecteur de la phrase j du document i , k la composante de l'index, $ft_{D_i P_j k}$ le nombre d'occurrences du lemme k dans la phrase j du document i . E correspond à une classe parmi les quatre valeurs possibles de l'échelle d'évaluation : TP, P, N, TN . C correspond à une classe critère : critère 1, ..., critère p , HCE .

On note :

- $E_{HC} = \{V_{D_i P_j} / C_{C, D_i P_j} = HCE\}$: l'ensemble tels que $C=HCE$;
- $E_{Cr} = \{V_{D_i P_j} / C_{C, D_i P_j} = C_k, \forall k \in \{1, 2, 3\}\}$: l'ensemble des vecteurs évaluatifs ;
- $E_{TP} = \{V_{D_i P_j} / C_{E, D_i P_j} = TP\}$: l'ensemble des vecteurs tels que $E=TP$,
- $E_P = \{V_{D_i P_j} / C_{E, D_i P_j} = P\}$: l'ensemble des vecteurs tels que $E=P$,
- $E_N = \{V_{D_i P_j} / C_{E, D_i P_j} = N\}$: l'ensemble des vecteurs tels que $E=N$,
- $E_{TN} = \{V_{D_i P_j} / C_{E, D_i P_j} = TN\}$: l'ensemble des vecteurs tels que $E=TN$.

Les ensembles d'apprentissage nécessaires à la détermination des classifieurs pour le processus automatique d'évaluation multicritère s'écrivent alors :

- $E_{Cr-HC} = E_{HC} \cup E_{Cr}$ pour le classifieur qui classe les phrases comme étant narratives ou évaluatives ;
- E_{Cr} pour le classifieur qui associe un critère à une phrase ;
- $E_{P-N} = E_{TP} \cup E_P \cup E_N \cup E_{TN}$ pour le classifieur qui sépare évaluations positives et négatives ;
- $E_{P-TP} = E_P \cup E_{TP}$ pour le classifieur qui quantifie les appréciations positives ;
- $E_{N-TN} = E_N \cup E_{TN}$ pour le classifieur qui quantifie les appréciations négatives.

3.3. Réduction de l'index dédié à un classifieur

Considérons un ensemble de vecteurs choisis pour effectuer le calcul d'un classifieur par apprentissage. L'espace vectoriel défini sur la base d'apprentissage globale et dans lequel sont définis ces vecteurs comporte un nombre important de dimensions. Par suite, les vecteurs sélectionnés pour l'apprentissage d'un classifieur donné peuvent avoir de nombreuses composantes toujours nulles selon certaines de ces dimensions. On peut donc considérer que ces dimensions n'ont aucune incidence dans le processus de classification et peuvent même ajouter du bruit dans le calcul du classifieur entraînant des performances moindres de la classification.

Pour pallier cet inconvénient, nous avons choisi d'effectuer une réduction de l'index afin d'améliorer les performances des classifieurs. Une méthode très connue présentée par Cover mesure l'information mutuelle associée à chaque dimension de l'espace vectoriel [Cover *et al.*, 1991].

Plusieurs méthodes comparables à celles de Cover ont été proposées pour effectuer une sélection des mots représentatifs du domaine. [Wang *et al.*, 2003] compare quatre méthodes de sélection du vocabulaire. Ces quatre méthodes utilisent respectivement la procédure de sélection IDF (Inverse Document Frequency), la mesure de gain d'information (IG), la mesure de l'information mutuelle et la méthode du χ^2 (chi 2).

La très répandue méthode IDF est basée sur la fréquence en documents [Salton, 1981]. La fréquence en documents d'un mot ou terme est le nombre de documents dans lesquels ce terme apparaît. La valeur IDF est $\log(N/N_i)$ où N est le nombre total de documents de la base de connaissance et N_i est le nombre de documents de la base de connaissance dans lesquels le terme i apparaît. Des seuils sur la valeur IDF permettent de sélectionner ou rejeter le mot analysé de l'index [Salton *et al.*, 1988]. La mesure du gain d'information est une mesure fondée sur l'entropie [Han *et al.*, 2001]. La méthode du χ^2 (chi 2) repose sur des principes voisins de ceux de l'information mutuelle [Yang *et al.*, 1997]. Wang conclut que la mesure de l'information mutuelle de Cover donne les meilleurs résultats [Wang *et al.*, 2003]. C'est elle que nous utilisons pour la réduction des index dédiés à un classifieur. En voici les principes.

Considérons un ensemble de documents. Soit C la variable aléatoire associée à l'ensemble des classes et M_t la variable aléatoire représentant l'absence ou la présence du mot m_t dans un document, où M_t prend les valeurs $f_t \in \{0,1\}$: $f_t = 0$ indique l'absence du mot m_t et $f_t = 1$ indique la présence du mot m_t . Introduisons les grandeurs suivantes :

- $P(c)$: le nombre documents de la classe $c \in C$ apparaissant divisé par le nombre total de documents ;
- $P(f_t)$: le nombre de documents contenant une ou plusieurs occurrences du mot m_t divisé par le nombre total de documents ;
- $P(c, f_t)$: le nombre de documents de la classe $c \in C$ et contenant également le mot m_t divisé par le nombre total de documents.

Nous utilisons l'entropie $H(C)$ de la variable classe C : $-\sum_{c \in C} P(c) \log(P(c))$ et l'entropie de cette variable conditionnée par la présence ou l'absence du mot m_t , $H(C|M_t)$:

$$H(C, M_t) = - \sum_{f_t \in \{0,1\}} P(f_t) \sum_{c \in C} P(c|f_t) \log(P(c|f_t)) \quad (V-1)$$

L'information mutuelle moyenne est la différence entre l'entropie de la variable aléatoire C et son entropie conditionnelle relativement au mot m_t :

$$I(C, M_t) = H(C) - H(C | M_t) \quad (V-2)$$

On a donc :

$$I(C, M_t) = - \sum_{c \in C} P(c) \log(P(c)) + \sum_{f_t \in \{0,1\}} P(f_t) \sum_{c \in C} P(c | f_t) \log(P(c | f_t)),$$

soit :
$$I(C, M_t) = \sum_{c \in C} \sum_{f_t \in \{0,1\}} P(c | f_t) \log\left(\frac{P(c, f_t)}{P(c)P(c | f_t)}\right) \quad (V-3)$$

Nous sélectionnons les mots dont l'information mutuelle est supérieure à un seuil donné (s_I). Dans le chapitre suivant, nous verrons dans les différentes expérimentations que ce calcul diminue grandement la taille de l'index d'un classifieur qui peut passer de 12000 à moins de 600 mots.

Cette technique permet donc, pour un classifieur donné, de calculer un index réduit dédié à ce classifieur. Nous effectuons ce calcul pour les ensembles d'apprentissage définis à la fin du paragraphe 3.2 : les index réduits de chaque classifieur définissent un nouvel espace vectoriel dédié au classifieur. Nous nommons ces index I_{Cr-HC} , I_{Cr} , I_{P-N} , I_{P-TP} , I_{N-TN} . Il est alors nécessaire de calculer pour chaque sous-ensemble d'apprentissage les vecteurs occurrences avec ces nouveaux index. Nous aurons donc des vecteurs dédiés à chaque classifieur, l'apprentissage s'effectuera toujours sur ces index dédiés.

3.4. Ajout des synonymes

Ce traitement n'est pas spécifique au calcul de l'espace vectoriel dédié à chaque classification. Il vient en complément. Il consiste à introduire la notion de synonyme dans les vecteurs calculés pour chaque index réduit de classifieur.

Prenons un exemple : considérons le thème de la photographie, l'index réduit contient par exemple le terme « cliché ». Mais, imaginons que l'on ne retrouve pas son synonyme « épreuve » dans cet index parce qu'il a été éliminé « par le calcul » à l'issue de la sélection par information mutuelle. Chaque fois que l'on rencontre une phrase où ce mot « épreuve » est présent, il n'est pas associé au thème de la « photographie »... Nous estimons que ce phénomène crée une perte de précision dans la représentation vectorielle des phrases, en partie liée au faible nombre de lemmes que l'on peut espérer trouver dans une phrase. L'introduction des synonymes nous permet de compléter et de donner plus de précision à la représentation vectorielle. Cette étape consiste donc à construire une table de correspondance entre les lemmes de l'index réduit du classifieur et leurs synonymes respectifs. Aujourd'hui, cette étape est manuelle dans notre processus. Nous discuterons plus loin de l'automatisation possible de cette activité. Le principe ne consiste donc pas à modifier l'index du classifieur mais à associer à chaque lemme de cet index ses synonymes dans le contexte analysé : les synonymes sont donc comptabilisés dans le vecteur des occurrences exprimé dans l'index réduit sans que la dimension de l'espace vectorielle soit modifiée.

Soit i la coordonnée du vecteur de phrase associée au mot clé m_i , définissons $\{s_{ij}\}$ les synonymes de m_i , j allant de 1 à l_i avec l_i étant le nombre de synonymes du mot m_i .

m1	S11	S12	S13	
m2	S21	S22	S23	S24
m3	S31	S32		
m4	S41	S42	S43	S44
m5	S51	S52	S53	S54
m6	S61			
.....				

Figure V-6 : Exemple de la structure de table de correspondance entre lemmes de l'index et les synonymes.

Les vecteurs d'occurrences des phrases se calculent donc ainsi :

$$V_{D_i P_j} = \left(ft_{D_i P_j 1} + \sum_{k'=1}^{l_1} ft_{D_i P_j s_{1k'}}, \dots, ft_{D_i P_j k} + \sum_{k'=1}^{l_k} ft_{D_i P_j s_{kk'}}, \dots, ft_{D_i P_j |V|} + \sum_{k'=1}^{l_{|V|}} ft_{D_i P_j s_{|V|k'}}, C_{C, D_i P_j}, C_{E, D_i P_j} \right) \quad (V-4)$$

Les vecteurs des ensembles d'apprentissage E_{Cr-HC} , E_{Cr} , E_{P-N} , E_{P-TP} , E_{N-TN} sont réévalués sur ce modèle.

Le calcul de chacun des classifieurs peut alors être effectué. Nous avons décrit plusieurs méthodes permettant ce calcul au chapitre précédent. Le principe de tout classifieur peut être formulé par l'identification d'une fonction f qui à tout vecteur associe une classe : $f : V_{D_i P_j} \rightarrow C_k$. La méthode que nous avons retenue est fondée à la base sur le modèle multinomial de Bayes. Les raisons de ce choix sont essentiellement dues aux performances de ce classifieur sur les expérimentations que nous avons conduites, en particulier pour l'application que nous exposons dans le chapitre suivant. Les mesures de F-score (voir chapitre précédent) ont toujours été largement en faveur du classifieur de Bayes Multinomial sur l'ensemble de nos expérimentations.

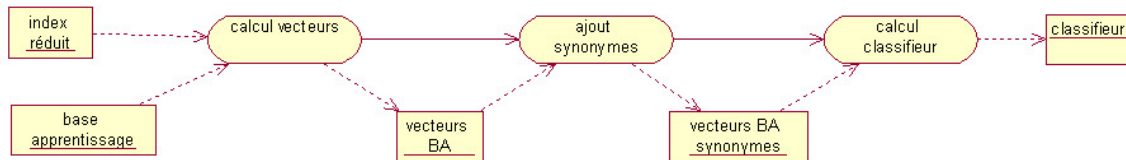


Figure V-7 : Diagramme d'activités générique du calcul d'un classifieur

4. La classification des textes

Pour un nouveau texte à analyser, il s'agit maintenant de 1) filtrer les parties narratives/évaluatives de ce texte, 2) affecter chaque fragment de texte à un critère et 3) Attribuer un score à chacun des fragments évaluatifs du texte. Chaque texte soumis à la classification est daté, cette date est attribuée à toutes les CAs qui sont générées depuis ce texte (date τ d'une CA).

4.1. le classifieur filtrage des textes évaluatifs et narratifs

Le diagramme d'activités de ce classifieur est proposé Figure V-8.

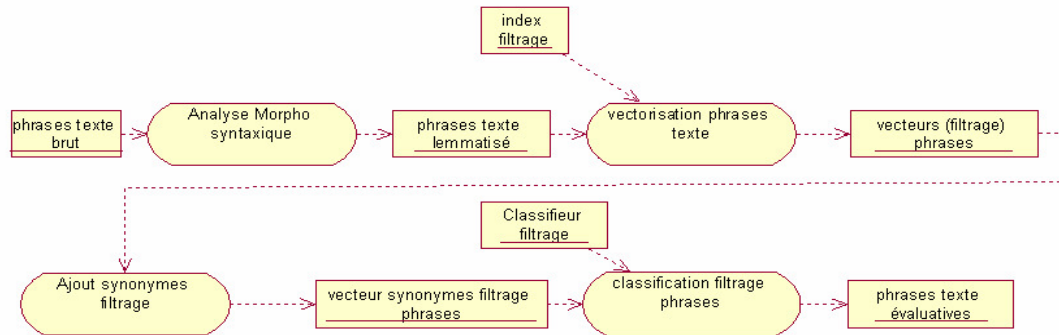


Figure V-8 : Diagramme d'activités du classifieur « filtrage »

Cette activité consiste à traiter le texte brut pour éliminer les phrases qui n'ont pas pour objectif de donner une évaluation ou une opinion sur l'une des alternatives concernées par l'évaluation.

Elle est composée de quatre activités :

- L'analyse morphosyntaxique qui permet d'obtenir un ensemble de phrases lemmatisées à partir du texte brut ;
- La vectorisation qui consiste à traduire sous forme de vecteur d'occurrences chacune des phrases lemmatisées dans l'index réduit du classifieur « filtrage » ;
- L'ajout de synonymes qui consiste à « enrichir » le vecteur d'occurrences précédent des synonymes ;
- La classification à proprement parler qui consiste à séparer les vecteurs d'occurrence de chaque phrase en deux ensembles : les phrases narratives et les phrases évaluatives. Le traitement des phrases narratives s'arrêtent ici : chaque phrase narrative est stockée dans la *synthèse* σ de la CA que l'on est en train de construire. Le traitement des phrases évaluatives se poursuit par leur association à l'un des critères d'évaluation.

4.2. Le classifieur « classification critère »

Cette activité consiste à traiter les phrases évaluatives du texte filtré provenant de l'activité décrite au paragraphe précédent et à les associer aux critères de jugement (Figure V-9 : Diagramme d'activités du classifieur « classification critère » (cas de 3 critères)).

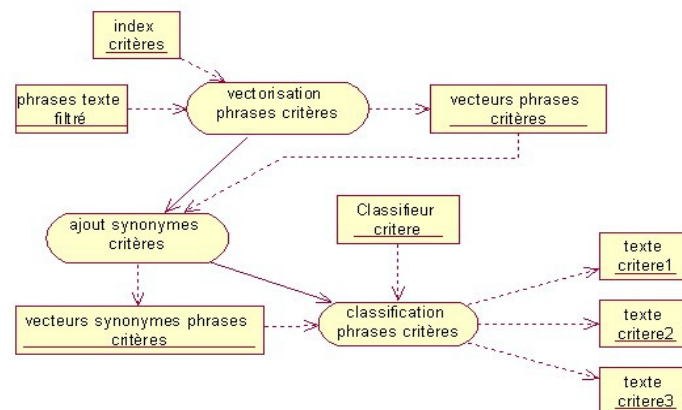


Figure V-9 : Diagramme d'activités du classifieur « classification critère » (cas de 3 critères)

Elle est composée de trois activités :

- La vectorisation des phrases analysées avec l'index réduit dédié à la classification en critères ;
- L'ajout de synonymes ;
- La classification à proprement parler qui consiste à associer les vecteurs d'occurrences aux critères du processus d'évaluation. Nous obtenons à la suite de ce traitement p ensembles de phrases, où chaque ensemble correspond à un critère. Les phrases issues d'un même texte attribuées à un même critère sont regroupées. Chacun de ces regroupements constitue le commentaire χ d'une CA ; l'alternative jugée et le critère identifié donnent les coordonnées (critère α , alternative φ) de la CA en question.

4.3. Attribution d'un score au commentaire d'une CA

Il s'agit de mettre en correspondance le commentaire χ d'une CA avec un score dans $\{TN, N, P, TP\}$. La procédure se fait en deux temps, on procède à une classification hiérarchique.

4.3.1. Evaluation hiérarchique

Si l'on demande à plusieurs individus de classer des phrases dans les quatre catégories : « Très Positif », « Positif », « Négatif », « Très Négatif », nous constatons que l'on peut dégager une majorité de phrases que tous les individus classeront de façon identique. Mais, il existera toujours un certain nombre de phrases pour lesquelles il y aura désaccord. Les catégories auxquelles l'on demande d'affecter les phrases sont des catégories sujettes à des nuances sémantiques parfois extrêmement subtiles. Si l'étiquetage manuel pose déjà problème, il va de soi que le classifieur aura d'autant plus de mal à obtenir un degré de performance parfaitement satisfaisant. C'est ce que nous avons constaté : si nous tentons de calculer un classifieur Naïve Bayes Multinomial pour ces quatre classes nous constatons que les performances sont très faibles, les F-scores prennent des valeurs inacceptables en terme de fiabilité. Pour cette raison nous avons choisi un traitement hiérarchique comme le montre la Figure V-10. Nous avons présenté nos travaux précédents sur ce sujet dans [Planté *et al.*, 2005b] qui exposait uniquement une évaluation selon deux modalités.

L'idée est la suivante : un traitement manuel des phrases évaluatives permettrait aisément de distinguer les jugements favorables des jugements défavorables. On construit alors un classifieur à deux classes $\mathbf{P} = \{P, TP\}$ et $\mathbf{N} = \{N, TN\}$. Les performances du classifieur s'améliorent très significativement pour atteindre un F-score tout à fait satisfaisant comme nous le verrons dans le chapitre d'application suivant. Ensuite pour chacune des deux classes $\mathbf{P}$ et $\mathbf{N}$, il faut quantifier le caractère favorable ou défavorable de la phrase. Là, le classifieur suivant repose sur un tout autre registre de mots pour la quantification des scores. Intuitivement, si dans la classification $\mathbf{P}$ versus $\mathbf{N}$, la distinction se faisait sur l'opposition de qualificatifs satiriques, négatifs (ingrat, décevant, etc.) et de qualificatifs élogieux, positifs (superbe, remarquable, etc.), pour le classifieur N/TN (resp. P/TP), ce sont les quantificateurs et les superlatifs qui vont autoriser les distinctions de quantification. Nous aurons donc conformément à la Figure V-10, trois classifieurs à deux catégories chacun (puisque l'on a décidé de se contenter de 4 niveaux d'évaluation dans notre étude).

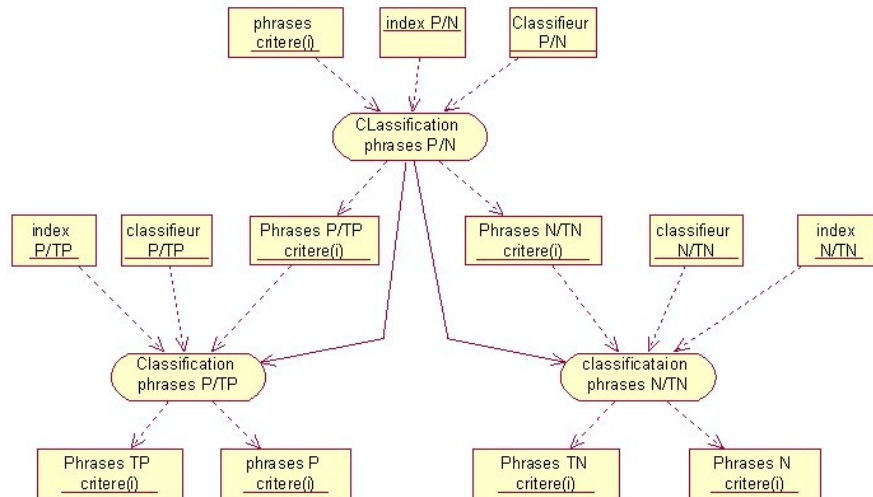


Figure V-10 : Evaluation hiérarchique des phrases

Nous allons donc décrire le contenu de chacune des trois activités du diagramme ci-dessus.

4.3.2. Le classifieur « classification phrases P/N »

Cette activité consiste à traiter chaque CA provenant de l'activité décrite au paragraphe précédent et à classer les phrases les constituant en deux groupes d'appréciations **P** ou **N** (Figure V-11).

Elle est composée de trois activités :

- La vectorisation d'un commentaire de CA associé à un critère avec l'index réduit dédié à la classification **P/N** ;
- L'ajout de synonymes ;
- La classification à proprement parler qui consiste à séparer les vecteurs d'occurrences de chaque phrase du commentaire de la CA en deux sous-ensembles correspondant aux deux évaluations **P** et **N**. Nous obtenons à la suite de ce traitement deux sous-ensembles de phrases, les phrases P et les Phrases N de la CA tout en conservant leur lien avec le commentaire dont elles sont originaires.

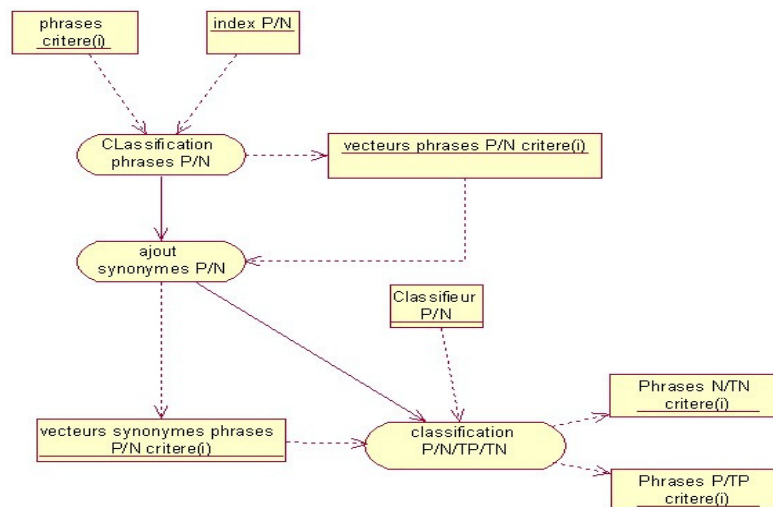


Figure V-11 : Diagramme d'activités du classifieur « classification **P/N** »

4.3.3. Le classifieur « classification N/TN » (resp. P/TP)

Cette activité consiste à traiter chacune des phrases classées N (resp. P) à l'étape précédente et à quantifier cet aspect défavorable (resp. favorable) du commentaire. Chaque phrase de N (resp. de P) doit être classée N ou TN (resp. P ou TP).

Elle est composée de trois activités :

- La vectorisation du vecteur analysé avec l'index réduit dédié à la classification TN/N (resp. TP/P) ;
- L'ajout de synonymes ;
- La classification à proprement parler qui consiste à classer les vecteurs d'occurrences de chaque phrase en deux sous-ensembles correspondant aux deux évaluations TN et N (TP et P). Nous obtenons à la suite de ce traitement deux ensembles de phrases, les phrases N et TN (respectivement les phrases P et TP).

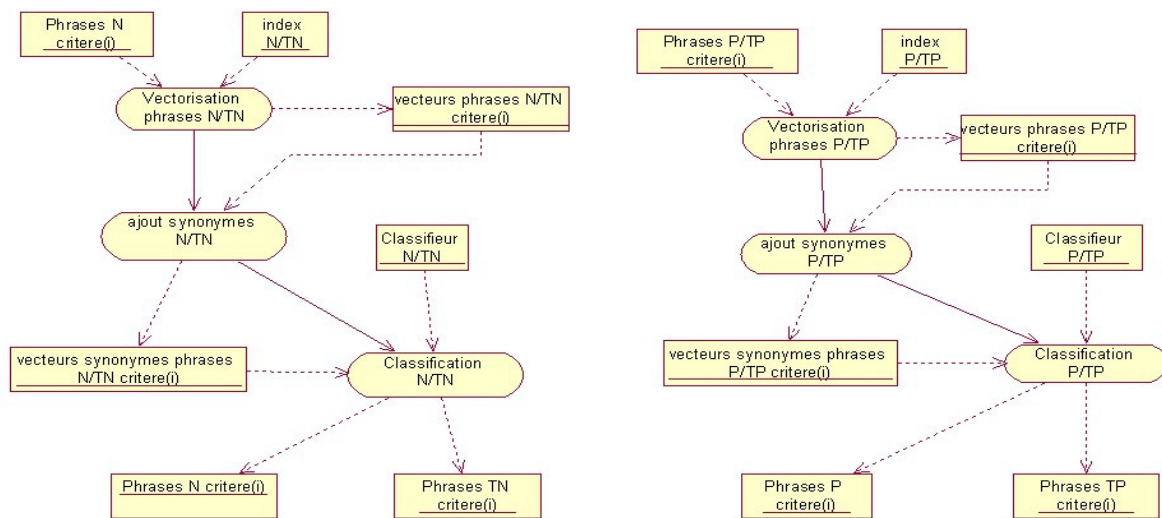


Figure V-12 : Diagramme d'activités du classifieur « classification N/TN »
(resp. « classification P/TP »)

4.4. Score d'une CA

A ce stade de la procédure, on dispose de regroupements de phrases relatives à un même critère et appartenant à un même texte à l'origine. Chacune de ces phrases a été affublée d'un score. Le score p de la CA dont le commentaire χ est constitué de ces phrases est la moyenne arithmétique des scores de chacune des phrases.

5. Quelques justifications

Nous allons décrire les raisons qui ont guidé le choix de la chaîne de traitement décrite précédemment.

L'introduction de la phase de tri par le classifieur « filtrage phrases narratives/évaluatives » correspond à une préoccupation récurrente dans le domaine de la fouille de textes : la nécessité d'extraire l'information ou la connaissance pertinente à partir d'un document. Dans notre contexte « connaissance pertinente » signifie la connaissance utile à l'analyse d'intentions dans les textes, c'est-à-dire la partie évaluative des textes. On élimine les éléments de connaissances non utilisables par l'évaluation car sans cela, ils perturberaient le processus. C'est le principe qui a guidé le défi DEFT 2005 auquel nous avons participé

[Planté *et al.*, 2005a]. Le challenge portait sur la distinction de phrases étrangères dans un discours. Plus exactement, les organisateurs du concours avaient introduit des phrases du président Mitterrand dans des discours du président Jacques Chirac. Le but du concours était d'identifier correctement les phrases de Mitterrand pour ensuite les éliminer [Planté *et al.*, 2005a].

La tâche de « classification critère » initie le processus de décision multicritère (voir à ce sujet [Planté *et al.*, 2005b]). Elle permet d'associer les parties évaluatives du texte liées aux différents critères de décision. L'efficacité de ce classifieur dépend beaucoup du filtrage précédent.

L'attribution d'un score à une phrase évaluative initie elle, l'évaluation. Une classification hiérarchique permet d'abord de distinguer les textes élogieux des textes critiques relativement à une alternative et un critère donnés. Dans un second temps, le recours aux superlatifs et quantificateurs sémantiques permet d'affiner l'évaluation. Chaque phrase constituant le commentaire d'une CA associée à un critère porte un score, le score de la CA est défini comme étant la moyenne arithmétique des scores des phrases. Il est bien évident qu'il est tout à fait possible de choisir un autre opérateur pour cette agrégation comme proposés dans le chapitre II, on pense en particulier aux majorités et unanimités restreintes qui ne tolèrent pas les compensations contrairement aux opérateurs moyennes.

La tâche d'ajout de synonymes enrichit la sémantique de la représentation vectorielle des textes sans augmenter la dimension de l'index d'un classifieur. Nous avons introduit les synonymes après la réduction de l'index d'un classifieur. Une autre logique aurait pu être l'introduction des synonymes avant le filtrage par calcul de l'information mutuelle... La granularité de nos analyses est la phrase, le faible nombre de lemmes potentiellement représenté dans l'index réduit d'un classifieur est une autre justification de l'introduction des synonymes dans nos analyses.

Le traitement automatique de la tâche d'ajout de synonymes peut être envisagé sans mettre en œuvre de processus trop lourds. Rappelons par exemple, la présentation des synonymes proposée par le thésaurus Larousse [Larousse, 1992] : pour un mot donné m_i , dans un premier temps nous sont proposés les différents sens du mot ou pour être plus précis les différentes définitions d_{ij} de m_i . Puis pour chaque définition d_{ij} des synonymes s_{ijk} sont listés.

Prenons un exemple. Considérons le mot m_1 : acteur. Et considérons le thème « cinéma ». Les différents sens du mot acteur que l'on trouve dans le thésaurus sont : *agent* et *comédien*. Si l'on regarde les définitions dans un dictionnaire, on peut trouver les définitions suivantes pour le mot *acteur* :

- d_{11} : *Personne qui prend une part déterminante dans une action ;*
- d_{12} : *Artiste qui joue dans une pièce de théâtre ou dans un film, comédien.*

A ces deux définitions correspondent deux listes de synonymes distinctes :

- s_{111} *Protagoniste*, s_{112} *membre*, s_{113} *organisateur*, s_{114} *leader*, etc.
- s_{121} *comédien*, s_{122} *interprète*, s_{123} *personnage*, s_{124} *vedette*, s_{125} *étoile*, s_{126} *star*, s_{127} *baladin*, etc.

Les synonymes que l'on retiendra seront ceux associés au domaine d'étude considéré, le cinéma, c'est-à-dire, les mots associés à la définition d_{12} : $s_{12,1}$ s_{122} , s_{123} , s_{124} , s_{125} , s_{126} , s_{127} .

Aujourd'hui, cette tâche relativement rapide (car même si elle n'est pas automatisée, elle est largement accompagnée par l'outil informatique, l'outil synapse proposant déjà cette liste de synonymes [Synapse, 2001]) est manuel. Il semble tout à fait envisageable de déterminer

automatiquement pour un mot donné la définition appropriée au contexte en effectuant une mesure de proximité entre les définitions et le thème d'étude (ici le « cinéma »). Il suffit alors de charger les synonymes associés à cette définition.

6. Conclusion

La chaîne de traitement que nous avons choisie permet de traiter les textes d'un corpus documentaire pour générer automatiquement les Connaissances Actionnables nécessaires au processus d'évaluation multicritère. La seule intervention humaine reste l'indexation manuelle de la base d'apprentissage (et aujourd'hui encore la détermination des synonymes pour les calculs des vecteurs d'occurrences).

A ce stade de traitement de l'information, nous sommes donc en mesure d'alimenter le processus d'évaluation multicritère avec des CAs. En effet, suite à ce traitement, la CA est repérée dans le référentiel du SGDC par le quintuplet $(\sigma, \chi, (\alpha, \varphi), \rho, \tau)$. La base de documents est alors indexée en CAs pour le processus de décision relativement au référentiel des n alternatives et p critères d'évaluation. La grille d'évaluation (critères X alternatives) à une date t donnée peut être renseignée :

- D'abord, en agrégeant case par case les scores des CAs (l'opérateur d'agrégation est à définir selon la sémantique de fusion que l'on désire) ;
- Puis, le score global de l'alternative est calculé en agrégeant les scores partiels des cases correspondantes (à ce niveau, l'opérateur d'agrégation modélise la stratégie de décision).

Ainsi, à chaque instant, les CAs cumulées et gérées par le SGDC permettent une évaluation complètement automatisée des alternatives concurrentes. Les fonctionnalités—justification des choix, contrôle de la dynamique décisionnelle par le risque—que nous avons explicitées au chapitre II, peuvent être mises en œuvre.

Dans le chapitre II, nous avons montré comment il était possible de déterminer les dimensions critiques du processus de décision pour lesquelles des CAs additionnelles devraient être rentrées dans le SGDC pour parvenir à une situation de décision stable. Le seul pré requis était qu'il y ait intervention humaine pour introduire ces CAs additionnelles au temps suivant. Le travail que nous avons exposé dans les trois derniers chapitres permet maintenant d'envisager l'automatisation de cette tâche. Nous avons donc construit l'actionneur (l'organe de commande) de la boucle de régulation d'un processus de décision multicritère comme exposé dans le chapitre II.

Nous allons illustrer et discuter ces propos sur une application dans le chapitre qui suit.

Chapitre VI.

Application : un outil d'aide à la décision pour un gestionnaire de vidéo-club

1. Gestionnaire d'un vidéo-club	122
1.1. Liste des films candidats	124
1.2. Critères d'évaluation	125
2. Collecte et évaluation des critiques	125
2.1. Constitution de la base de connaissance en vue de l'apprentissage	125
2.2. Calcul de l'index global	125
2.3. Écrans de saisie d'une nouvelle critique	125
3. Agrégation multicritère, justification et identification	132
3.1. Présentation générale de l'IHM	133
3.2. Utilisation des grilles multicritères	134
3.3. La grille Film/Date	135
3.4. La grille Film/Critère	139
3.5. La grille Critère/Date	140
3.6. Report de nouvelles critiques sur la grille	141
4. Evaluation du risque décisionnel	142
4.1. L'IHM de la fonction « risque »	142
4.2. Calcul du risque	143
4.3. Signal de contrôle	144
4.4. Calcul du nombre de critiques minimum pour modifier le classement	145
5. Argumentation	146
5.1. Argumentation en absolu	146
5.2. Argumentation en relatif	149
6. Scénario d'évolution dans le temps et stratégies de contrôle	152
6.1. Scénario CONFIRMER	153
6.2. Scénario EProuver	154
6.3. Scénario LIBRE ou BOUCLE OUVERTE	156
7. Conclusion	157

Dans ce chapitre, nous présentons une plate-forme logicielle d'aide à la décision qui met en œuvre l'ensemble de nos propositions sur un exemple à vocation pédagogique mais dont le caractère générique laisse envisager de nombreuses applications de notre démarche intégrée. L'accent sera bien sûr mis sur les techniques d'extraction automatique de connaissances actionnables (Chapitre III, IV et V), mais l'analyse multicritère, les notions de justification des choix et de risque décisionnel du chapitre II seront aussi largement illustrées. Nous verrons ainsi que notre plate-forme logicielle répond à l'ensemble des spécifications du chapitre I.

Nous avons choisi d'utiliser des outils existants et si possible standards et open source pour la plateforme. Ainsi, le prototype a été développé en langage java, la base de connaissance est une base de données gérée par le SGBD MySQL. Les outils d'apprentissage et de classification que nous avons utilisés reposent sur la plate forme standard Weka qui propose un ensemble de classifieurs et en particulier le classifieur Bayes Multinomial [Weka Project, 2002-2005]. Cette plateforme est elle aussi implémentée en java, ce qui a facilité l'intégration.

1. Gestionnaire d'un vidéo-club

Le problème que nous avons choisi de traiter pour illustrer notre démarche et nos outils est celui du gestionnaire d'un vidéo-club. L'idée est la suivante : le problème décisionnel auquel est confronté le gestionnaire d'un vidéo-club est de gérer au mieux ses nouvelles acquisitions afin de répondre aux demandes de sa clientèle tout en s'appuyant sur la critique disponible dans la presse spécialisée ou sur les sites web pour cinéphiles. Sa décision porte essentiellement sur le nombre d'exemplaires d'un film qu'il doit commander en fonction du succès de celui-ci lors de sa sortie sur les écrans et sur les films à privilégier dans ses commandes en fonction de son budget nécessairement limité.

Le gestionnaire s'en remet à l'évaluation de masse par les cinéphiles sur le web pour repérer les investissements les plus judicieux au regard des préférences identifiées de sa clientèle (supposée stable sur l'exercice !). Il ajuste sa stratégie d'achats en fonction de l'adéquation entre une bonne cote de la critique et les goûts de ses clients. Les connaissances manipulées seront donc rattachées au thème du cinéma, les connaissances actionnables se rapporteront à l'évaluation de films.

Ce domaine présente plusieurs caractéristiques intéressantes qui illustrent parfaitement notre approche.

Tout d'abord, la décision repose largement sur la capacité du gestionnaire à trouver et gérer les flux informationnels relatifs à la critique cinématographique. En effet, il existe aujourd'hui de très nombreux sites sur le net qui permettent de stocker et partager l'opinion des cinéphiles sur tous les films à l'affiche. Le gestionnaire du vidéo-club peut donc disposer de cette base et l'analyser pour projeter à six mois (il faut environ ce délai pour que le film passe de l'affiche aux bacs) quels seront les DVD les plus attendus à la location compte tenu du retour de la critique et la spécificité de sa clientèle. Il y a donc clairement un problème de traitement de la quantité d'information. Cette masse d'informations permet par ailleurs de construire une base d'apprentissage conséquente pour la calibration des classifieurs de l'application.

Ensuite, les critiques se rapportent presque toujours à un référentiel d'évaluation multicritère : les critiques se réfèrent à de multiples perspectives d'analyse sur les films, comme le scénario, la réalisation ou l'esthétique des images, etc. Sur certains sites, il peut être exigé que les cinéphiles attribuent des scores soit numériques soit par le biais d'« étoiles » ou de « barres » (de 0 à 5 avec une granularité d'une demi-étoile ou d'une demi-barre) en fonction d'une liste de critères préétablie sur le site d'accueil. Néanmoins, la majorité des sites laissent le

cinéphile libre d'exprimer son avis sans trame prédéfinie. Exiger qu'une critique soit formulée sous la forme d'une évaluation multicritère quantitative peut en effet être vue comme une fastidieuse tâche d'indexation qui pourrait nuire à la liberté d'expression et surtout au bon vouloir des critiques amateurs...qui font vivre et animent le site. Il faut donc se doter d'outils qui permettent de transcrire des critiques en langue naturelle en des scores relatifs à des critères d'évaluation afin de décharger complètement l'auteur d'une critique d'une formalisation trop astreignante. Enfin, toute la problématique autour de la stabilité d'une évaluation et sa justification dans un univers de discours multidimensionnel se met en place dans ce contexte de recommandation comme nous l'avons définie dans le chapitre II.

Le principe est donc de proposer au gérant un tableau de bord multicritère qui assiste le gestionnaire de vidéo-club dans son recueil d'évaluations. L'agrégation multicritère, qui modélise le comportement décisionnel de la clientèle du vidéo-club et permet au gérant d'ordonner les films les plus attendus par ses clients en fonction de leur score global, nécessite l'identification d'un opérateur de fusion des scores partiels ; cette identification devra être menée au préalable sur une sélection de films soumis à évaluation des adhérents du vidéoclub et que nous précisons plus loin.

Le problème du gestionnaire en termes de traitement de l'information s'illustre aisément. En effet, considérons ne serait-ce qu'une dizaine de films et pour chaque film environ 50 critiques disponibles. Chaque critique peut être le support de plusieurs jugements de valeur en référence à différents critères d'évaluation. Il s'agit donc de « découper » les critiques selon les p critères. Ce qui donne rapidement une quantité de fragments de textes, certes courts, mais extrêmement nombreux dont le gérant ne saurait tirer toute l'information évaluative sans une aide informatique. L'outil doit donc, en particulier, lui offrir une synthèse de l'expression de ces opinions et automatiser autant que ce peut l'évaluation multicritère.

La figure ci-après reprend le modèle général sur lequel nous avons conclu au chapitre II, instancié au problème du gérant de vidéo-club (Figure VI.1). Pour un classement à un instant donné obtenu par comparaison des scores agrégés des films en compétition, le système estime l'incertitude épistémique (mesurée par la stabilité du classement) au vu des critiques disponibles à cet instant et détermine les critères les plus susceptibles de pouvoir perturber ce classement si de nouvelles critiques venaient à être rentrées dans la base (calcul du risque décisionnel). La batterie de classifieurs que nous avons définie dans le chapitre précédent reçoit alors ce « signal de commande » et recherche automatiquement, sur un ensemble de bases présélectionnées, les CAs correspondant aux dimensions identifiées par le calcul du risque décisionnel.

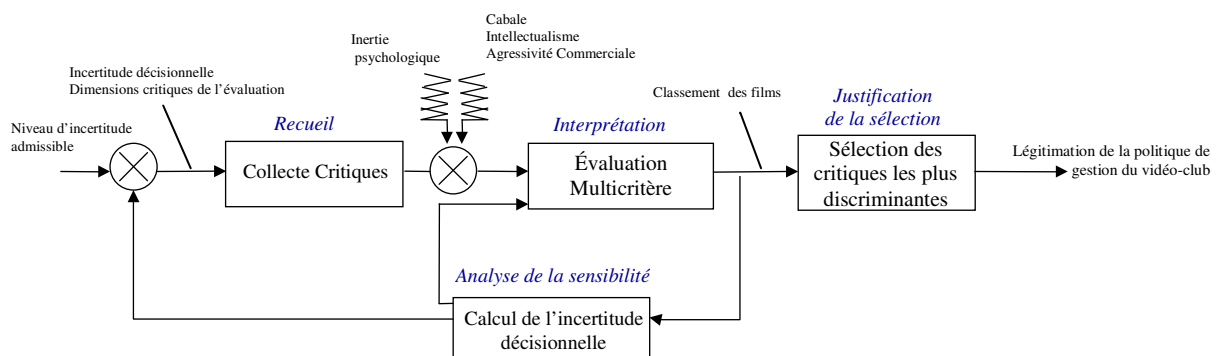


Figure VI.1 : Le processus de décision multicritère du gestionnaire de vidéo-club

Le processus de sélection de films s'apparente donc à un système dynamique qui peut être piloté par le risque décisionnel, mesuré comme nous l'avons expliqué au chapitre II par la stabilité du classement. Tant que le risque décisionnel reste au-delà d'un seuil admissible, l'incertitude épistémique est jugée trop importante pour que le gérant puisse arrêter sa sélection « sereinement ». La base de critiques dont il dispose ne contient pas les arguments nécessaires et suffisants pour que son arbitrage puisse être justifié aux yeux de ses clients. Il a besoin de nouvelles critiques : le calcul du risque décisionnel indique les dimensions du processus d'évaluation qui prédisposent le système à l'instabilité. Le système d'information recherche alors de nouvelles critiques qui sont soumises aux classifieurs de CAs pour produire une évaluation au temps suivant. La boucle de contrôle est active jusqu'à ce que le risque décisionnel devienne acceptable, le système a alors construit un argumentaire du nouveau classement : le gérant peut faire valoir les critiques qui ont eu le rôle le plus discriminant dans sa sélection.

Notre problématique diffère donc des *films recommenders* classiques du net [Akharraz *et al.*, 2003; Perny *et al.*, 2001]. Ces classifieurs sont généralement basés sur des techniques de content-mining ou de profil-mining à destination de l'individu. Ici, le gestionnaire a la responsabilité d'acquérir les films susceptibles de satisfaire au mieux la communauté des adhérents du vidéo-club [Balabanovic *et al.*, 1997; Basu *et al.*, 1998; Herlocker *et al.*, 2000; Hill *et al.*, 1995; Mukherjee *et al.*, 2001].

Les films sont évalués à travers un certain nombre de critères que nous allons décrire au paragraphe 1.2. Une note partielle est attribuée à chaque film candidat selon chaque critère sur la base des critiques collectées sur ce film et sous l'angle d'analyse de ce critère : elle est obtenue par agrégation des scores p des CAs calculés par les classifieurs. Au vu de ces évaluations partielles sur chaque critère, une note globale peut ensuite être attribuée au film avec un second niveau d'agrégation. L'opérateur d'agrégation synthétise alors la spécificité, les préférences et les priorités des adhérents. Les films sont ensuite classés en fonction de leur score global et les scores sont justifiés par des extraits de critiques (les commentaires χ des CAs). On retrouve bien dans ce processus les concepts de justification des chapitres I et II : les scores sont des informations certes synthétiques mais qui ne sauraient remplacer les arguments rhétoriques que sont les textes des critiques. Le gestionnaire doit préparer un argumentaire de sa sélection pour répondre à toute réclamation : il pourra, par exemple, éditer à cet effet une gazette annonçant les nouveautés assorties des élogieuses critiques que le SIAD aura « construites » et gérées pour lui.

Le problème du gestionnaire du vidéo-club va nous permettre d'illustrer la chaîne de l'information complète comme proposée sur la Figure VI.1 : traitement automatique des critiques collectées sur le net, scoring automatique des critiques, évaluation multicritère, calcul du risque décisionnel et justification de la logique de décision. Notre contribution technique ayant plus spécifiquement portée sur la construction automatique de CAs, une place privilégiée est réservée à cette activité.

1.1. Liste des films candidats

Pour illustrer nos propos, nous nous sommes limités à une sélection de 5 films :

- Le pacte des loups (film fantastique) ;
- Love actually (comédie romantique) ;
- Immortel (film de science fiction et d'animation) ;
- Meurs un autre jour (film d'action) ;
- Cube (thriller et film de science fiction).

1.2. Critères d'évaluation

Les critères utilisés pour évaluer les films sont déterminés au préalable par le gestionnaire du vidéo-club qui s'est livré à une étude des goûts de sa clientèle (cf paragraphe 3.3.2). Toujours hors application, une base d'apprentissage propre à chaque critère est élaborée. Cette base contient un ensemble de critiques qui ont été attribuées « manuellement » au critère. Dans notre application, les critères retenus sont les suivants :

- Acteur : qualité du jeu d'acteur, interprétation, etc.
- Scénario : qualité du scénario, structure de l'histoire, originalité du scénario, etc.
- Réalisation : mise en scène, effets spéciaux, cohérence entre les séquences, mise en image, traduction du scénario, etc.

2. Collecte et évaluation des critiques

Le calcul du risque décisionnel indique sur quels couples (critère α ; alternative φ), de nouvelles critiques doivent être intégrées dans la base dédiée à l'évaluation. Les classifieurs sont alors lancés pour rechercher, construire et noter les CAs. On peut donc considérer que le moteur de recherche qui rapatrie les textes depuis le net et les classifieurs de CAs constituent l'actionneur ou organe de commande de la boucle de contrôle de la Figure VI.1. Cet actionneur collecte les critiques, les décompose en fragments textuels relatifs aux critères et enfin leur attribue un score pour créer les CAs. L'actionneur repose sur la procédure de traitement de l'information du chapitre V.

2.1. Constitution de la base de connaissance en vue de l'apprentissage

Ce processus nécessite la constitution d'une base de connaissance permettant l'apprentissage et le calibrage des classifieurs nécessaires à l'évaluation multicritère. Nous avons collecté un ensemble de critiques sur Internet : 398 critiques se rapportant à 137 films, soit 4533 phrases. Le processus manuel d'indexation de la base d'apprentissage a été simplifié. Nous avons décomposé un sous-ensemble de textes de critiques en deux parties : « texte_synopsis » (fragments de texte non évaluatifs), « texte_évaluation ». Chaque « texte_évaluation » a ensuite été décomposé en trois extraits correspondant aux trois critères. Enfin, les notes P, N, TP, TN ont été attribuées à chaque extrait de texte et non à chaque phrase. Nous avons choisi cette procédure afin de simplifier la tâche d'indexation manuelle.

2.2. Calcul de l'index global

Dans le processus d'apprentissage de la méthode du chapitre V, figure également le calcul de l'index complet qui permet d'établir la représentation vectorielle des textes à traiter. Ce calcul sur la base d'apprentissage de films a donné lieu à un index total de 12131 mots lemmatisés. Tous les traitements que nous allons décrire utilisent cette base de connaissance et l'index complet.

2.3. Écrans de saisie d'une nouvelle critique

L'écran de la Figure VI.2 illustre la chaîne de traitement des critiques de films, préparatoire à l'évaluation multicritère. Nous allons illustrer le traitement d'une critique de film que nous souhaitons insérer dans la base de connaissance. Ce processus, décrit dans le chapitre précédent, comporte plusieurs étapes que nous allons reprendre sur notre application ci-après.

Saisir une nouvelle critique

texte original

L'auteur raconte une très belle histoire. Christophe Gans réussit avec ce film, un génial mélange des genres. Film historique à la base inspiré d'une légende ancestrale, mais filmé d'une manière totalement novatrice, mêlant à la fois du suspense lié à une intrigue bien présente, et des scènes de combats dignes des plus grands films d'action, le tout en costumes d'époque. Les acteurs sont irréprochables, on retrouve un Samuel Le Bihan obstiné jusqu'à découvrir la vérité, un énigmatique Mark Dacascos, une gracieuse Emilie Dequenne, une mystérieuse Monica Bellucci, et un Vincent Cassel aussi louche qu'à son habitude. C'est du grand spectacle ! Le film décrit de façon romancée la traque et la capture d'une bête féroce certainement un loup qui terrorisait la région du gévaudan en France.

partie critique

L'auteur raconte une très belle histoire. Christophe Gans réussit avec ce film, un génial mélange des genres. Film historique à la base inspiré d'une légende ancestrale, mais filmé d'une manière totalement novatrice, mêlant à la fois du suspense lié à une intrigue bien présente, et des scènes de combats dignes des plus grands films d'action, le tout en costumes d'époque. Les acteurs sont irréprochables, on retrouve un Samuel Le Bihan obstiné jusqu'à découvrir la vérité, un énigmatique Mark Dacascos, une gracieuse Emilie Dequenne, une mystérieuse Monica Bellucci, et un Vincent Cassel aussi louche qu'à son habitude. C'est du grand spectacle !

texte critère : Réalisation

Christophe Gans réussit avec ce film, un génial mélange des genres. Film historique à la base inspiré d'une légende ancestrale, mais filmé d'une manière totalement novatrice, mêlant à la fois du suspense lié à une intrigue bien présente, et des scènes de combats dignes des plus grands films d'action, le tout en costumes d'époque. C'est du grand spectacle !

texte critère : Scénario

L'auteur raconte une très belle histoire.

texte critère : Acteur

Les acteurs sont irréprochables, on retrouve un Samuel Le Bihan obstiné jusqu'à découvrir la vérité, un énigmatique Mark Dacascos, une gracieuse Emilie Dequenne, une mystérieuse Monica Bellucci, et un Vincent Cassel aussi louche qu'à son habitude.

lepectedesloups

30 05 2004

charger critique

Extraire partie critique

décomposer en critères

NOTER

Note R : 18.0

Note S : 18.0

Note A : 18.0

REPORTER NOTES SUR GRILLE

SORTIE

Figure VI.2 Ecran de saisie d'une nouvelle critique

2.3.1. Choix automatique du type de classifieur

Nous disposons d'une bibliothèque de classifieurs. Les classifieurs sont choisis automatiquement par notre système en phase d'apprentissage parmi ceux que nous avons évoqués au chapitre précédent (Bayes Multinomiale, Arbre de Décision, SVM). Nous avons également intégré dans cette liste des classifieurs dérivés du classifieur de Bayes Multinomiale (que nous n'avons pas mentionnés dans ce document) qui permettent de corriger certains biais de celui-ci. Il ne s'agit pas dans notre étude de démontrer la justification théorique mathématique d'un classifieur car le traitement automatique du langage naturel (TALN) constitue un domaine où l'usage de classifieurs bénéficie déjà d'un certain retour

d'expérience. Les classifieurs les plus utilisés en TALN sont disponibles sur la plate forme standard Weka. Nous avons donc procédé de manière pragmatique et systématique en testant tous les classifieurs disponibles sur la base d'apprentissage qui devient un benchmark. Les classifieurs sont jugés sur deux caractéristiques principales :

- la réussite de la classification mesurée par la mesure F-score ;
- le temps de calcul.

Comme discuté au chapitre IV, le F-score constitue le meilleur compromis de mesure dans notre problématique [Rijsbergen, 1979].

La rapidité des calculs est a priori un facteur de moindre importance. Néanmoins, si l'on imagine que les classifieurs doivent être recalculés régulièrement pour des raisons d'évolution de la définition du domaine, autrement dit de l'index et l'ajout de nouvelles critiques intégrées dans le processus d'apprentissage (ne serait-ce que l'adaptabilité du système aux nouveaux noms du cinéma, à l'évolution des critères), les temps de calcul doivent être inférieurs à quelques minutes. Quand le nombre de vecteurs est important et que l'index complet est de grande dimension, les temps de calculs peuvent dépasser pour certains types de classifieurs l'heure de calcul.

Ainsi, le calcul de F-score est effectué pour plusieurs types de classifieurs en phase d'apprentissage et le classifieur avec le meilleur F-score est implémenté dans le SIAD. Cependant, dans la grande majorité des cas, le classifieur Naïve Bayes Multinomial a donné les meilleurs résultats.

2.3.2. Première étape de filtrage de texte

- La première action est de charger un fichier texte relatif à un film donné et réputé contenir une critique par l'intermédiaire du bouton « charger critique ». Le texte apparaît alors dans la zone « texte original ». Prenons ici l'exemple d'une critique associée au film « Le pacte des Loups ». Le texte que nous voyons affiché sur la Figure VI.2 comprend un ensemble de phrases à filtrer conformément à la chaîne de traitement du chapitre précédent ;
- La deuxième action consiste à séparer le texte en deux sous-ensembles de phrases : les phrases évaluatives et les phrases hors contexte d'évaluation. Cette action est commandée par l'intermédiaire du bouton « Extraire partie critique ». Dans notre application d'évaluation de films, cette étape consiste donc à séparer :
 - les phrases ne portant aucun jugement de valeur, c'est-à-dire les phrases simplement destinées à décrire le film sujet de la critique, ce que l'on appelle classiquement dans le domaine du cinéma : le synopsis ;
 - les phrases contenant la partie évaluative de la critique.
- Le résultat du filtrage est un texte affiché dans la zone « partie critique » : il contient toutes les phrases classées évaluatives. Les phrases réputées constituer le synopsis ont été supprimées dans cette zone d'affichage. Dans notre exemple, le filtrage a supprimé la phrase : « *Le film décrit de façon romancée la traque et la capture d'une bête féroce certainement un loup qui terrorisait la région du Gévaudan en France* » qu'il a donc classée « synopsis ». Ce synopsis constituera la synthèse (σ) de toutes les CAs générées depuis cette critique.

Pour le calcul de ce classifieur, nous avons utilisé l'ensemble d'apprentissage E_{Cr-HC} dans lequel chaque critique de film est scindée manuellement en deux parties : la partie « synopsis » et la partie critique. Nous constituons ainsi les deux sous-ensembles :

- (i) E_{HC} contenant tous les extraits de textes associés aux synopsis des critiques ;
- (ii) E_{Cr} contenant tous les extraits de textes associés à la partie évaluative des critiques.

Chaque texte de E_{HC} comme de E_{Cr} est ensuite considéré comme un ensemble de phrases conformément à la démarche décrite dans le chapitre V. Nous effectuons en premier lieu une réduction de l'index complet. Ce calcul sélectionne un ensemble de 520 lemmes discriminants sélectionnés dans l'index complet de 12131. Le nombre de phrases de la base d'apprentissage est de 1978, réparties en 978 phrases « synopsis » et 1000 phrases « évaluatives ». Nous calculons le classifieur Naïve Bayes Multinomial pour cette base d'apprentissage. Nous l'évaluons par une procédure de validation croisée. Les résultats sont tout à fait satisfaisants et récapitulés dans le tableau suivant :

	Précision	Rappel	F-score
Classe critique	84,8%	82,3	83,5
Classe synopsis	83,2%	85,6	84,4

Au total le classifieur classe correctement 84% des phrases.

2.3.3. Catégorisation « critère »

A cette étape, les phrases évaluatives sont associées aux critères *scénario*, *réalisation* et *acteur*. Cette action est lancée via le bouton « Décomposer en critères ». Associé à un critère, un texte évaluatif est le commentaire χ d'une CA. Le résultat de cette décomposition est affiché dans les zones : « texte critère : scénario », « texte critère : réalisation » et « texte critère : acteur » de la Figure VI.2. Sur l'exemple, c'est le critère « réalisation » qui est le mieux renseigné après décomposition de la critique initiale.

Pour le calcul de ce classifieur, nous avons utilisé l'ensemble d'apprentissage E_{Cr} dans lequel l'utilisateur associe chaque phrase à un critère. Puis, nous utilisons la procédure du chapitre V pour calculer le classifieur. Nous construisons l'index réduit. Ce calcul donne un sous-ensemble de 130 lemmes discriminants sélectionnés dans l'index complet. Les 978 phrases de E_{Cr} sont réparties de la façon suivante : 278 phrases appartenant au critère acteur, 491 phrases appartenant au critère réalisation et 209 phrases appartenant au critère scénario. Le classifieur Naïve Bayes Multinomial s'avère être le meilleur classifieur sur la base d'apprentissage. On met en place une procédure de validation croisée pour évaluer ses performances récapitulées dans le tableau ci-contre :

	Précision	Rappel	F-score
Classe acteur	89,2%	62,6	73,6
Classe réalisation	70,8%	95,1	81,1
Classe scénario	89,4%	52,6	66,3

Au total, le classifieur classe correctement 77% des phrases. Soit :

	Acteur	Réalisation	Scénario
Acteur	174	99	5
Réalisation	16	467	8
scénario	5	94	110

On constate que les scores obtenus peuvent paraître plus mitigés que pour le classifieur de la tâche de filtrage. Ils sont dus principalement au fait que la classe « réalisation » a tendance à *recupérer* des individus des autres classes. Ceci s'explique par le fait que le concept de *réalisation* recouvre en partie les idées des deux autres critères.

Prenons un autre exemple de ce traitement pour le texte suivant :

« *Le film est très maladroit. Hélas pour ajouter à cela, l'histoire est très mal construite. L'interprétation est malheureusement décevante.* »

Le classifieur classe les phrases selon le tableau suivant :

Critère Réalisation	Le film est très maladroit.
Critère Scénario	Hélas pour ajouter à cela, l'histoire est très mal construite.
Critère Acteur	L'interprétation est malheureusement décevante.

A l'issue des classifications *filtrage* et *critère*, les textes évaluatifs ont été associés aux critères. Autrement dit, à ce stade du traitement les champs (*synthèse* (σ), *commentaire* (χ), *critère x film* (α, φ)) des CAs sont renseignés.

2.3.4. Attribution automatique d'un score

Cette étape consiste à attribuer automatiquement un score à chaque commentaire χ de CA. Chaque phrase est analysée et notée conformément au processus hiérarchique décrit au chapitre 5. Les scores des phrases de χ sont ensuite agrégés afin d'associer un score à χ . Ce score partiel ρ est l'évaluation du film φ au regard du critère α .

Cette tâche est lancée par l'intermédiaire du bouton « NOTER ». Le résultat est l'affichage de trois notes dans les zones : « Note R », « Note S », « Note A » correspondant aux trois critères (si le texte évaluatif a donné lieu à la création de trois CAs se rapportant aux trois critères bien sûr). Le classifieur a évalué chacune des phrases du texte évaluatif comme indiqué ci-dessous.

Phrase	Note
Critère Réalisation :	
Christophe Gans réussit avec ce film, un génial mélange des genres.	18
Film historique à la base inspiré d'une légende ancestrale, mais filmé d'une manière totalement novatrice, mêlant à la fois du suspense lié à une intrigue bien présente, et des scènes de combats dignes des plus grands films d'action, le tout en costumes d'époque.	18
C'est du grand spectacle !	18
Critère Scénario :	
L'auteur raconte une très belle histoire.	18
Critère Acteur :	
Les acteurs sont irréprochables, on retrouve un Samuel Le Bihan obstiné jusqu'à découvrir la vérité, un énigmatique Mark Dacascos, une gracieuse Emilie Dequenne, une mystérieuse Monica Bellucci, et un Vincent Cassel aussi louche qu'à son habitude.	18

Dans cet exemple toutes les phrases ont été associées à 18 correspondant à l'évaluation TP « très positive ». La note du critère *réalisation* au regard de cette critique est la moyenne des notes de chaque phrase du texte associé à ce critère. Les textes des deux autres critères ne comprennent qu'une seule phrase.

Sur l'exemple, la critique initiale a donc donné lieu à la création de trois CAs relatives aux trois critères, toutes trois extrêmement positives. Dans cette description, nous avons expliqué à des fins pédagogiques l'enchaînement des tâches via les boutons de commande, il va de soi que cette séquence peut faire l'objet d'un script lancé sur une base de critiques qui générera de façon complètement automatique toutes les CAs.

Revenons sur les performances des classifieurs qui attribuent les scores.

Classification P/N. Dans un premier temps, on cherche simplement à savoir si le commentaire χ de la CA est positif ou négatif. Nous utilisons la procédure générique décrite au chapitre V. L'index réduit ramène la dimension de l'espace vectoriel associé au classifieur P/N à 580 lemmes discriminants. Le nombre de phrases de l'ensemble d'apprentissage $E_{P-N} = E_{TP} \cup E_P \cup E_N \cup E_{TN}$ est de 4533, soit 2692 phrases appartenant à la catégorie **P**, 1841 à la catégorie **N**. Le classifieur Naïve Bayes Multinomial est retenu. Les performances du classifieur sont testées par validation croisée. Le résultat est très encourageant puisqu'au total le classifieur classe correctement 75% des phrases.

	Précision	Rappel	F-score
Classe positif	74,2%	88,2	80,6
Classe négatif	76,2%	55,2	64,1

Certaines performances mériteraient d'être améliorées. La qualité du classement dépend de la constitution de la base de connaissance. Actuellement, chaque commentaire χ de CA a été jugé globalement positif ou négatif (pour l'apprentissage, les phrases de χ héritent donc d'un même score), car ce jugement n'a pas été effectué au niveau de chaque phrase. En effet cela représente un travail d'indexation manuelle important qu'il n'est pas envisageable de transposer dans un environnement réel. Une façon d'améliorer les performances serait un raffinement de l'indexation de la base d'apprentissage.

Ce classifieur permet donc d'obtenir un classement de chaque phrase en positif ou négatif. Le résultat de cette étape est donc une note globale qualitative **P** ou **N** correspondant à la moyenne des notes de chaque phrase de χ , qui est affinée à l'étape suivante.

Classification N/TN. La deuxième étape du processus hiérarchique d'évaluation est de quantifier le score qualitatif précédent. Ici, on souhaite simplement distinguer deux niveaux de quantification : négatif N et très négatif TN.

L'index réduit ramène l'espace de représentation dédié au classifieur à 330 lemmes contre 12131 dans l'index complet. $E_{N-TN} = E_N \cup E_{TN}$ comprend 1841 phrases, dont 1207 TN et 634 N. Sur la base d'une validation croisée, nous obtenons des résultats très encourageants pour un Bayes multinomial :

	Précision	Rappel	fscore
Classe positif	80,4%	95,9	87,5
Classe négatif	87,8%	55,4	67,9

Soit au total le classifieur classe correctement 75% des phrases.

On constate dans le tableau suivant que certains scores obtenus sont mitigés (le rappel de la classe N). Les mauvais scores sont essentiellement dus aux critiques N classées TN. Cette

classification est plus difficile que les précédentes parce qu'elle fait appel à des nuances de langage plus subtiles. L'idée de base est qu'une critique est classée TN lorsqu'il y ait fait usage de superlatifs ou bien que les termes employés sont « au-delà de la norme » (nul, à oublier au plus vite, navrant, etc). Malheureusement, nombre de critiques simplement négatives font appel à ces mots, mais dans une tournure, un contexte qui leur retirent leur « virulence ». La qualité de la base d'apprentissage doit être particulièrement soignée pour que cette classification prétende au même niveau de performance que les précédentes. Mais, ce n'était pas l'objectif de cette étude. Nous souhaitions montrer dans ce chapitre dédié à une application où l'erreur n'est pas cruciale, que la mise en œuvre de la méthode était accessible à tous et qu'une phase d'apprentissage à « faible coût » permettait déjà d'avoir des résultats souvent amplement suffisants. En effet, l'une des raisons majeures qui dédramatise l'erreur de classification est que l'objectif est de faire remonter au gestionnaire du vidéo-club les critiques qui soutiennent sa politique d'acquisition, puis, seuls 2 à 5 % de ces critiques seront utilisés pour justifier de cette politique auprès de ses adhérents et publiées dans la gazette du vidéo-club.

	Très négatif	Négatif
Très négatif	1158	49
Négatif	283	351

Ce classifieur permet donc d'obtenir un classement de chaque phrase en très négatif ou négatif. Le résultat de cette étape est donc une note définitive TN ou N qui sera transformée pour les calculs d'agrégation, en note sur 20 : 3 pour TN et 8 pour N.

Classification P/TP. La procédure est parfaitement symétrique à la précédente. L'index du classifieur est de 170 lemmes contre les 12131 de l'index complet. L'ensemble d'apprentissage $E_{P-TP} = E_P \cup E_{TP}$ comprend 283 phrases constituées de 97 phrases appartenant à la catégorie très positif TP, 186 phrases appartenant à la catégorie positif. Les performances sont tout à fait favorables puisqu'au total le classifieur classe correctement 73% des phrases.

	Précision	Rappel	F-score
Classe très positif	56,5%	89,7	69,3
Classe positif	92,2%	64	75,6

	Très positif	positif
Très positif	87	10
positif	67	119

Ce classifieur permet donc d'obtenir un classement de chaque phrase en très positif et positif. Le résultat de cette étape est donc une note définitive très positif TP ou positif P qui sera transformée pour la suite des traitements du SIAD en note sur 20 : 18 pour TP et 13 pour P.

2.3.5. Variation des notes calculées par la notation automatique

Pour illustrer la variation du résultat de la notation hiérarchique nous montrons ci-dessous les notations pour le texte suivant déjà mentionné précédemment :

« Le film est très maladroit. Hélas pour ajouter à cela, l'histoire est très mal construite. L'interprétation est malheureusement décevante. ».

La notation automatique est commandée par l'intermédiaire du bouton « NOTER ». Le résultat est composé de trois notes affichés dans les zones : « Note R », « Note S », « Note A ». Le classifieur a évalué chaque phrase comme indiqué ci-dessous.

Phrase	Note
Critère Réalisation :	3
Le film est très maladroit.	
Critère Scénario :	3
Hélas pour ajouter à cela, l'histoire est très mal construite.	
Critère Acteur :	8
L'interprétation est malheureusement décevante.	

2.3.6. Synthèse

L'évaluation précédente termine le processus d'acquisition d'une nouvelle critique. Il ne manquait que le champ ρ des CAs générées depuis la critique, le voici maintenant renseigné. Les processus mis en œuvre ne requiert pas d'intervention humaine une fois les bases d'apprentissage renseignées. Un simple script assure l'enchaînement des classifieurs sur des bases de critiques disponibles sur le net.

La dernière étape est alors de reporter ces éléments sur la grille d'analyse multicritère par le bouton « reporter notes sur grille ». Nous verrons cette fonction au paragraphe 3.6.

3. Agrégation multicritère, justification et identification

Chaque CA est associée à un score ρ et un couple *critère x alternative* (α, φ) . L'évaluation de φ sous l'angle d'analyse de α est le résultat de l'agrégation de tous les scores ρ des CAs relatives au couple (α, φ) . C'est le score partiel ρ_α de φ au regard de α . Les commentaires de ces CAs reflètent et expliquent le résultat numérique de l'agrégation. Dans un second temps, l'agrégation des ρ_α selon un opérateur défini au préalable donne accès à l'évaluation globale de φ . Les fonctions d'agrégation que nous avons choisies ici sont la moyenne arithmétique au niveau d'un couple (α, φ) et une moyenne pondérée pour l'agrégation qui rend le score global. Nous reviendrons sur l'identification des poids de la moyenne pondérée sur la base d'une enquête auprès des adhérents du vidéo-club au paragraphe 3.3.1. D'autres fonctions d'agrégation plus élaborées ont été proposées dans le chapitre II de ce manuscrit, pour lesquelles nous avons étendu les concepts de justification et de risque décisionnel, mais l'objet de ce chapitre est davantage de mettre en avant la chaîne de traitement complète depuis l'acquisition des critiques cinématographiques jusqu'au contrôle par l'incertitude épistémique que de manipuler des outils mathématiques complexes.

Le concept d'Interface Homme/machine au centre de cette étape est la grille multicritère introduite au chapitre II et présentée également dans [Planté *et al.*, 2005b]. Reprenons ce concept d'IHM pour l'application au vidéo-club. Nous allons montrer dans cette section comment nous utilisons les CAs dans l'analyse multicritère.

Notre application met en œuvre plusieurs grilles :

- la grille Film/Date constituant la grille générale de notre application : elle donne les scores des films en fonction du temps ;
- la grille Film/Critère à une date donnée : elle donne les scores partiels des films à une date fixée ;

- la grille Critère/Date pour un film donné : elle donne les scores partiels du film en fonction du temps.

Ces trois grilles constituent l'ossature de l'IHM du SIAD supportant toute la chaîne de traitement, de l'introduction et la gestion des critiques dans la base de connaissance à l'argumentation automatique et la sélection des meilleurs candidats. Plusieurs scénarii en illustreront l'utilisation.

3.1. Présentation générale de l'IHM

La Figure VI.3 est la grille Film/Critère au 15 mai 2004. Elle croise les films en « compétition » et les critères de jugement.

Fichier	classifieur	Risque	lepactedesloups	loveactually	immortel	meursunaautrejour	cube
			20	11	8	7	9
Réalisation			9.7	13.0	18.0	13.0	18.0
			10	3	2	3	3
Acteur			10.4	15.0	16.333334	13.0	16.333334
			5	5	3	3	3
Scénario			12.0	14.666667	16.333334	13.0	14.666667
			5	3	3	1	3

Figure VI.3 : Grille multicritère Film/Critère

Chaque colonne de la grille représente un film. Les lignes de la grille représentent les critères de jugement permettant l'évaluation du film.

L'affichage contenu dans chaque case de la grille propose trois indicateurs :

- une note au centre de la cellule qui est l'agrégation de tous les scores des CAs correspondant au critère désigné par la ligne et au film désigné par la colonne. Il s'agit ici de la moyenne arithmétique des scores des CAs. C'est un nombre compris entre 0 et 20 qui donne l'évaluation du film désigné par la colonne pour le critère désigné par la ligne ;
- une couleur qui d'un point de vue purement informatif est redondante avec la valeur moyenne des scores des CAs de la case. La correspondance entre les couleurs attribuées et les scores est explicitée plus loin à la Figure VI.4 ;
- un nombre en bas à droite de la case indique le nombre de CAs cumulées pour le critère considéré à la date d'édition de la grille.

Les couleurs permettent très rapidement d'avoir une évaluation visuelle de chaque candidat et une vision plus synthétique de la position relative de chaque candidat et ce, pour chaque critère.

3.1.1. Explication des codes de couleurs

Chaque case présente une couleur qui varie graduellement du rouge au vert en fonction du score (voir Figure VI.4). La variation graduelle est structurée autour de 5 couleurs :

- la couleur rouge pour la note 0,
- la couleur orange pour la note 5,
- la couleur jaune pour la note 10,
- la couleur verte pour la note 15,
- la couleur vert foncé, pour la note 20.

Les autres couleurs ne sont que des barycentres de deux couleurs adjacentes.



Figure VI.4 : Notes associées aux couleurs de la grille

3.1.2. Généralisation de la grille Film/Critère

La grille multicritère Film/Critère est en fait une vision détaillée d'une grille plus générale qui présente de façon plus synthétique les différents candidats à évaluer en fonction du temps. Ce véritable tableau de bord est l'objet du prochain paragraphe qui présente l'enchaînement des différentes grilles multicritères.

3.2. Utilisation des grilles multicritères

La grille Film/Date donne les scores des films en fonction du temps. Elle constitue le tableau de bord de plus haut niveau pour suivre la cote des films analysés. D'un coup d'œil, il est facile pour le gestionnaire du vidéo-club de savoir en temps réel quels sont les films les plus appréciés de la critique, mais aussi d'être renseigné sur les tendances d'évolution des scores (par exemple, un film boudé par la critique officielle à sa sortie, mais qui, grâce au bouche à oreille, finit par recueillir un franc succès en salle, aura un score initial médiocre mais qui ne cessera de s'améliorer ; pour le gestionnaire, cela est synonyme de gros succès populaire).

La grille Film/Date permet d'aller de la vision du processus la plus synthétique qui soit jusqu'à la vision détaillée de chaque élément participant à la décision via les fonctionnalités d'explication. Ainsi, si l'on considère le scénario de la Figure VI.5, sur la première grille (la grille Film/Date), on peut observer l'évolution des scores des cinq films en compétition sur 6 semaines avec une période d'échantillonnage d'une semaine. *Immortel* et *Cube* se détachent au temps courant (15 mai 2004), alors que *Le pacte des loups* se tient bon dernier. Le gestionnaire du vidéo-club peut alors chercher à savoir pourquoi un film aussi médiatisé que *Le pacte des loups* accuse un tel retard. Il demande alors une explication au SIAD qui lui retourne la grille Critère/Date pour *Le pacte des loups*. La réalisation apparaît alors comme le point faible du film. Depuis la grille Film/Date, le gestionnaire peut aussi s'interroger sur les raisons qui ont fait que le classement au 15 mai 2004 est tel qu'il est. Le SIAD renvoie cette fois-ci une grille d'explication au 15 mai 2004 sous la forme d'une grille Film/Critère qui explicite les scores des films par leurs scores partiels relativement aux trois critères d'évaluation.



Figure VI.5 : Les grilles multi-critères et l'explication

3.3. La grille Film/Date

Cette grille (Figure VI.6) donne au gestionnaire la vision globale de l'évolution des scores des films. C'est la tableau de bord ou l'IHM de supervision. La cartographie visuelle est toujours plus facile à interpréter que le seul résultat numérique. Cette grille se lit de la façon suivante :

- Les critiques ont été intégrées dans la base de connaissance de façon continue sur une période du 10 avril au 15 mai 2004 ;
- Les calculs d'agrégation des scores sont effectués sur une période d'échantillonnage d'une semaine.

SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS								
Fichier	classifieur	Risque	10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04
lepackedeslousps			8.0	8.0	8.75	8.75	9.916667	10.316667
			6	6	13	13	15	20
loveactually			13.0	13.0	13.0	13.416667	13.416667	13.944444
			5	5	5	9	9	11
immortel			18.0	18.0	18.0	18.0	18.0	17.166668
			3	3	3	3	3	8
meursunautrejour			13.0	13.0	13.0	13.0	13.0	13.0
			3	3	3	3	3	7
cube			15.5	15.5	15.5	15.5	15.5	16.88889
			3	3	3	3	3	9

Figure VI.6 : Grille film/date

Les colonnes de la grille sont les dates d'évaluation, les lignes de la grille les films à analyser.

L'affichage contenu dans chaque case contient trois informations :

- une note au centre de la cellule qui est l'agrégation des scores partiels relativement aux trois critères pour le film considéré. Nous avons retenu la moyenne pondérée à

cette étape de la fusion des CAs. Nous détaillons l'expérimentation mise en œuvre par le gestionnaire du vidéo-club pour déterminer les paramètres de cette moyenne pondérée dans le paragraphe qui suit ;

- une couleur qui ne fait que reprendre visuellement le résultat de la moyenne pondérée ;
- un nombre en bas à droite de la case qui indique le nombre de CAs cumulées jusque la date indiquée par la colonne.

3.3.1. Mode de calcul de la moyenne pondérée

Le score partiel d'un film relativement au critère c_i est la moyenne arithmétique des scores de ses CAs relativement à c_i .

Soit u_i^j le score partiel du film f^j selon le critère c_i :

$$u_i^j = \frac{1}{n_i^j} \sum_{k=1}^{n_i^j} CA_{i,k}^j(\rho) \quad \text{VI-1}$$

où n_i^j est le nombre de critiques cumulées à une date donnée par le film f^j selon le critère c_i et où $CA_{i,k}^j(\rho)$ est le score de la $k^{ième}$ critique du film f^j selon le critère c_i .

Chaque case de la grille Film/date donne l'évaluation d'un film f^j à une date donnée. Cette évaluation est la moyenne pondérée des scores partiels u_i^j .

Soit u^j le score de f^j : $u^j = \sum_{i=1}^{|C|} p_i \cdot u_i^j$, avec $|C|$ le cardinal de l'ensemble des critères, p_i le poids associé au critère c_i , avec $\sum_{i=1}^{|C|} p_i = 1$.

Les poids p_i sont supposés avoir été définis au préalable. On peut imaginer que ces poids sont imposés par les canons du cinéma. Mais, on peut également penser que ces poids permettent d'appréhender les goûts et priorités des clients du vidéo-club. Il s'agit alors d'un problème d'identification paramétrique.

3.3.2. Indentification du comportement collectif

Nous proposons une méthode fondée sur la prise en compte des avis des utilisateurs pour identifier les paramètres (les poids) de la moyenne pondérée. En effet, le gestionnaire du vidéo-club cherche à déterminer combien d'exemplaires de chaque film sélectionné il doit commander, il semble évident que cela dépend des préférences cinématographiques de sa clientèle. Les poids de chacun des critères se doivent de représenter les priorités de la clientèle. Sur la base d'une campagne d'évaluation, nous proposons le calcul automatique de ces poids. Nous proposons également d'identifier les membres les plus représentatifs du comportement de la clientèle. En effet, ces « prototypes » permettent d'anticiper les tendances de la consommation par la clientèle du vidéo-club et donc de mieux acheter.

Pour constituer la base de données nécessaire à cet apprentissage, le gestionnaire du vidéo-club propose à quelques uns de ses adhérents de participer à une enquête. Pour un ensemble de films-tests donné, il demande à chacun d'attribuer pour les films qui lui sont soumis :

- Un score partiel par critère, soit trois notes par film ;

- Un score global d'appréciation du film.

3.3.2.1. Identification des poids de la moyenne pondérée

On suppose donc que le modèle qui lie un score global aux scores partiels des critères est linéaire et on cherche à déterminer les coefficients de pondération de chacun des critères. Pour résoudre ce problème nous utilisons le principe d'approximation des moindres carrés, en cherchant à minimiser l'erreur quadratique.

Nous définissons les symboles suivants :

- i l'indice du critère,
- $|C|$ le cardinal de l'ensemble des critères,
- m le nombre de tests, c'est-à-dire le nombre d'évaluations retournées par les clients,
- p_i le poids associé au critère c_i ,
- x_i^j le score attribué selon le critère i lors de l'évaluation j ;
- y^j le score global attribué lors de l'évaluation j .

Nous utilisons la méthode des multiplicateurs de Lagrange. Pour résoudre le problème, on forme donc le Lagrangien suivant : $L = (\sum_{j=1}^m ((\sum_{i=1}^{|C|} p_i x_i^j) - y^j)^2) - \lambda ((\sum_{i=1}^{|C|} p_i) - 1)$. On cherche à

minimiser l'erreur quadratique $(\sum_{j=1}^m ((\sum_{i=1}^{|C|} p_i x_i^j) - y^j)^2)$ sous la contrainte $\sum_{i=1}^{|C|} p_i = 1$.

On annule les dérivées partielles de L , c'est-à-dire les $\frac{\partial L}{\partial p_i}$ et $\frac{\partial L}{\partial \lambda}$:

$$\frac{\partial L}{\partial p_i} = 0 \Rightarrow \frac{\partial L}{\partial p_i} = 2 \cdot p_i \cdot \sum_{j=1}^m x_i^{j2} + 2 \cdot (\sum_{i \neq k} (p_k \cdot (\sum_{j=1}^m x_i^j x_k^j))) - 2 \cdot \sum_{j=1}^m x_i^j \cdot y^j - \lambda = 0 \text{ pour tout } i.$$

$$\text{et } \frac{\partial L}{\partial \lambda} = 0 \Rightarrow (\sum_{i=1}^{|C|} p_i) - 1 = 0.$$

On a donc $m+1$ équations pour les $m+1$ inconnues $p_i, \forall i$ et λ .

Ce qui se traduit sous la forme matricielle :

$$\begin{bmatrix} \sum_{j=1}^m x_1^{j2} & \sum_{j=1}^m x_1^j x_2^j & \dots & \sum_{j=1}^m x_1^j x_i^j & \dots & \sum_{j=1}^m x_1^j x_{|C|}^j & -\frac{1}{2} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \sum_{j=1}^m x_i^j x_1^j & \sum_{j=1}^m x_{ij} x_{2j} & \dots & \sum_{j=1}^m x_{ij}^2 & \dots & \sum_{j=1}^m x_i^j x_{|C|}^j & -\frac{1}{2} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \sum_{j=1}^m x_{|C|}^j x_1^j & \sum_{j=1}^m x_{|C|}^j x_2^j & \dots & \sum_{j=1}^m x_{|C|}^j x_i^j & \dots & \sum_{j=1}^m x_{|C|}^j x_{|C|}^j & -\frac{1}{2} \\ 1 & 1 & \dots & 1 & \dots & 1 & 0 \end{bmatrix} \begin{bmatrix} p_1 \\ p_2 \\ \cdot \\ \cdot \\ p_{|C|} \\ \lambda \end{bmatrix} = \begin{bmatrix} \sum_{j=1}^m x_1^j y^j \\ \sum_{j=1}^m x_2^j y^j \\ \cdot \\ \cdot \\ \sum_{j=1}^m x_{|C|}^j y^j \\ 1 \end{bmatrix} \quad \text{VI-2}$$

Que l'on notera $\hat{A}\hat{P} = B$ où $\hat{P}$ étant le vecteur estimé des poids.

La résolution de ce système linéaire fournit le modèle linéaire recherché, reliant scores partiels et scores globaux. Ce modèle est la meilleure approximation au sens des moindres carrés que l'on puisse obtenir pour rendre compte de la façon dont les adhérents du vidéo-club agrègent les scores partiels pour obtenir leur évaluation globale.

Dans notre application, certains clients se sont donc prêtés à une évaluation complète sur une base de films-tests. Les résultats sont les suivants : le poids de *Réalisation* est 0.5, celui de *Acteur* est 0.333 et celui de *Scénario* 0.167.

3.3.2.2. Identification des prototypes

A ce stade, nous avons montré qu'il existe un opérateur d'agrégation qui synthétise de façon satisfaisante la façon dont les clients attribuent une note globale à un film depuis les notes partielles de ce film. Maintenant, pour identifier les individus les plus représentatifs de la population du vidéo-club, nous procédons en deux temps. Tout d'abord nous tentons d'identifier qualitativement s'il existe des groupes homogènes d'adhérents, puis nous utilisons une méthode d'identification plus précise de ces groupes, et enfin nous cherchons les individus les plus représentatifs de ces groupes.

Pour un film donné, un individu est représenté par les trois scores partiels qu'il attribue au film (on n'utilise pas le score global puisque, de par l'étape précédente, on considère que celui-ci n'est qu'une combinaison linéaire des trois scores partiels indépendamment de l'individu). Sur cet espace à trois dimensions et toujours pour un film donné, on procède alors à une analyse en composante principale (ACP). Elle permet de se fixer simplement sur l'existence ou non de groupes de clients ayant des goûts cinématographiques voisins.

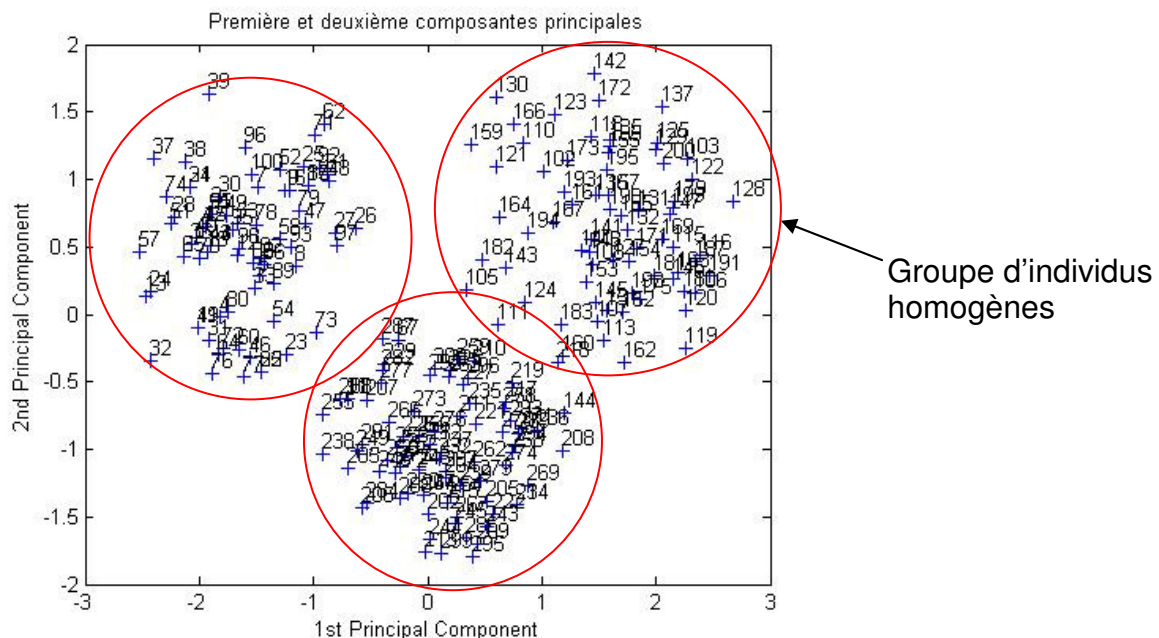


Figure VI.7 : Une vue de l'ACP sur l'espace initial des évaluations partielles.

L'ensemble des vues de l'ACP, pour notre application, montre qu'il existe trois comportements évaluatifs bien distincts pour le film soumis à la critique. Pour affiner notre catégorisation, nous avons recours à la technique de « Clustering Hiérarchique » [Hartigan, 1975; Jain *et al.*, 1999].

Cette méthode consiste à calculer la distance entre les n individus à regrouper. Les deux individus les plus proches sont remplacés par un barycentre (par défaut leur milieu). Il reste alors $n-1$ points, les distances sont calculées à nouveau. On réitère l'opération autant que nécessaire, l'ACP ayant fourni le nombre de clusters à rechercher. Une hiérarchie de regroupement se constitue ainsi comme indiqué sur la Figure VI.8.

A la suite de ces calculs, nous obtenons les groupes homogènes d'individus dont on connaît par construction le barycentre (« individu fictif »). Il suffit alors de choisir les individus les plus proches de ce barycentre pour prototypes de la clientèle.

Sur l'exemple, pour chacun des trois clusters, quelques individus seront retenus par le gérant du vidéo-club parce qu'ils ont des goûts caractéristiques d'une sous population de la clientèle. Lorsque la campagne est menée sur m films, on peut faire rentrer en ligne de compte le nombre de fois (m au maximum) qu'un individu apparaît parmi les individus les plus proches des barycentres des clusters (renforcement du prototype).

Ainsi, on peut imaginer que le gérant du vidéo-club puisse offrir des places de cinéma à « ses » prototypes en échange de leur avis sur les films. On peut alors imaginer que parmi les critiques gérées par notre système, une attention toute particulière soit donnée à celles rédigées par les « prototypes ».

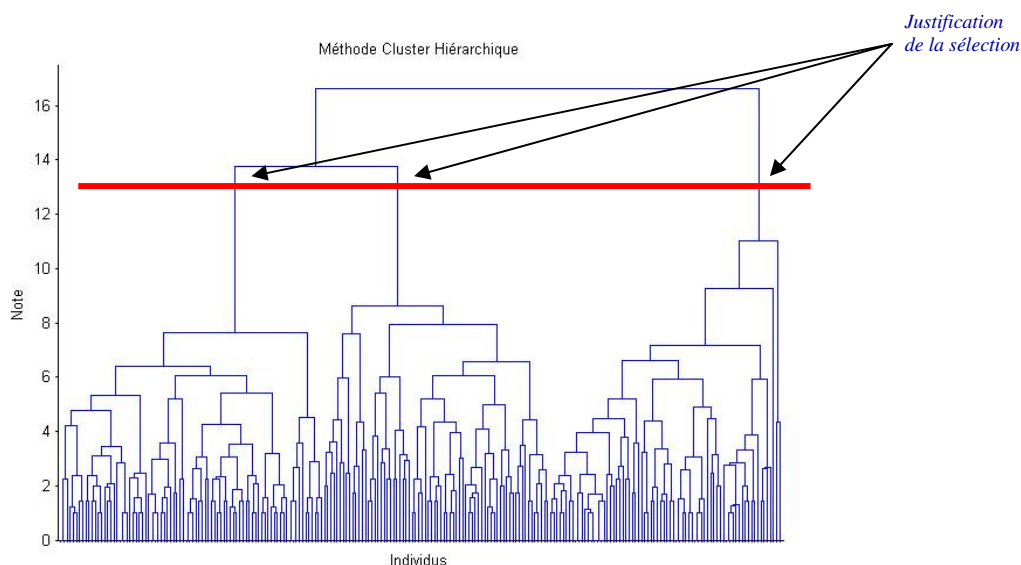


Figure VI.8 : La méthode du « clustering hiérarchique » sur la base test des clients du vidéo-club.

3.4. La grille Fim/Critère

La grille Film/Critère (Figure VI.9) à une date donnée donne les scores partiels des films à cette date. Les scores partiels affichés dans une case de la grille correspondent à la moyenne arithmétique des scores élémentaires des CAs cumulées jusque cette date, de cette même case.

Fichier	classifieur	Risque
	le pacte des loups	20
Réalisation	9.7	10
Acteur	10.4	5
Scénario	12.0	5

Figure VI.9 : La grille Film/critère au 15 mai 2004

Cette grille peut être vue comme l'explication de la grille film/date à une date donnée : elle permet d'expliquer le score global d'un film par ses scores partiels. La couleur de la case où s'inscrit le nom du film correspond à son score global. Le score global est la somme des contributions marginales de chacun des critères (voir chapitre II). Les indicateurs contenus dans chaque case ont la même sémantique que dans la grille précédente (score, couleur, nombre de critiques cumulées).

A partir de chaque case, on peut obtenir en cliquant sur le bouton droit de la souris, le menu de la Figure VI.11 qui affiche la liste des scores des CAs cumulées disponibles dans cette case et le titre des CAs (par défaut les premiers mots des commentaires des CAs). Il est bien sûr possible à partir de ce dernier menu d'afficher la critique entière.

Autrement dit, le gestionnaire du vidéo-club peut s'intéresser à l'évaluation d'un film en particulier. La grille Film/date lui donne l'évaluation globale du film et son classement. S'il veut en savoir plus sur le mode d'évaluation du film, il appelle la grille Film/critère ci-dessus. Il peut ainsi voir graphiquement les points forts et faibles du film qu'il analyse à un premier niveau d'explication. Dans un second temps, il peut pousser son investigation en consultant les critiques qui ont valu au film tel ou tel score partiel. Il peut faire en parallèle cette analyse sur plusieurs films et voir ainsi quels sont les atouts et les faiblesses d'un film par rapport à un autre.

3.5. La grille Critère/Date

La grille Critère/Date pour un film donné donne les scores partiels d'un film en fonction du temps. Elle peut être vue comme une explication de la grille Film/date. Sur la grille film/date le gestionnaire du vidéo-club cherche à comprendre l'évolution du score global d'un film dans le temps. Cette grille lui fournit cette information critère par critère. La notion de contribution marginale des critères est la base de l'explication. Là encore, un simple clic droit sur la souris lui donne accès aux critiques cumulées disponibles pour approfondir sa compréhension de l'évolution du score du film. La sémantique des indicateurs de chaque case est la même que dans les grilles précédentes.

Fichier	classifieur	Risque
	10.04.04	17.04.04
Réalisation	5.5	5.5
Acteur	10.5	10.5
Scénario	10.5	10.5

Figure VI.10 : La grille Critère/date

A partir de chaque case on peut obtenir en cliquant sur le bouton droit le menu de la Figure VI.11 qui affiche la liste des notes des CAs disponibles dans cette case avec le titre de la CA (par défaut le début du commentaire). Il est bien sûr possible d'avoir accès à chacune des CAs (Figure VI.12).

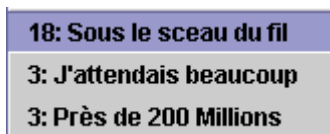


Figure VI.11 : Affichage des CA pour une case

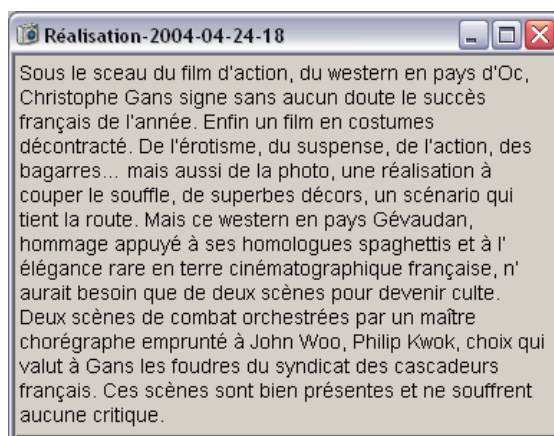


Figure VI.12 : Affichage d'une CA suite à une demande d'explication

3.6. Report de nouvelles critiques sur la grille

Lorsque l'on utilise la fonction de saisie d'une nouvelle critique vue au paragraphe 2.3. Cette critique est découpée en CAs. Elle peut générer jusque trois (nombre de critères) CAs avec leur score associé. Ces notes sont reportées dans la grille Film/date par le bouton « reporter notes sur grille » de la Figure VI.2. Dans notre exemple nous obtenons la grille mise à jour à la date du 29 mai 2004 de la Figure VI.13. Par exemple, on voit donc que trois nouvelles critiques ont été introduites dans la base entre le 15 mai 2004 et le 29 mai 2004 pour le film *Le pacte des loups*. L'affichage de la grille explicative Critère/Date nous permettrait de savoir sur quels critères ces CAs ont été reportées. Les 3 nouvelles critiques ont fait évoluer le score du film de 10.3 à 11.3. Le gestionnaire peut avoir accès comme expliqué précédemment aux trois critiques qui ont permis cette amélioration.

SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS									
Fichier	classifieur	Risque							
		10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04	22.05.04	29.05.04
lepactedesloups		8.0	8.0	8.75	8.75	9.916667	10.316667	10.316667	11.282828
		6	6	13	13	15	20	20	23
loveactually		13.0	13.0	13.0	13.416667	13.416667	13.944444	13.944444	13.944444
		5	5	5	9	9	11	11	11
immortel		18.0	18.0	18.0	18.0	18.0	17.166668	17.166668	17.166668
		3	3	3	3	3	8	8	8
meursunaautrejour		13.0	13.0	13.0	13.0	13.0	13.0	13.0	13.0
		3	3	3	3	3	7	7	7
cube		15.5	15.5	15.5	15.5	15.5	16.88889	16.88889	16.88889
		3	3	3	3	3	9	9	9

Figure VI.13 : La grille Film/Date au 29 mai 2004

4. Evaluation du risque décisionnel

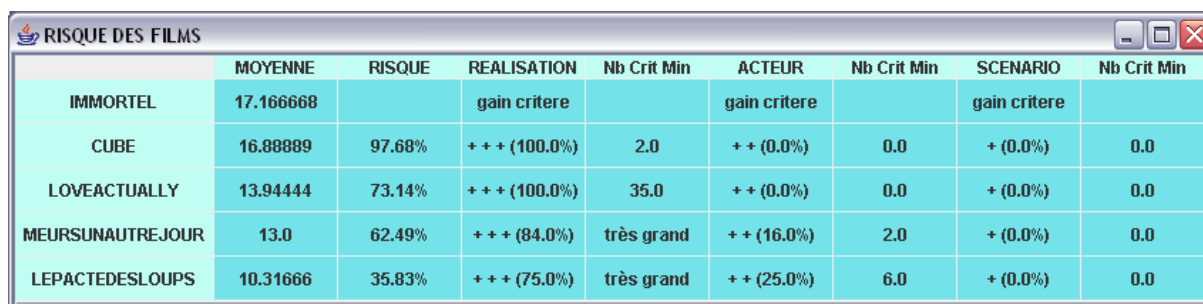
Le risque décisionnel que nous avons défini au chapitre II

($r = 1 - \min_{\substack{k=2..n \\ k \neq 1}} \frac{d(f^1, f^k)}{p}$, f^1 l'alternative préférée) est lié à la stabilité du classement des

alternatives : plus l'alternative préférée distance les autres concurrents, plus il sera difficile d'inverser le classement en ajoutant de nouvelles CAs. Le risque ainsi défini peut également être assimilé à une mesure de l'incertitude épistémique inhérente à la situation de décision : compte tenu des connaissances actionnables (extraites des critiques cinématographiques) dont dispose le gérant du vidéo-club à un instant donné, dans quelle mesure celui-ci peut-il être certain de faire le bon choix, c'est-à-dire qu'un revirement de tendance est des plus improbables ?

4.1. L'IHM de la fonction « risque »

L'écran « risque » à une date donnée rappelle d'abord le classement des films à cette date puisque dans la formule de calcul du risque, les distances sont toujours calculées par rapport à l'alternative préférée (*Immortel* sur la Figure VI.14). Les films sont rangés dans l'ordre décroissant des scores globaux.



	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTEL	17.166668		gain critere		gain critere			
CUBE	16.88889	97.68%	+++ (100.0%)	2.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LOVEACTUALLY	13.94444	73.14%	+++ (100.0%)	35.0	++ (0.0%)	0.0	++ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	62.49%	+++ (84.0%)	très grand	++ (16.0%)	2.0	++ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	35.83%	+++ (75.0%)	très grand	++ (25.0%)	6.0	++ (0.0%)	0.0

Figure VI.14 : Ecran « risque » au 15 mai 2004

La description des colonnes du tableau de la Figure VI.14 est la suivante :

- La première colonne propose les films classés de haut en bas par ordre décroissant des scores ;
- La deuxième colonne donne le score du film ;
- La troisième colonne donne le risque à choisir le premier film plutôt que le $k^{ième}$. Le calcul est un simplex élémentaire donné au paragraphe 4.2 ;
- La quatrième (resp. 6^{ème} et 8^{ème}) colonne donne l'importance relative du critère *Réalisation* (resp. *Acteur* et *Scénario*) dans l'évaluation du risque. Nous reviendrons sur l'indicateur double (« + » versus « % ») ;
- Les cinquièmes, septièmes et neuvièmes colonnes indiquent « qualitativement » le nombre minimal de CAs positives sur l'ensemble des critères qu'il serait nécessaire d'acquérir pour que le classement puisse s'en trouver modifié ;

Ce tableau de bord fixe à un instant donné l'état de la situation décisionnelle. Le gérant du vidéo-club peut d'un seul coup d'œil savoir si le classement est stable ou non, quelles sont les dimensions sensibles entre deux films, etc.

4.2. Calcul du risque

Compte tenu que nous avons choisi comme opérateur d'agrégation une moyenne pondérée pour illustrer de manière pédagogique les concepts théoriques du chapitre II, les calculs et algorithmes sont extrêmement simplifiés.

Le calcul du risque repose sur la notion de distance entre l'alternative préférée et les autres concurrents.

Soit f^1 le film en tête du classement et $f^j, \forall j$ les challengers à une date donnée. Chaque film f^j est défini par son vecteur de scores partiels $\mathbf{u}^j = [u_1^j \dots u_i^j \dots u_n^j]^T$ et son score global

$$u^j = \sum_{i=1}^{|C|} p_i \cdot u_i^j.$$

Le calcul de la distance entre le film f^1 et le film f^j s'énonce sous la forme du problème d'optimisation suivant :

Objectif: $F\Delta^j = \min \|\delta^j\|_1 = \min \left(\sum_{i=1}^n \delta_i^j \right)$

où δ_i^j est l'amélioration selon la dimension i de l'évaluation, l'indice j fait référence à la solution j , $\mathbf{u}^j = [u_1^j, \dots, u_p^j]^T, \delta^j = [\delta_1^j, \dots, \delta_p^j]^T$ pour $j \neq 1$

Contraintes: $MP(\mathbf{u}^j + \delta^j) = MP(\mathbf{u}^1)$, où MP est la moyenne pondérée

$\forall i, 0 \leq \delta_i^j \leq 1 - u_i^j.$

Soit $\bar{\delta}^j = [\bar{\delta}_1^j, \dots, \bar{\delta}_p^j]^T$ la solution du problème (soit $F\Delta^j = \|\bar{\delta}^j\|_1$).

Lorsque MP désigne la moyenne pondérée, les contraintes sont linéaires et l'algorithme n'est autre qu'un simplex.

Le risque décisionnel s'écrit alors : $r = 1 - \min_j \frac{F\Delta^j}{|C|}$ VI-3

La distance $\min_j F\Delta^j$ est d'autant plus faible que le risque est grand. La division par

$|C|$ permet de ramener $\min_j F\Delta^j$ dans $[0 ; 1]$. Ainsi, $r = 1 - \min_j \frac{F\Delta^j}{|C|} \in [0;1]$. Par convention,

nous avons choisi d'afficher le risque en pourcentage (colonne 3 de la Figure VI.14).

La Figure VI.15 montre une évolution de la grille risque après avoir ajouté des critiques positives concernant le film *Le pacte des loups*. On s'aperçoit que le classement n'a pas changé, mais que le score global du film *Le pacte des loups* ayant significativement augmenté, la distance $F\Delta^{\text{Pacte des loups}}$ entre *Immortel* et *Le pacte des loups* a diminué.

RISQUE DES FILMS								
	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTEL	17.166668		gain critere		gain critere		gain critere	
CUBE	16.88889	97.68%	+++ (100.0%)	2.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LOVEACTUALLY	13.94444	73.14%	+++ (100.0%)	35.0	++ (0.0%)	0.0	++ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	62.49%	+++ (84.0%)	très grand	++ (16.0%)	2.0	++ (0.0%)	0.0
LEPACTEDESLOUPS	12.59935	60.37%	+++ (92.0%)	très grand	++ (8.0%)	2.0	++ (0.0%)	0.0

Figure VI.15 : Ecran risque après adjonction de critiques positives sur *Le pacte des loups*

4.3. Signal de contrôle

$F\Delta^j$ est le nombre minimal de points à gagner par le film f^j pour qu'il puisse remonter à hauteur du film f^1 ; $\bar{\delta}_i^j$ indique le nombre de points à gagner selon le critère i pour une amélioration optimale (au sens du coût informationnel).

Dans cette application, nous considérons que les critères qui ont le plus « coûté » au film f^j sont les critères pour lesquels $p_i \cdot \bar{\delta}_i^j$ est important. Autrement dit, si f^j n'est pas classé premier, c'est principalement à cause des critères pour lesquels $p_i \cdot \bar{\delta}_i^j$ est significatif.

Le ratio $p_i \cdot \bar{\delta}_i^j / \sum_{i=1}^{|C|} p_i \cdot \bar{\delta}_i^j$ donne l'influence (ou la contribution) du critère i sur le déficit accumulé par le film f^j vis-à-vis de f^1 . Il est exprimé en pourcentage.

Ainsi, sur l'exemple de la Figure VI.16, le risque à choisir *Immortel* plutôt que *Love Actually* est de 72.91%. En terme de contribution, le critère *Réalisation* explique prioritairement le déficit de $f^{\text{Love actually}}$ puisque sa contribution s'élève à 60% alors que le critère *Acteur* explique les 40 derniers %. Remarquons néanmoins que si l'on raisonne uniquement sur les $\bar{\delta}_i^j$ (et non pas sur les contributions $p_i \cdot \bar{\delta}_i^j$), alors on s'aperçoit que le nombre de points déficitaires est plus important sur le critère *Acteur* (+ + +) que sur le critère *Réalisation* (+ +). C'est le poids accordé à la *Réalisation* (0.5 contre 0.333 à *Acteur*) qui inverse l'influence des critères. Nous avons choisi de reporter dans la grille risque les deux informations : les contributions au déficit en % (les $p_i \cdot F\Delta_i^j / \sum_{i=1}^{|C|} p_i \cdot F\Delta_i^j$) et les déficits bruts ($F\Delta_i^j$) sous la forme d'un nombre de « + ».

RISQUE DES FILMS								
	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTEL	18.0		gain critere		gain critere		gain critere	
LOVEACTUALLY	13.0	72.91%	+ + (60.0%)	2.0	+ + + (40.0%)	45.0	+ (0.0%)	0.0
LEPACTEDESLOUPS	8.75	52.08%	+ + + (70.0%)	22.0	+ + (30.0%)	5.0	+ (0.0%)	0.0

Figure VI.16 : exemple de grille risque affichée

Le calcul du risque permet donc de déterminer s'il y a lieu ou non d'acquérir de nouvelles critiques : si le risque calculé est au-delà du seuil d'acceptabilité fixé au préalable (la consigne de la boucle de régulation de la Figure VI.1), le classement n'est pas stable (l'incertitude épistémique est trop grande) et de nouvelles critiques sont nécessaires pour valider, éprouver ou étayer le classement courant. Le calcul du risque détermine par ailleurs les critères pour lesquels il est le plus pertinent d'acquérir des compléments d'information : il faut vérifier et valider les scores relatifs aux critères qui expliquent prioritairement le déficit des challengers.

Autrement dit, le calcul du risque fournit le « signal de contrôle » qui permet de lancer les classifieurs du SIAD sur la recherche d'informations pertinentes et nécessaires pour converger vers une situation décidable (c'est-à-dire dont le risque est acceptable). Comme nous l'avons expliqué formellement au chapitre II, les classifieurs, associés à un moteur de recherche, jouent le rôle de l'actionneur de notre boucle de régulation (Figure VI.1).

Le système est donc capable de déterminer de lui-même les dimensions sur lesquelles il n'est pas capable de donner les justifications nécessaires à sa prise de position à une date donnée et pour lesquelles il va donc chercher automatiquement de nouvelles critiques, construire de nouvelles CAs. La procédure se répète jusqu'à ce que le risque tombe en deçà du seuil d'acceptabilité...sans qu'il n'y ait nécessairement d'intervention humaine... C'est le paroxysme de l'automatisation cognitive... Le SIAD apprend pour décider.

4.4. Calcul du nombre de critiques minimum pour modifier le classement

Il s'agit de proposer ici une formule « théorique » pour calculer « qualitativement » le nombre minimal de CAs positives sur l'ensemble des critères qu'il serait nécessaire d'acquérir pour que le classement puisse s'en trouver modifié.

Pour un « challenger » donné, on cherche donc le nombre minimal de CAs positives tous critères confondus pour que le score de ce challenger soit au moins égal à celui du film préféré (en supposant qu'aucune CA ne soit attribuée par ailleurs au leader dans cet intervalle de temps).

Soit f^j le challenger. Au temps $t+1$, on cherche à avoir : $MP(\mathbf{u}^j + \bar{\delta}^j) = MP(\mathbf{u}^1)$ où $\bar{\delta}^j$ est dû à l'apport de CAs positives entre t et $t+1$ pour f^j (δ^j est calculé par l'algorithme précédent).

On note $u_i^j(t)$ le score partiel de f^j selon le critère i au temps t .

Soit n_i^j le nombre de critiques disponibles selon le critère i pour f^j .

On cherche δn_i^j le nombre minimal de critiques positives supplémentaires nécessaires à f^j relativement au critère i pour que :

$$u_i^j(t+1) = u_i^j(t) + \bar{\delta}_i^j = \frac{1}{n_i^j + \delta n_i^j} \sum_{i=1}^{n_i^j + \delta n_i^j} CA_{i,k}^j(\rho) \quad \text{VI-4}$$

En supposant que les δn_i^j CAs recueillies entre t et $t+1$ soient telles que $CA_{i,k}^j(\rho) = 1$ (pour que le nombre de CAs soit minimum), on a alors :

$$u_i^j(t) + \bar{\delta}_i^j = \frac{1}{n_i^j + \delta n_i^j} (\sum_{i=1}^{n_i^j} CA_{i,k}^j(\rho) + \delta n_i^j \cdot 1) = \frac{1}{n_i^j + \delta n_i^j} \cdot (n_i^j \cdot u_i^j(t) + \delta n_i^j).$$

$$\text{Soit : } u_i^j(t) + \bar{\delta}_i^j = \frac{1}{n_i^j + \delta n_i^j} \cdot (n_i^j \cdot u_i^j(t) + \delta n_i^j) \Leftrightarrow \delta n_i^j \cdot (1 - u_i^j(t) - \bar{\delta}_i^j) = n_i^j \cdot \bar{\delta}_i^j$$

$$\text{Finalement, il vient : } \delta n_i^j = \frac{n_i^j \cdot \bar{\delta}_i^j}{1 - u_i^j(t) - \bar{\delta}_i^j} \quad \text{VI-5}$$

Il conviendra d'ajouter $E(\delta n_i^j) + 1$ CAs.

A titre d'exemple, sur la Figure VI.14, il suffit de 2 critiques très favorables à *Cube* sur le critère *Réalisation* pour que le score de *Cube* soit identique à celui du leader *Immortel* à cette date. Sur la Figure VI.16, il faut non seulement 22 critiques sur le critère *Réalisation* mais encore 5 sur *Acteurs* pour que *Le Pacte des Loups* puisse revenir dans la course. La valeur quantitative de $E(\delta n_i^j) + 1$ est d'une utilité limitée, c'est son ordre de grandeur qui est à retenir pour le gérant du vidéo-club.

5. Argumentation

Nous avons expliqué que le gestionnaire doit préparer un argumentaire de sa sélection pour répondre à toute réclamation de sa clientèle (on pourrait dans une situation à enjeu plus marqué, parler de principe de précaution !) et qu'il pouvait pour cela éditer une gazette annonçant les nouveautés acquises, assorties des CAs les plus élogieuses que le SIAD aura « construites » et gérées pour lui. Le SIAD va pouvoir lui fournir automatiquement ces CAs. Justifier la sélection d'un film (ou plus précisément le nombre de cassettes du film acquises par le vidéo-club), c'est expliquer l'excellence du score global du film par les critères qui ont eu une contribution majeure à ce score (argumentation en absolu), la « supériorité » du film sur d'autres concurrents (argumentation en relatif), puis aller rechercher dans le SGDC les commentaires des CAs correspondant à ces arguments. L'édition de la gazette est donc très largement supportée par le SIAD et le SGDC.

5.1. Argumentation en absolu

L'écran d'argumentation en absolu peut être obtenu à partir de la grille Film/Date pour expliquer le bon score d'un film. On y est renseigné sur les critères qui ont le plus contribué au choix d'un film.

L'écran (Figure VI.17) montre plusieurs informations :

- Le paramètre *pourcentage d'explication* : il permet de sélectionner le niveau de détail que l'on souhaite obtenir dans l'explication du score. Plus il est élevé, plus l'explication sera détaillée (attention, le critère de pertinence de l'information prévaut sur celui de quantité, par conséquent une valeur élevée de ce taux n'est pas nécessairement conseillée ! Le gérant ne souhaite faire apparaître dans sa gazette que les faits les plus marquants et ne pas s'attacher au détail). Les explications concernant la détermination de ce paramètre sont données au paragraphe 5.1.1 ;
- Le paramètre *Tolérance* : il permet au gérant du vidéo-club de contrôler le nombre de critiques dont il souhaite disposer pour étayer un choix. Plus la tolérance est grande, moins on est exigeant sur les scores obtenus par un film (par exemple, une tolérance de 60 pourra faire que l'on considère que 12 est un bon score alors qu'avec une tolérance de 25, ce même 12 sera considéré comme décevant). Les explications concernant la détermination de ce paramètre sont données au paragraphe 5.1.2 ;
- Le reste de l'écran est évidemment consacré à l'affichage des CAs sélectionnées triées par critère. Les critères sont affichés par ordre d'importance de leur contribution dans la sélection du film. Sur l'écran, on peut donc lire les critères ayant joué un rôle majeur dans la sélection du film, puis pour chacun de ces critères les CAs associées les plus significatives. Le nombre de critères affichés dépend du paramètre *pourcentage d'explication* et le nombre de CAs du paramètre *Tolérance*.

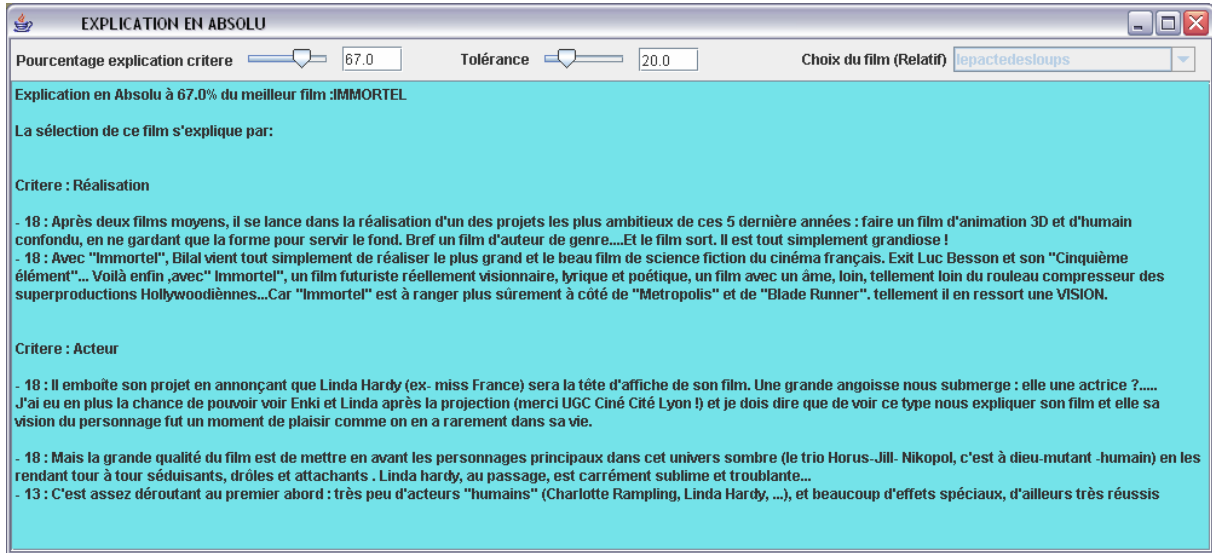


Figure VI.17 : Ecran d'une explication en absolu

5.1.1. Niveau de précision de l'explication absolue

Considérons le film f^1 en tête du classement, u^1 son score global à une date t donnée.

Le paramètre k , « pourcentage d'explication critère », permet de déterminer les critères qui suffisent à expliquer $k\%$ du score de f^1 .

Par définition : $u^1 = MP(f^1) = MP(\mathbf{u}^1) = \sum_{i=1}^{|C|} p_i \cdot u_i^1$.

Après avoir réordonné les termes $p_i \cdot u_i^j$, tels que $\forall i, p_i \cdot u_i^j \geq p_{i+1} \cdot u_{i+1}^j$, on cherche p_0 tel que :

$$\sum_{i=1}^{p_0 \leq |C|} p_i \cdot u_i^1 = k\% \sum_{i=1}^{|C|} p_i \cdot u_i^1 = k\% \cdot MP(f^1) = k\% \cdot u^1 \quad \text{VI-6}$$

Cette équation se lit : « Les p_0 premiers critères expliquent $k\%$ du score obtenu par f^1 ».

Par conséquent, plus k est grand, plus le nombre de critères utilisés pour justifier le score du film est important. Tous les critères n'ont pas forcément de contribution significative sur le score de f^1 et ce n'est vraisemblablement pas ces explications « anecdotiques » que le gérant du vidéo-club souhaite faire apparaître dans son argumentaire.

Sur l'exemple de la Figure VI.17, les critères *Réalisation* et *Acteur* expliquent à 67% le score obtenu par *Immortel*. Sur la Figure VI.18, *Immortel* jouit d'une excellente réalisation, or réalisation est le critère de poids maximal (0.5), par conséquent une explication « grossière » à 40% se résume au seul critère *Réalisation*.

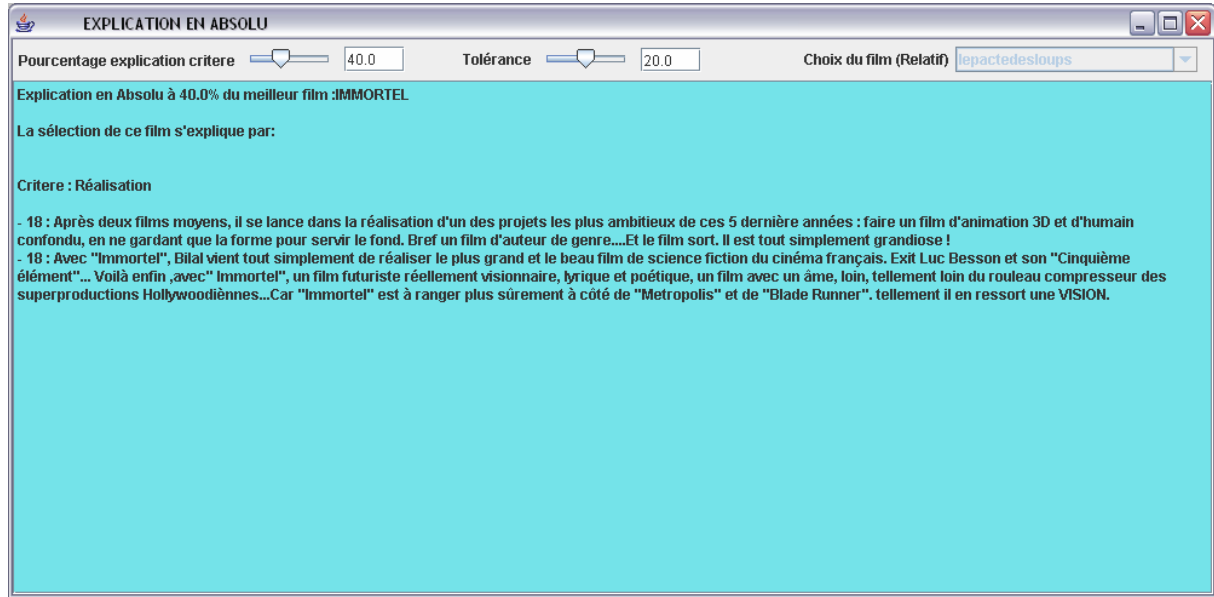


Figure VI.18 : Explication en absolu et paramètre pourcentage d'explication

5.1.2. Niveau de tolérance sur le score d'une CA

Le paramètre *Tolérance* ou ε permet de contrôler le nombre de critiques que l'on veut faire valoir pour un critère. Plus la tolérance est grande, moins on est exigeant sur les scores des CAs relatives au film.

On utilise une version simplifiée du raisonnement aux ordres de grandeur proposé dans le paragraphe 3.2.3 du chapitre II et en particulier à la figure 5 de ce chapitre.

Pour justifier d'un bon score partiel u_i^1 , les CAs que l'on va retenir sont telles que :

$$\forall k, \frac{1}{(1+\varepsilon)^2} \cdot \max_k CA_{i,k}^1(\rho) \leq CA_{i,k}^1(\rho) \quad \text{VI-7}$$

Une valeur commune de ε est 0.1. ε ne peut dépasser 0.4655 (sinon on ne respecterait plus $\varepsilon < \frac{1}{(1+\varepsilon)^2}$ de la figure V du chapitre II).

Ainsi, sur la Figure VI.17, le paramètre tolérance ε est à 20%. Nous observons que pour le critère *Acteur*, une CA dont le score est 13 est considérée comme une appréciation favorable. Si nous diminuons la tolérance ε à 10%, nous obtenons l'écran de la Figure VI.19 dans lequel les notes des critiques affichées pour le critère *Acteur* ne comprennent plus que des CAs dont le score est de 18 : 13 n'est plus considéré comme une bonne évaluation dans ce système d'appréciation plus sévère.

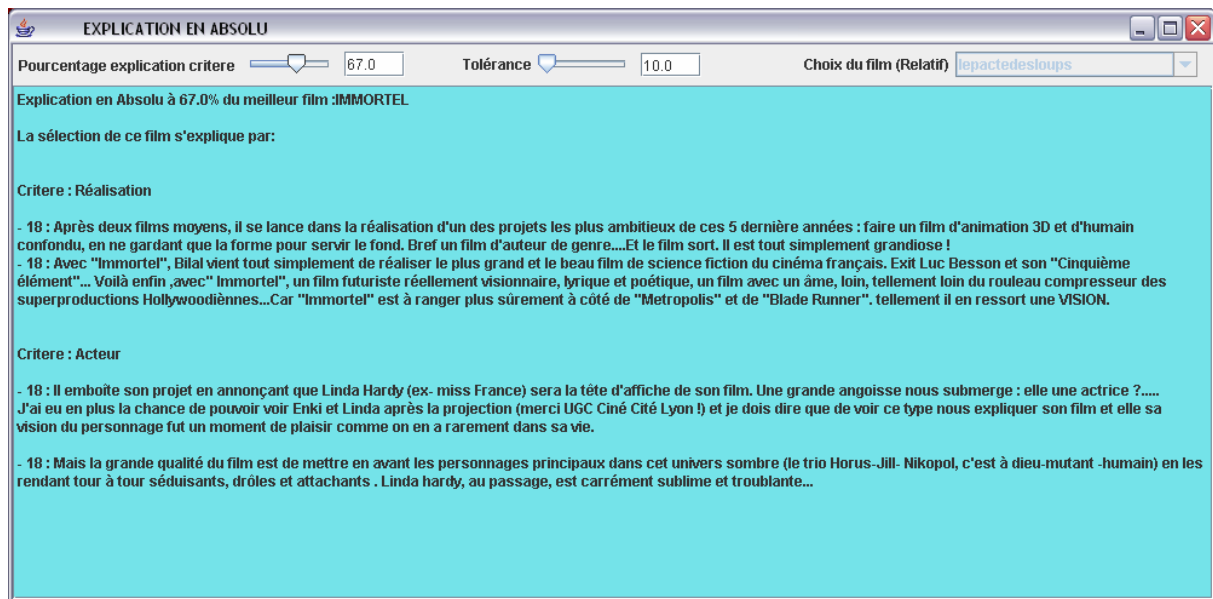


Figure VI.19 : Ecran d'explication et paramètre tolérance

5.2. Argumentation en relatif

L'écran d'argumentation en relatif peut être obtenu à partir de la grille Film/Date pour expliquer la supériorité de f^1 sur ses challengers. On y est renseigné sur les critères qui ont le plus contribué à distinguer f^1 de ses challengers.

Une fois les critères les plus discriminants identifiés, le SIAD recherche dans le SGDC les CAs qui supportent ces arguments.

Sur la vue argumentation en relatif :

- Un champ permet de choisir le film que l'on souhaite comparer à f^1 ;
- Comme pour l'explication en absolu, on peut jouer sur les paramètres $k\%$ et ε pour avoir une explication plus ou moins fine et un nombre de CAs plus ou moins grand ;
- Le reste de l'écran est consacré à la consultation des CAs sélectionnées et triées par critère puis par film. En premier lieu, elles sont triées par critère : le premier critère est celui grâce auquel f^1 s'est le plus significativement distingué de son challenger et ainsi de suite. Puis pour chaque critère sont affichées les CAs associées à f^1 , c'est-à-dire les CAs qui sont favorables à f^1 . Ensuite viennent les CAs du challenger, mais cette fois-ci, le SIAD propose les CAs qui ont le plus pénalisé le challenger sur chacun des critères. Viennent ensuite les mêmes éléments pour les autres critères.

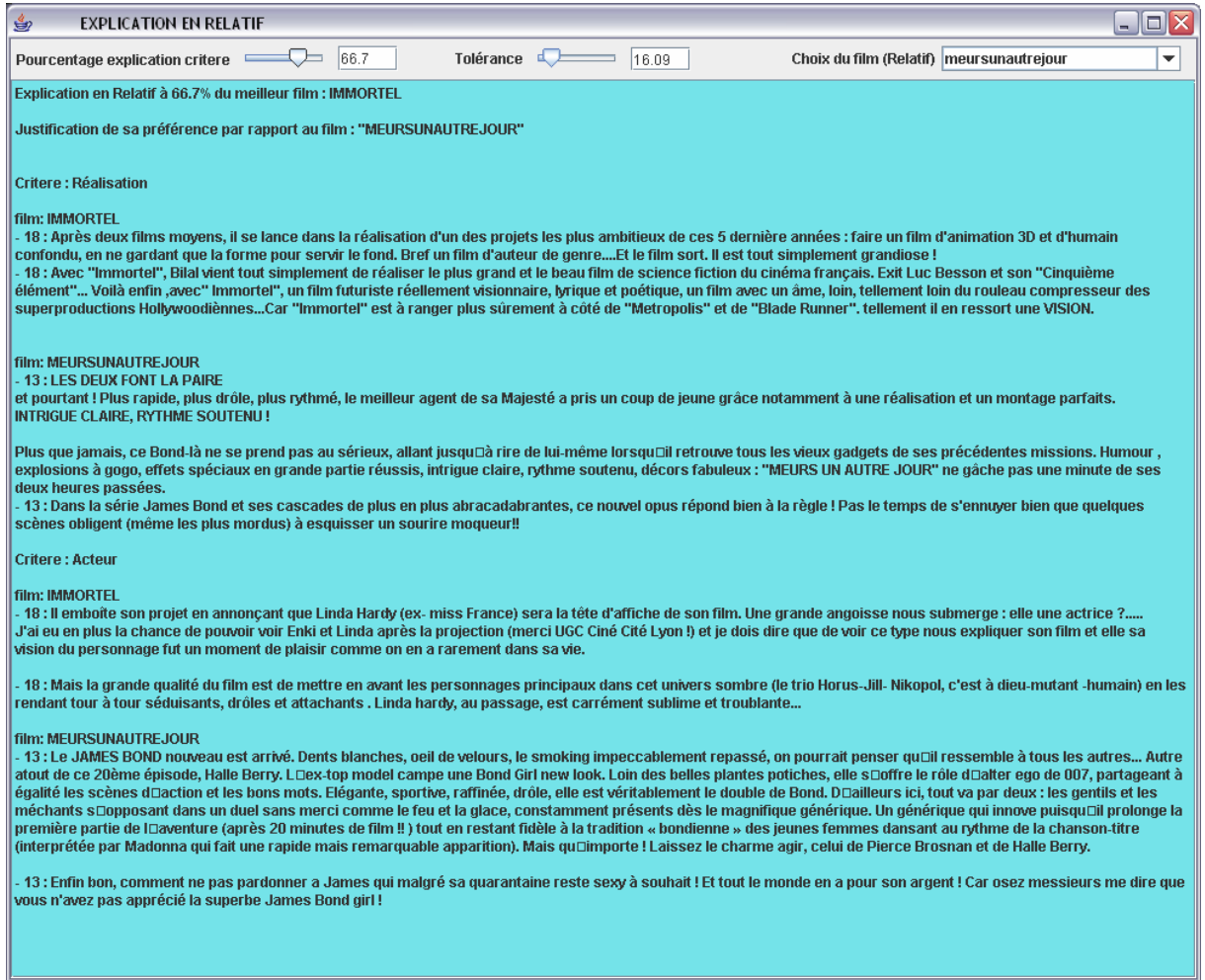


Figure VI.20 : écran d'explication en relatif.

5.2.1. Pourcentage d'explication dans l'explication en relatif

Le principe est exactement le même que pour l'explication absolue, seule l'expression analysée change. Considérons le film f^1 en tête du classement, u^1 son score global à une date t donnée, et f^j le challenger qu'on souhaite lui comparer.

Le paramètre k , « pourcentage d'explication critère », permet de déterminer les critères qui suffisent à expliquer $k\%$ de la différence de score des films f^1 et f^j , soit $k\%$ de $u^1 - u^j$.

Par définition : $MP(f^1) = MP(\mathbf{u}^1) = \sum_{i=1}^{|C|} p_i \cdot u_i^1$.

Après avoir réordonné les termes $p_i \cdot (u_i^1 - u_i^j)$, tels que $\forall i, p_i \cdot (u_i^1 - u_i^j) \geq p_{i+1} \cdot (u_{i+1}^1 - u_{i+1}^j)$, on cherche p_0 tel que :

$$\sum_{i=1}^{p_0 \leq |C|} p_i \cdot (u_i^1 - u_i^j) = k\% \sum_{i=1}^{|C|} p_i \cdot (u_i^1 - u_i^j) = k\% \cdot MP(\mathbf{u}^1 - \mathbf{u}^j) \quad \text{VI-8}$$

Cette équation se lit : « Les p_0 premiers critères expliquent $k\%$ de la différence de score obtenu par f^1 sur f^j ».

Par conséquent, plus k est grand, plus le nombre de critères utilisés pour justifier cette différence est important.

Sur l'exemple de la Figure VI.21, l'interprétation est la suivante : la raison essentielle pour laquelle *Immortel* est préféré à *Meurs un autre jour* est la *Réalisation*.

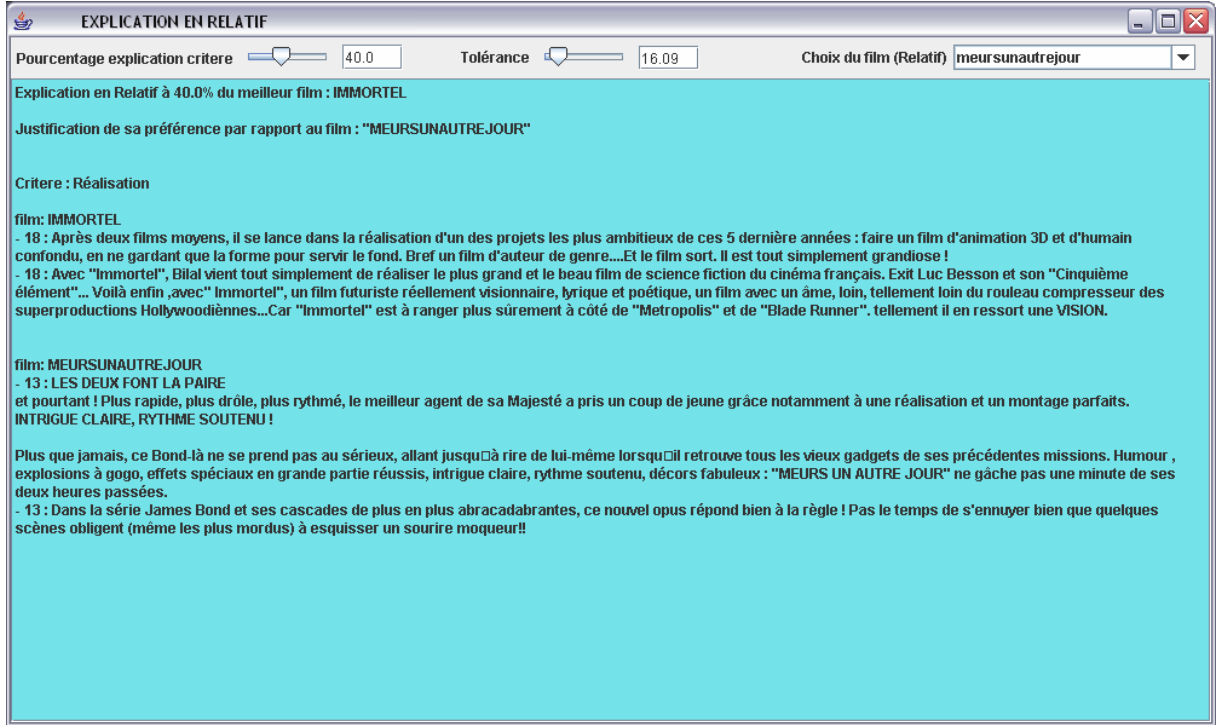


Figure VI.21 : écran d'explication en relatif lorsque $k\%$ est faible

5.2.2. Détail du calcul de la Tolérance dans l'explication en relatif

Le paramètre *Tolérance* ou ε permet de contrôler le nombre de critiques que l'on veut faire valoir pour un critère comme pour l'argumentation absolue ; mais cette fois-ci, il est utilisé pour les CAs de f^j et de f^1 .

Pour justifier d'un bon score partiel u_i^1 , les CAs que l'on va retenir sont telles que :

$$\forall k, \frac{1}{(1+\varepsilon)^2} \cdot \max_k CA_{i,k}^1(\rho) \leq CA_{i,k}^1(\rho) \quad \text{VI-9}$$

Pour justifier d'un mauvais score partiel u_i^j , les CAs que l'on va retenir sont telles que :

$$\forall k, CA_{i,k}^j(\rho) \leq \frac{1}{(1+\varepsilon)^2} \cdot \max_k CA_{i,k}^j(\rho) \quad \text{VI-10}$$

Lorsque la tolérance est très petite ($\varepsilon = 10\%$ par exemple), alors seules les notes excellentes de f^1 seront confrontées aux notes basses ou moyennes de f^j par le SIAD.

Ainsi sur la Figure VI.20, la tolérance est à 16.09 % et nous constatons les affichages suivants :

- Pour le film préféré *Immortel* relativement aux critères *Réalisation* et *Acteur*, le SIAD ne remonte que les CAs dont le score est supérieur ou égal à 18 ;
- Pour le challenger *Meurs un autre jour*, des scores de 13 apparaissent encore comme des résultats mitigés en comparaison de ceux obtenus par *Immortel*.

Si nous augmentons la tolérance pour l'amener à 20%, nous obtenons l'écran de la Figure VI.22 et nous constatons les affichages suivants :

- Pour le film préféré *Immortel* relativement aux critères *Réalisation* et *Acteur*, le SIAD remonte maintenant les CAs dont le score est supérieur ou égal à 13 ;
- Pour *Meurs un autre jour*, le SIAD ne propose aucune CA. En effet, en relaxant la tolérance, le SIAD ne cherche plus que des scores « franchement » mauvais pour le challenger, or *Meurs un autre jour* n'est pas un mauvais film et aucune CA vraiment désobligeante n'est disponible.

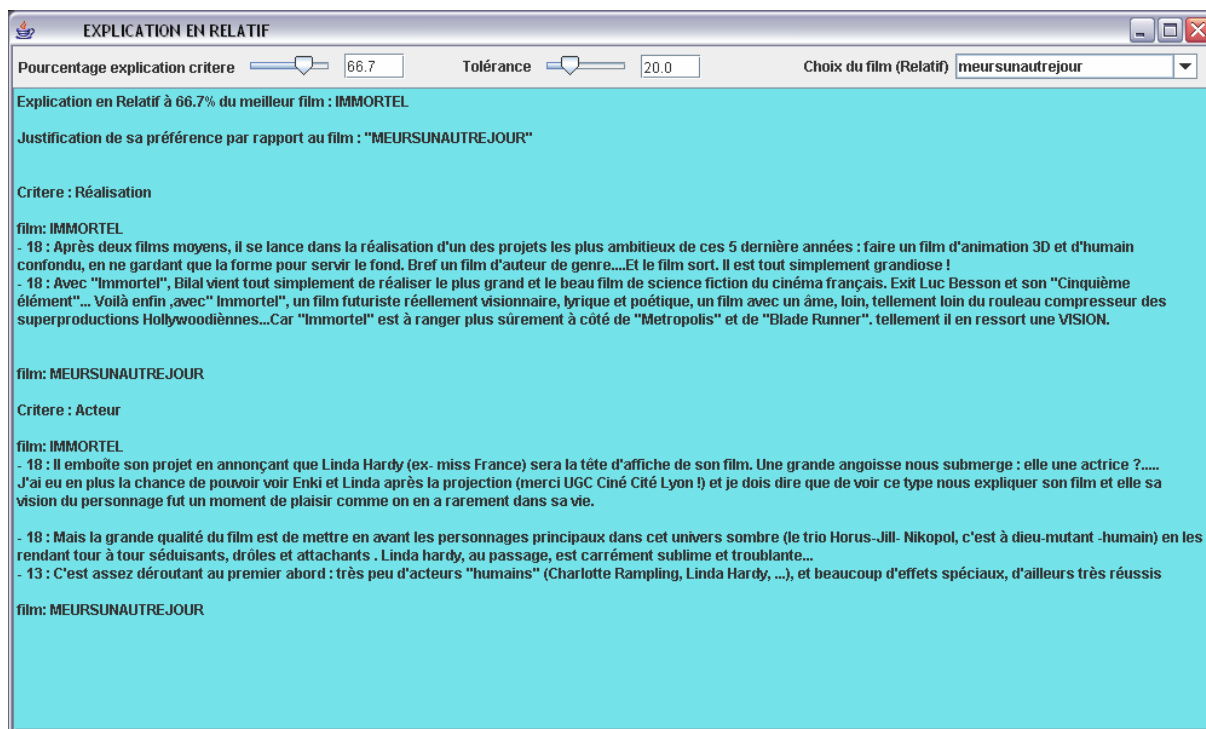


Figure VI.22 : Explication en relatif avec une tolérance plus grande

6. Scénario d'évolution dans le temps et stratégies de contrôle

Nous montrons dans ce paragraphe comment utiliser le SIAD en fonction de la stratégie de contrôle que l'on veut exercer sur la dynamique du processus décisionnel.

Repartons de la grille multicritère de la Figure VI.23.

Soit t_0 : la date du 15/05/2004. A t_0 , la situation est donc la suivante :

Fichier	classifieur	Risque	10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04
lepactedeslousps			8.0	8.0	8.75	8.75	9.916667	10.316667
loveactually			12.166667	12.166667	12.166667	12.722222	12.722222	13.277778
immortel			14.25	14.25	14.25	14.25	14.25	14.666668
meursunautrejour			13.0	13.0	13.0	13.0	13.0	13.0
cube			12.305556	12.305556	12.305556	12.305556	12.305556	14.660354

Figure VI.23 : Grille Film/Date arrêtée au 15 mai 2004

Le SIAD fournit l'écran du risque décisionnel à cette date.

RISQUE DES FILMS								
	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTTEL	14.666668		gain critere		gain critere		gain critere	
CUBE	14.66035	99.94%	+++ (100.0%)	1.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LOVEACTUALLY	13.27777	88.42%	+++ (100.0%)	2.0	++ (0.0%)	0.0	++ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	86.11%	+++ (100.0%)	3.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	63.74%	+++ (100.0%)	55.0	++ (0.0%)	0.0	++ (0.0%)	0.0

Figure VI.24 : Risque décisionnel associé à chaque film candidat

Immortel est le mieux placé.

Nous allons étudier à travers trois scénarii différents, trois types de contrôle qui sont proposés au gérant du vidéo-club via la boucle de régulation de la Figure VI.1. Nous avons baptisé ces scénarii CONFIRMER, EPROUVER et LIBRE (ou BOUCLE OUVERTE).

Pour chacun des types de contrôle, nous allons examiner la situation à l'instant suivant c'est-à-dire le 22/05/2004.

6.1. Scénario CONFIRMER

L'attitude CONFIRMER correspond à la situation suivante. Le classement proposé par le SIAD convient parfaitement au gérant du vidéo-club (il correspond par exemple tout à fait à son intuition). Cependant le niveau de risque est trop élevé. Le gérant souhaite que le SIAD élabore un argumentaire plus fiable et par conséquent que l'actionneur du SIAD (l'ensemble moteur de recherche et classifieurs) acquiert des CAs additionnelles venant conforter ce classement. Autrement dit, le SIAD ne construira que des CAs favorables à *Immortel*, défavorables à ses challengers... Ce comportement peut apparaître surprenant, à la limite du *malhonnête*, mais le gérant ne cherche pas une vérité absolue, il veut juste se construire un argumentaire qui montre que ses choix sont bel et bien fondés sur des critiques que tout à chacun peut consulter... (Cette attitude est beaucoup plus « classique » qu'on ne le pense... on la retrouve aussi bien dans les décisions quotidiennes d'un couple... que dans la décision politique !).

Pour la clarté de l'exemple, ne raisonnons que sur les deux premiers films. Au 15 mai 2004, *Immortel* est certes leader mais suivi de près par *Cube*. Le risque est donc très grand (99.94%). Le SIAD va donc chercher à réduire le niveau de risque tout en confirmant le classement courant. Le calcul du risque nous indique par ailleurs que c'est le critère *Réalisation* qui a la plus forte influence sur le niveau de risque.

La boucle de rétroaction va donc initier la recherche de critiques qui seront à l'origine de la création de CAs positives pour *Immortel* relativement au critère *Réalisation* et des CAs négatives pour *Cube* toujours relativement à *Réalisation*. Cela ne signifie pas qu'il n'existe pas de CAs négatives sur *Immortel* relativement à *Réalisation*, mais ce n'est pas un problème... il faut de tout pour faire un monde ! Bien sûr que l'on peut détester *Immortel*, le gérant du vidéo-club ne remet pas cela en doute, mais il existe néanmoins tout un pendant de la critique qui lui donne raison...

Entre le 15 et le 22 mai 2004, on imagine donc que deux CAs soient ajoutées dans le SGDC :

$$CA_{réalisation,k_{t=22}}^{immortel}(\rho = ++)\text{ et }CA_{réalisation,k_{t=22}}^{cube}(\rho = --).$$

La situation décisionnelle au 22 mai 2004 est alors :

SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS								
Fichier	classifieur	Risque						
		10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04	22.05.04
lepactedesloups		8.0	8.0	8.75	8.75	9.916667	10.316667	10.316667
		6	6	13	13	15	20	20
loveactually		12.166667	12.166667	12.166667	12.722222	12.722222	13.277778	13.277778
		7	7	7	11	11	13	13
immortel		14.25	14.25	14.25	14.25	14.25	14.666668	15.291668
		4	4	4	4	4	9	10
meursunautrejour		13.0	13.0	13.0	13.0	13.0	13.0	13.0
		3	3	3	3	3	7	7
cube		12.305556	12.305556	12.305556	12.305556	12.305556	14.660354	13.535354
		12	12	12	12	12	18	19

RISQUE DES FILMS								
	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTEL	15.291668		gain critere		gain critere		gain critere	
CUBE	13.53535	85.36%	+++ (100.0%)	4.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LOVEACTUALLY	13.27777	83.21%	+++ (100.0%)	5.0	++ (0.0%)	0.0	++ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	80.90%	+++ (100.0%)	6.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	58.54%	+++ (100.0%)	285.0	++ (0.0%)	0.0	++ (0.0%)	0.0

Figure VI.25 : Situation décisionnelle au 22 mai 2004 –Scénario CONFIRMER

Le simple ajout d'une CA positive sur *Réalisation* en faveur du leader *Immortel* au 15 mai et d'une CA négative toujours sur *Réalisation*, mais pour *Cube* a significativement diminué le risque décisionnel. Le SIAD va répéter l'opération dans le temps jusqu'à ce que le niveau de risque devienne acceptable, le gérant du vidéo-club disposera d'un argumentaire de ses choix d'investissement qu'il pourra faire valoir auprès de sa clientèle dans la gazette du vidéo-club. Notons que dans ce type de contrôle, la dynamique du système en boucle fermée est la plus rapide qui soit.

6.2. Scénario EPROUVER

C'est a priori l'usage du SIAD le plus classique (parce qu'il apparaît « plus honnête » !). Donc repartons de la situation décisionnelle au 15 mai 2004. *Immortel* est leader mais suivi de près par *Cube* avec un risque décisionnel très grand (99.94%). Le SIAD indique par ailleurs que c'est le critère *Réalisation* qui a la plus forte influence sur le niveau de risque.

Dans le scénario EPROUVER, l'attitude consiste à tester (éprouver et non plus confirmer) cette dimension critique de la comparaison de *Immortel* et *Cube*. Le contrôle délivre donc cette fois-ci comme unique message à l'actionneur de rechercher des critiques relatives à *Immortel* et *Cube* sur le critère *Réalisation* sans préjuger de ce que les CAs correspondantes seront positives ou négatives pour l'un comme pour l'autre des deux films. On est cette fois-ci dans une logique de vérification d'hypothèse.

Entre le 15 et le 22 mai 2004, le SIAD génère les CAs suivantes : $CA_{réalisation,k_{i=22}}^{immortel} (\rho = ++)$, $CA_{réalisation,k_{i=22}}^{cube} (\rho = --)$ (comme dans le scénario précédent), mais auxquelles viennent s'ajouter 8 autres CAs, $CA_{réalisation,k_{i=22}}^{cube} (\rho = ++), k = 1..8$.

La situation décisionnelle au 22 mai 2004 est alors :



SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS

Fichier	classifieur	Risque						
		10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04	22.05.04
lepactedesloups		8.0 6	8.0 6	8.75 13	8.75 13	9.916667 15	10.316667 20	10.316667 20
loveactually		12.166667 7	12.166667 7	12.166667 7	12.722222 11	12.722222 11	13.277778 13	13.277778 13
immortel		14.25 4	14.25 4	14.25 4	14.25 4	14.25 4	14.666668 9	15.291668 18
meursunautrejour		13.0 3	13.0 3	13.0 3	13.0 3	13.0 3	13.0 7	13.0 7
cube		12.305556 12	12.305556 12	12.305556 12	12.305556 12	12.305556 12	14.660354 18	15.381507 27



RISQUE DES FILMS

	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
CUBE	15.381507		gain critere		gain critere		gain critere	
IMMORTEL	15.29166	99.25%	+++ (100.0%)	1.0	++ (0.0%)	0.0	+ (0.0%)	0.0
LOVEACTUALLY	13.27777	82.46%	+++ (100.0%)	5.0	++ (0.0%)	0.0	+ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	80.15%	+++ (100.0%)	7.0	++ (0.0%)	0.0	+ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	57.79%	+++ (100.0%)	595.0	++ (0.0%)	0.0	+ (0.0%)	0.0

Figure VI.26 : Situation décisionnelle au 22 mai 2004 –Scénario EProuver

Le classement a été changé suite à l'ajout de ces dix nouvelles CAs entre le 15 et le 22 mai. *Cube* est le nouveau leader mais ce changement n'a pas amélioré la stabilité de la situation décisionnelle, le risque reste très élevé.

Au temps suivant, le 29 mai 2004, le SIAD remonte une nouvelle CA, $CA_{réalisation, k_{r=29}}^{immortel}$ ($\rho = ++$).

La nouvelle situation au 29 mai est :



SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS

Fichier	classifieur	Risque							
		10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04	22.05.04	29.05.04
lepactedesloups		8.0 6	8.0 6	8.75 13	8.75 13	9.916667 15	10.316667 20	10.316667 20	10.316667 20
loveactually		12.166667 7	12.166667 7	12.166667 7	12.722222 11	12.722222 11	13.277778 13	13.277778 13	13.277778 13
immortel		14.25 4	14.25 4	14.25 4	14.25 4	14.25 4	14.666668 9	15.291668 18	15.666668 11
meursunautrejour		13.0 3	13.0 3	13.0 3	13.0 3	13.0 3	13.0 7	13.0 7	13.0 7
cube		12.305556 12	12.305556 12	12.305556 12	12.305556 12	12.305556 12	14.660354 18	15.381507 27	15.381507 27



RISQUE DES FILMS

	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTEL	15.666668		gain critere		gain critere		gain critere	
CUBE	15.38150	97.62%	+++ (100.0%)	2.0	++ (0.0%)	0.0	+ (0.0%)	0.0
LOVEACTUALLY	13.27777	80.09%	+++ (100.0%)	7.0	++ (0.0%)	0.0	+ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	77.77%	+++ (100.0%)	10.0	++ (0.0%)	0.0	+ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	54.58%	+++ (96.0%)	très grand	++ (4.0%)	1.0	+ (0.0%)	0.0

Figure VI.27 : Situation décisionnelle au 29 mai 2004 –Scénario EProuver

Immortel est à nouveau en tête, et le risque même s'il reste fort a diminué. La boucle de rétroaction continue à être active jusqu'à ce que la situation devienne décidable, c'est-à-dire que le risque devienne acceptable.

Dans le scénario EPRROUVER, on n'optimise pas la dynamique d'évolution vers une situation décidable, on régule simplement l'incertitude épistémique jusqu'à ce que le niveau de risque soit devenu acceptable.

6.3. Scénario LIBRE ou BOUCLE OUVERTE

Ce cas de figure revient à supprimer la fonction de contrôle par le risque décisionnel de l'information entrante. Il n'a donc aucun intérêt pour notre système si ce n'est que de voir comment la situation décisionnelle du 15 mai aurait évolué en boucle ouverte. Autrement dit, le moteur de recherche soumet toutes les critiques trouvées (le contrôle ne filtre rien puisqu'on est en boucle ouverte) entre le 15 et le 22 mai aux classifieurs du SIAD

Les CAs sont les suivantes :

$CA_{réalisation, k_{t=22}}^{immortel}(\rho = ++)$, $CA_{réalisation, k_{t=22}}^{cube}(\rho = --)$, $CA_{réalisation, k_{t=22}}^{cube}(\rho = ++)$, $k = 1..8$ comme dans le scénario EPROUVER auxquelles s'ajoutent $CA_{acteur, k_{t=22}}^{cube}(\rho = ++)$, $k = 1..4$ et $CA_{acteur, k_{t=22}}^{cube}(\rho = +)$ qui n'ont pas été filtrées cette fois-ci.

La nouvelle situation décisionnelle au 22 mai 2004 est :



SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS

-

□

✖

Fichier	classifieur	Risque						
		10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04	22.05.04
lepactedesloups		8.0 6	8.0 6	8.75 13	8.75 13	9.916667 15	10.316667 20	10.316667 20
loveactually		12.166667 7	12.166667 7	12.166667 7	12.722222 11	12.722222 11	13.277778 13	13.277778 13
immortel		14.25 4	14.25 4	14.25 4	14.25 4	14.25 4	14.666668 9	15.291668 10
meursunautrejour		13.0 3	13.0 3	13.0 3	13.0 3	13.0 3	13.0 7	13.0 7
cube		12.305556 12	12.305556 12	12.305556 12	12.305556 12	12.305556 12	14.660354 18	15.561431 32



RISQUE DES FILMS

-

□

✖

	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
CUBE	15.561431		gain critere		gain critere		gain critere	
IMMORTEL	15.29166	97.75%	+ + + (100.0%)	1.0	+ + (0.0%)	0.0	+ (0.0%)	0.0
LOVEACTUALLY	13.27777	80.96%	+ + + (100.0%)	6.0	+ + (0.0%)	0.0	+ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	78.65%	+ + + (100.0%)	9.0	+ + (0.0%)	0.0	+ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	55.89%	+ + + (98.0%)	très grand	+ + (2.0%)	1.0	+ (0.0%)	0.0

Figure VI.28 : Situation décisionnelle au 22 mai 2004 –Scénario LIBRE ou BO

Cube l'emporte mais le risque décisionnel reste élevé. Le classement de départ a donc été bouleversé.

L'actionneur continue à alimenter le SIAD sans aucun contrôle des classifieurs et au 29 mai 2004, la situation décisionnelle, après ajout de $CA_{réalisation, k_{t=29}}^{immortel}(\rho = ++)$ (comme dans le scénario précédent à la date du 29 mai) et des $CA_{acteur, k_{t=29}}^{immortel}(\rho = ++)$, $k = 1..8$ devient :

SUPERVISION APPLIQUEE A LA RECOMMANDATION DE FILMS

Fichier

classifieur

Risque

	10.04.04	17.04.04	24.04.04	01.05.04	08.05.04	15.05.04	22.05.04	29.05.04
lepactedesloups	8.0 6	8.0 6	8.75 13	8.75 13	9.916667 15	10.316667 20	10.316667 20	10.316667 20
loveactually	12.166667 7	12.166667 7	12.166667 7	12.722222 11	12.722222 11	13.277778 13	13.277778 13	13.277778 13
immortel	14.25 4	14.25 4	14.25 4	14.25 4	14.25 4	14.666668 9	15.291668 10	16.070707 13
meursunaautrejour	13.0 3	13.0 3	13.0 3	13.0 3	13.0 3	13.0 7	13.0 7	13.0 7
cube	12.305556 12	12.305556 12	12.305556 12	12.305556 12	12.305556 12	14.660354 18	15.561431 32	15.561431 32

RISQUE DES FILMS

	MOYENNE	RISQUE	REALISATION	Nb Crit Min	ACTEUR	Nb Crit Min	SCENARIO	Nb Crit Min
IMMORTTEL	16.070707		gain critere		gain critere		gain critere	
CUBE	15.56143	95.75%	+++ (100.0%)	5.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LOVEACTUALLY	13.27777	76.72%	+++ (100.0%)	12.0	++ (0.0%)	0.0	++ (0.0%)	0.0
MEURSUNAUTREJOUR	13.0	74.41%	+++ (100.0%)	22.0	++ (0.0%)	0.0	++ (0.0%)	0.0
LEPACTEDESLOUPS	10.31666	49.53%	+++ (90.0%)	très grand	++ (10.0%)	2.0	++ (0.0%)	0.0

Figure VI.29 : Situation décisionnelle au 29 mai 2004 –Scénario LIBRE ou BO

Le film *Immortel* est repassé premier et le risque a baissé. Les critiques positives à l'égard d'*Immortel* relativement aux critères *Acteur* qui avaient été filtrées par le contrôleur dans le scénario EProuver à la date du 29 mai, font ici pencher la situation en faveur de *Immortel*.

7. Conclusion

Nous avons montré dans ce chapitre la mise en œuvre de notre vision de l'aide à la décision sur une application à fin pédagogique. Cette application a néanmoins permis de mettre en avant l'ensemble des fonctionnalités de notre SIAD. Nous avons insisté sur les notions de contrôle et d'actionneur de notre SIAD en montrant comment le calcul du risque décisionnel permettait de lancer les classifieurs sur des dimensions de l'information bien identifiées par le système.

Ce chapitre a permis de montrer que la démarche et la modélisation que nous avons adoptées pour définir ce que doit être une aide à la décision de groupe est aujourd'hui entièrement instrumentée sur notre plate-forme logicielle, des outils de traitement automatique de la langue naturelle jusqu'aux outils du multicritère. De plus l'application est disponible en ligne sur Internet en mode client/serveur et multi-utilisateurs.

L'application que nous avons choisie pour illustrer ce chapitre a l'avantage de ne présenter aucun problème de compréhension des situations décisionnelles en jeu afin de ne pas nuire à la présentation de notre démarche intégrée. De plus, elle présente un caractère générique qui permet d'envisager de nombreuses transpositions vers des problèmes industriels (management des performances) ou économiques (gestion de projet).

Conclusion générale

Cette thèse, sans prendre parti, aborde le sujet délicat qu'est l'automatisation cognitive. Elle propose la mise en place d'une chaîne informatique complète de chacune des étapes de la décision. Elle traite en particulier de l'automatisation de la phase d'apprentissage en faisant de la connaissance actionnable—la connaissance utile à l'action—une entité informatique manipulable par des algorithmes. L'application des techniques implémentées dans le SIADG à la gestion des bacs d'un vidéo-club ne saurait prêter à polémique. Cependant, cette illustration nous laisse entendre que si l'explosion des technologies de l'information nous ouvre les portes d'un monde de l'information prodigieux et sans limite, c'est aussi l'ouverture à un monde inexplorable et incontrôlable en terme de quantité et de fiabilité des informations délivrées... Qu'est-ce que l'informatique décisionnelle ? Automatiser la décision humaine, manipuler ou contrôler les choix d'un individu ou d'un collectif, susciter un achat via une « e-recommandation »... est-ce là l'incommensurable bénéfice de tout un chacun à disposer du WWW où il veut, quand il le veut ? Cette prouesse technologique a vraisemblablement déplacé le paradigme « le pouvoir, c'est le savoir ». Le savoir s'affiche, se brade sur la toile sans contrôle : l'affabulation peut prendre valeur de lemme, rien n'en garantit la fiabilité.

Pourtant notre approche n'échappe pas à cette fascination de l'inépuisable source d'information. Le modèle qui supporte notre SIADG s'appuie largement sur des traitements automatiques de la connaissance. « Datamining », « text-mining », multicritère et optimisation sont autant de techniques qui viennent se compléter pour élaborer un artefact de décision. Dans une première lecture, c'est en tout cas ce que pourrait laisser penser notre interprétation cybernétique du modèle de H.A. Simon. Cependant, la boucle de rétroaction de notre système a justement pour objet d'analyser la stabilité de la décision, elle met à profit l'abondance des avis et opinions pour discerner quelle en est la tendance résumée. L'incertitude épistémique autour d'un débat public ou d'un enjeu politique est mesurée par le risque décisionnel qui analyse les facteurs discriminants entre les alternatives envisagées.

Plusieurs attitudes dans le contrôle du risque décisionnel peuvent être envisagées : le SIADG peut être utilisé pour valider, vérifier ou infirmer un point de vue. Dans tous les cas, le contrôle exercé sur l'incertitude épistémique n'est pas neutre quant à la dynamique du processus de décision, mais il s'agit davantage d'accélérer l'apprentissage inhérent à une décision que de manipuler ce savoir. Initié par les travaux d'Akharraz [Akharraz, 2004], notre modèle apporte un éclairage plus formel et moins restrictif des liens entre incertitude épistémique, risque décisionnel et stabilité de la décision.

Par ailleurs, notre intention première n'était pas d'aborder la problématique de l'automatisation cognitive. Il y a quelques années, avant que les NTIC ne soient omniprésentes, notre discours ne se serait sans doute jamais égaré sur cette voie. Initialement, nous avons effectivement conçu notre SIADG en tant qu'assistance, support au collectif des décideurs. Nous nous sommes basés sur une typologie des erreurs humaines qui guettent un collectif dans la prise de décision. Cette analyse s'est appuyée largement sur les concepts introduits par les sciences humaines dans une réflexion pluridisciplinaire autour de la décision. Sur la catégorisation de l'erreur humaine dans un processus de décision, nous avons identifié les fonctionnalités attendues d'un SIADG. Le SIADG doit soutenir l'activité cognitive des décideurs tant dans leur apprentissage de la situation que dans les phases d'évaluation, de comparaison et de sélection des alternatives. En particulier, l'Interface Homme/Machine d'un SIADG se doit d'être synthétique, spécifique et claire dans la présentation de l'information, pertinente par rapport à la situation et adaptée aux modes de raisonnement des décideurs. Cette démarche s'inscrit donc bel et bien dans une assistance à la décision plus que dans une automatisation de celle-ci. Il s'agit d'accompagner le décideur dans chacune des phases du processus de la décision, d'instrumenter informatiquement autant que se peut chaque activité cognitive de celui-ci, mais l'homme peut intervenir à chaque étape sur l'information, la représentation de la situation ou le critère de sélection. Au-delà de toute quête idéologique, notre SIADG est avant tout une chaîne de traitement automatique de l'information pour la décision unissant divers outils des mathématiques décisionnelles pour systématiser, fiabiliser et optimiser l'action humaine sur la chaîne de l'information.

En dehors des aspects de formalisation et de généralisation, notre principal apport au modèle cybernétique originel d'Akharraz consiste à en avoir élaboré l'actionneur. Sous les traits d'un moteur de recherche et d'une batterie de classifieurs, l'organe de commande de la boucle de rétroaction prend la forme d'un système d'information : le calcul du risque décisionnel détermine les dimensions stratégiques de l'information guidant ainsi une quête itérative des savoirs utiles à la décision. Le calcul du risque décisionnel délivre le signal de commande utilisé par nos outils de traitement automatique de la langue naturelle (TALN) pour acquérir de nouvelles connaissances et pallier ainsi aux déficits informationnels les plus critiques du problème. Il s'agit d'un apprentissage contrôlé par l'objectif de parvenir à une sélection stable au plus vite. Seuls les éléments qui peuvent avoir un rôle discriminant sont recherchés par notre système d'information.

L'essentiel de ce manuscrit est donc consacré à l'élaboration de cet actionneur. Le chapitre III introduit les concepts fondamentaux de connaissance actionnable (CA) et d'indexation automatique sur lesquels reposent nos modèles et outils de TALN. La notion de connaissance actionnable, généralement utilisée pour définir une connaissance utile à l'action, trouve ici une autre justification : c'est aussi la connaissance manipulée par l'actionneur. Le chapitre IV propose une synthèse rapide des techniques d'apprentissage les plus éprouvées en vue de l'extraction automatique de CAs pour la décision. Le chapitre V décline l'ensemble de ces notions et techniques sur la problématique spécifique d'extraction automatique des

connaissances dans un processus d'évaluation multicritère. Le chapitre VI reprend et illustre le processus informatisé dans sa globalité sur l'exemple du gérant d'un vidéoclub cherchant à optimiser ses investissements en fonction des préférences de sa clientèle.

Notre chaîne de TALN transforme des documents textuels en CAs pour notre SIADG. La première spécificité de notre outil concerne l'élimination des phrases sans apport informationnel au regard des critères de sélection. Ce traitement permet de diminuer la complexité des processus de classification ultérieurs. La seconde innovation consiste à décomposer l'élaboration d'une CA en processus de classification élémentaires : affectation à un critère de sélection et attribution d'une évaluation numérique en rapport avec le jugement de valeur exprimé en langue naturelle. Pour chacun de ces processus de classification élémentaires est calculé un index dédié de la base de connaissance qui permet de contrôler la complexité de l'apprentissage et d'améliorer les performances du classifieur. Un troisième point original consiste à introduire la synonymie pour « densifier » le vocabulaire manipulé par le classifieur. L'affectation de synonymes est ici innovante par la réduction des définitions d'un mot, pour lesquelles nous recherchons ces synonymes, en fonction de la classification concernée. Enfin la dernière particularité de notre chaîne de traitement relève du traitement hiérarchique pour l'attribution d'évaluations numériques à des jugements de valeurs exprimés en langue naturelle, tâche de classification extrêmement complexe car sensible à toutes les subtilités très subjectives de la langue.

Au-delà de cette réflexion technique sur le TALN, le chapitre d'application s'attache à démontrer le fonctionnement de la chaîne de traitement de l'information complète. Notre démarche intégrée pour l'aide à la décision met en œuvre nombre de techniques mathématiques ou informatiques se réclamant usuellement toutes du domaine de la décision et pourtant s'ignorant généralement les unes et les autres parce qu'issues de communautés scientifiques bien distinctes. Ainsi, cohabitent dans notre SIADG, calculs d'optimisation, analyse multicritère, datamining et « text-mining », le tout dans un contexte formel inspiré de l'automatique. C'est l'intégration de toutes ces techniques dans le modèle cybernétique qui permet d'aller du document à la sélection d'une alternative de manière fortement automatisée, sans oublier la conception d'une IHM compatible avec les modes cognitifs des utilisateurs.

Cette thèse s'est donc donné pour objet :

- une formalisation plus aboutie du modèle cybernétique d'Akharraz en proposant en particulier la conception de l'actionneur du modèle ;
- la proposition d'une démarche intégrée, scientifique et technique, et instrumentée pour le traitement de la connaissance utile à la décision ;
- la mise en œuvre d'une plate-forme informatique supportant l'ensemble des concepts et modèles.

La mixité des outils utilisés par notre SIADG donne un caractère atypique et prototypique à celui-ci ! Chacune des étapes doit être maintenant affinée. Par exemple, en ce qui concerne les outils du multicritère, l'algorithme du risque décisionnel nécessite dans le cas général d'avoir recours à des outils d'optimisation non linéaire. Sur le thème de la gestion des incertitudes, la multiplicité des avis devrait être représentée par des distributions probabilistes ou possibilistes, mais pas par un avis agrégé (« moyen ») trop réducteur car on perd toute l'information liée à la diversité des opinions qui a pourtant bien souvent un impact non négligeable sur la sélection finale. L'idée est donc de faire du multicritère non plus sur des

réels, mais sur des distributions [Denguir *et al.*, 2005]. En ce qui concerne les outils du TALN, une application grandeur nature aurait permis de construire des bases d'apprentissage significatives et de construire vraisemblablement des classifieurs performants qui auraient pu être évalués autrement que par validation croisée. La classification hiérarchisée qui permet d'attribuer un score numérique à un jugement de valeur exprimé en langage naturel doit être généralisée pour des échelles d'évaluation plus fines que celles proposées dans l'application. Les classifieurs utilisent des bases d'apprentissage indexées manuellement ; ces processus sont coûteux en terme d'intervention humaine. Nous travaillons donc à concevoir des outils et techniques visant à diminuer le temps consacré à l'indexation manuelle des bases d'apprentissage soit en informatisant les tâches répétitives ou systématiques, soit en substituant à l'homme un traitement algorithmique lorsque c'est possible (filtres de contexte, ontologies de domaine existantes, par exemple). Enfin, la prise en compte des profils utilisateurs est encore une dimension non intégrée dans notre démarche.

Le couplage du SIADG et d'un système d'information automatisé comme nous l'avons proposé ici est facilement transposable à différents secteurs d'activités : il concerne tout processus qui fait intervenir l'évaluation multicritère, la comparaison et la sélection d'alternatives. On peut citer la veille technologique, l'utilisation intelligente de la jurisprudence dans le domaine du droit, la conduite de projets stratégiques [Montmain *et al.*, 2002], la gestion d'appels d'offres, la maintenance préventive de patrimoines, la gestion de portefeuilles de projets, le contrôle de la productivité d'une organisation, le marketing et le benchmarking pour le e-commerce [Denguir-Rekik *et al.*, 2006b], etc. L'approche peut concerner tout processus décisionnel où l'on peut agir dynamiquement sur la phase d'information et qui, in fine, doit être légitimé.

Bibliographie

[Akharraz, 2004]

Akharraz A. *Acceptabilité de la décision et risque décisionnel : Un système explicatif de fusion d'information par l'intégrale de Choquet*. Thèse de Doctorat, Université de Savoie, Grenoble, Nîmes. 28 09 2004, 2004

[Akharraz et al., 2003]

Akharraz A., Montmain J., Troussel F., & Mauris G. *Acceptabilité d'un processus décisionnel dynamique : argumentation et contrôle d'une décision basée sur une agrégation par l'intégrale de Choquet*. in *Proceedings of LFA'2003*, Tours, France. 2003

[Allison, 1971]

Allison D. *The essence of Decision. Explaining the Cuban Missile Crisis*. Boston: Little Brown. 1971

[Argyris et al., 1978]

Argyris C., & Schön D. *Organizational learning : a theory of action perspective*. Addison Wesley., 1978

[Bainbridge, 1991]

Bainbridge L. *Les systèmes experts résoudre-t-ils tous les problèmes des opérateurs. Introductory lecture: Facteurs humains de la fiabilité et de la sécurité des systèmes complexes* (ed. INRS). Vandœuvre (F). 1991

[Baizet, 2004]

Baizet Y. *La Gestion des Connaissances en Conception, Application à la simulation numérique chez Renault - DIEC*. Thèse de Doctorat, Université Joseph Fourier - Grenoble 1, GRENOBLE. 2004

[Balabanovic et al., 1997]

Balabanovic M., & Shoham Y. Content-Based, Collaborative Recommendation. *Communications of the ACM*, 40(3). 1997

[Ballay, 1997]

Ballay J. F. *Capitaliser et transmettre les savoir-faire de l'entreprise* Electricité de France. 1997

[Basu et al., 1998]

Basu C., Hirsh H., & Cohen W. *Recommendation as Classification: Using Social and Content-Based Information in Recommendation*. in *Proceedings of Fifteenth National Conference on Artificial Intelligence*, Madison, Wisconsin, USA. 1998

[Benkhannouche, 1996]

Benkhannouche S. *Aide à la supervision des processus industriels : vers une méthodologie de conception*. Thèse de Doctorat, Université Pierre et Marie Curie, Paris 6. 1996

[Besançon, 2001]

Besançon R. *Intégration de connaissances syntaxiques et sémantiques dans les représentations vectorielles de textes - Application au calcul de similarités sémantiques dans le cadre du modèle DSIR*. Thèse de Doctorat, EPFL, Lausanne, CH. Décembre, 2001

[Besançon et al., 2000]

Besançon R., & Rajman M. Le modèle DSIR : une approche à base de sémantique distributionnelle pour la recherche documentaire. *Revue TAL'2000*, 41(2), pp 1-27. 2000

[Bonjour et al., 2002]

Bonjour E., & Dulmet M. *Articulation entre pilotage des systèmes de compétences et gestion des connaissances*. in *Proceedings of Gestion des Compétences et des Connaissances en Génie Industriel*, Nantes. 12-13 décembre, 2002

[Brooking, 1998]

Brooking A. Corporate Memory : Strategies For Knowledge Management, pp 4-5. 1998.

[Burgess et al., 1995]

Burgess C. J. C., Schölkopf B., & Vapnik V. *Extracting Support Data for a Given Task. in Proceedings of First International Conference on Knowledge Discovery & Data Mining*, AAAI Press, Menlo Park, CA. 1995

[Caropreso et al., 2001]

Caropreso M., Matwin S., & Sebastiani F. Independent evaluation of the usefulness of statistical phrases for automated text categorization. In A. G. Chin (Ed.), *Text Databases and Document Management : Theory and Practice* (ed. pp. 78-102): Idea Group Publishing, Hershey, US. 2001

[Caverni et al., 1990]

Caverni J.-P., Fabre J.-M., & Gonzalez M. *Cognitive biases*. Amsterdam: North Holland. 1990

[Cavner et al., 1994]

Cavner W., & Trenkle J. *N-gram-based text categorization. in Proceedings of SDAIR-94, the 3rd Annual Symposium on Document Analysis and Information Retrieval*, Las Vegas, US. pp. 161-175. 1994

[Cellier, 1990]

Cellier J. M. L'erreur humaine dans le travail, *Les Facteurs Humains de la Fiabilité dans les Systèmes Complexes, ouvrage collectif sous la direction de J. Leplat et G. de Terssac* (ed. OCTARES Entreprises, pp. 193-209). 1990

[Chabert-Ranwez, 2000]

Chabert-Ranwez S. *Composition Automatique de Documents Hypermédia Adaptatifs à partir d'Ontologies et de Requêtes Intentionnelles de l'Utilisateur*. Thèse de Doctorat en informatique, Université Montpellier II, Montpellier. 2000

[Chauché, 1990]

Chauché J. Détermination sémantique en analyse structurale : une expérience basée sur une définition de distance. *TA Information*, 31(1), pp 17-24. 1990

[Cholewa, 1985]

Cholewa W. Aggregation of fuzzy opinions an axiomatic approach. *Fuzzy sets and systems*(17), pp 249-258. 1985

[Christianini et al., 2000]

Christianini N., & Shawe-Taylor J. *An introduction to support vector machines*. Cambridge UK: Cambridge University Press. 2000

[Clavien et al., 2003]

Clavien L., & Bétrancourt M. *Animations multimédia : quels dispositifs pour réduire la charge cognitive ? in Proceedings of Deuxièmes Journées d'étude en Psychologie ergonomique*, Boulogne-Billancourt 2-3 octobre 2003, 2003

[Cover et al., 1991]

Cover, & Thomas. *Elements of Information Theory*. John Wiley. 1991

[Crampes et al., 2000]

Crampes M., Ranwez S., & Plantié M. *Ontology-Supported and Ontology-Driven Conceptual Navigation on the World Wide Web. in Proceedings of HyperText 2000, ACM HT2000, San Antonio, Texas USA*. pp. pp 191-199. June, 2000

[Crampes et al., 2002]

Crampes M., Ranwez S., Plantié M., & Vaudry C. Analyse de la pertinence d'une indexation portée par XML. *Editions Hermes Documentation numérique*. 2002

[Crawford et al., 1991]

Crawford S., Fung R., Appelbaum L., & Tong R. *Classification trees for information retrieval*. in *Proceedings of 8th International workshop Machine Learning*, Northwestern university, Illinois. 1991

[De_Keyser, 1990b]

De_Keyser V. Fiabilité humaine et la gestion du temps dans les systèmes complexes, *Les Facteurs Humains de la Fiabilité dans les Systèmes Complexes, ouvrage collectif sous la direction de J. Leplat et G. de Terssac* (ed. OCTARES Entreprises, pp. 85-108). 1990b

[De_Keyser et al., 1990a]

De_Keyser V., & Woods D. D. Fixation errors in complex systems. In A. G. Colombo, et al. (Eds.), *Advanced Systems Reliability Modelling* (ed.). Dordrecht: Kluwer Academic Publisher. 1990a

[Deerwester et al., 1990]

Deerwester S., Dumais S., Landauer T., Furnas G., & Harshman R. Indexing by latent semantic analysis. *Journal of the American Society for Information Science*, 41(6), pp 391-407. 1990

[Denguir-Rekik et al., 2006a]

Denguir-Rekik A., Mauris G., & Montmain J. Propagation of Uncertainty by the Possibility Theory in Choquet Integral based Decision making for an application to an E-commerce Website Choice Support. *IEEE Transactions on Instrumentation and Measurement*. 2006a

[Denguir-Rekik et al., 2006b]

Denguir-Rekik A., Montmain J., & Mauris G. A fuzzy-valued Choquet-integral-based multi-criteria decision-making support for marketing and benchmarking activities in e-commerce organizations. in *Proceedings of MCDM'2006*, Chania, Greece. 2006b

[Denguir et al., 2005]

Denguir A., Mauris G., & Montmain J. *Handling uncertainty in a multi-criteria decision-making support system: Application to the E-recommendation website*. Eusflat-LFA 2005, Joint 4th Eusflat & 11th LFA Conference, Barcelona, Spain. 2005

[Denis et al., 2000]

Denis F., & Gilleron R. *Apprentissage à partir d'exemples - Notes de cours*. <http://www.grappa.univ-lille3.fr/polys/apprentissage/>. Université de Lille. 2000

[Desjeux et al., 1988]

Desjeux D., Orhant I., & Taponier. *L'édition en sciences humaines, la mise en scène des sciences de l'homme et de la société* (ed. L'Harmattan). 1988

[Després, 1991]

Després F. M. *Automatisation des systèmes de production, du besoin à l'utilisation*. (ed. KIRK). 1991

[Doise et al., 1992]

Doise W., & Moscovi S. *Dissensions et consensus* (ed. PUF). 1992

[Dubois, 1983]

Dubois D. *Modèles mathématiques de l'imprécis et de l'incertain en vue d'applications aux techniques d'aide à la décision*. Thèse de Doctorat, Institut National Polytechnique, Grenoble. 1983

[Dubois et al., 1984]

Dubois D., & Prade H. Criteria aggregation and ranking of alternatives in the framework of fuzzy set theory, *Fuzzy Sets and Decision Analysis*, (ed. H-J. Zimmermann, L.A. Zadeh and B. Gaines, Vol. 5, pp. 209-240): TIMS Studies in the Management Sciences. 1984

[Dubois et al., 1985]

Dubois D., & Prade H. A review of fuzzy set aggregation connectives. *Information Sciences*(36), pp 85-121. 1985

[Ermine J.L. et al., 96]

Ermine J.L., Chaillot M., Bignon P., B. C., & D. M. MKSM : Méthode pour la gestion des connaissances, Ingénierie des systèmes d'information, *AFCET-Hermès*, 4, pp 541-575. 1996, 96

[Ermine, 1996]

Ermine J. L. *Les systèmes de connaissances* (ed. Hermès). Paris. 1996

[Gilbert, 1992]

Gilbert C. *Le pouvoir en situation extrême. Catastrophes et politiques* (ed. L'Harmattan). 1992

[Gilleron et al., 2000]

Gilleron R., & Tommasi M. *Découverte de connaissances à partir de données - notes de cours*. <http://www.grappa.univ-lille3.fr/polys/fouille/>: Université de Lille. 2000

[Grabisch et al., 1995]

Grabisch M., Nguyen H. T., & Walker E. A. (Ed.). (Eds.). *Fundamentals of Uncertainty Calculi, with Applications to Fuzzy Inference* (Vol. 8): Kluwer Academic. 1995

[Grabisch et al., 1998]

Grabisch M., Orłowski S. A., & Yager R. R. Fuzzy aggregation of numerical preferences. In R. Slowinski (Ed.), *Fuzzy Sets in Decision Analysis, Operations Research and Statistics*. (ed.): Kluwer Academic. 1998

[Grabisch et al., 1996]

Grabisch M., & Roubens M. The application of fuzzy integrals in multicriteria decision making. *European Journal of Operational Research*(89), pp 445-456. 1996

[Grabisch et al., 2000]

Grabisch M., & Roubens M. Application of the Choquet Integral in Multicriteria Decision Making, *Fuzzy Measures and Integrals : Theory and Applications* (ed.): (Grabisch, Murofushi and Sugeno), Physica-Verlag. 2000

[Grundstein, 1995]

Grundstein M. *La capitalisation des connaissances de l'entreprise, système de production de connaissances*. in *Proceedings of Colloque L'Entreprise Apprenante et les Sciences de la Complexité*, Aix en Provence. 22-24 Mai, 1995, 1995

[Grundstein, 2000]

Grundstein M. *Repérer et mettre en valeur les connaissances cruciales pour l'entreprise*. in *Proceedings of Actes du 10ème Congrès International de l'AFAPV*, Paris. novembre, 2000

[Han et al., 2001]

Han J., & Kamber M. *Data Mining : Concepts and Techniques*. San Francisco, California: Morgan Kaufmann Publishers. 2001

[Hartigan, 1975]

Hartigan J. A. *Clustering Algorithms*. New York: John Wiley & Sons., 1975

[Hearst, 1994]

Hearst M. A. *Multi-paragraph Segmentation of Expository Text*. in *Proceedings of 32nd Annual meeting of the Association for Computational Linguistics*, Las Cruces. 1994

[Hearst, 1997]

Hearst M. A. Text-tiling : segmenting text into multi-paragraph subtopic passages. *Computational Linguistics*, pp 59-66. 1997

[Herlocker et al., 2000]

Herlocker J. L., Konstan J. A., & Riedl J. *Explaining collaborative filtering recommendations*. in *Proceedings of ACM 2000 Conference on Computer Supported Cooperative Work*, Philadelphia, USA. 2000

[Hill et al., 1995]

Hill W., Stead L., Rosenstein M., & Furnas G. W. *Recommending and Evaluating Choices in a Virtual Community of Use*. in *Proceedings of ACM CHI'95 Conference on human factors in computing systems*, Denver, CO., USA. pp. 194-201. 1995

[Hogart, 1987]

Hogart R. *Judgement and choice : the psychology of decision*. Chichester: Wiley. 1987

[Iksal, 2002]

Iksal S. *Spécification Déclarative et Composition Sémantique pour les Documents Virtuels Personnalisables*. Thèse de Doctorat en informatique, ENST Bretagne, Nantes. 2002

[J-L. Koning, 1990]

J-L. Koning. *Un mécanisme de gestion de règles de décision antagonistes pour les systèmes à base de connaissances*. Thèse de Doctorat, Université Paul Sabatier, Toulouse. 1990

[Jaillet, 2005]

Jaillet S. *Catégorisation automatique de documents textuels : d'une représentation basée sur les concepts aux motifs séquentiels*. Thèse de Doctorat en Informatique, Université des Sciences et Techniques du Languedoc, Montpellier. 7 mars 2005, 2005

[Jaillet et al., 2004]

Jaillet S., Teisseire M., Dray G., & Plantié M. *Comparing Concept-Based and Statistical Representations for Textual Categorization*. in *Proceedings of IPMU'04: 10th Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems*, Perugia (Italy). pp. 91-98. 2004

[Jain et al., 1999]

Jain A. K., Murty M. N., & Flynn P. J. Data clustering : A review. *ACM Comput. Surv.*, 31(3), pp 264 - 323. 1999

[Jarrosson, 1994]

Jarrosson B. *Décider ou ne pas décider ?* (ed. Maxima, Laurent du Mesnil). 1994

[Joachims, 1998]

Joachims T. *Text Categorisation with Support Vector Machines : Learning with Many Relevant Features*. in *Proceedings of ECML*. 1998

[Kacprzyk, 1987]

Kacprzyk J. Towards "human-consistent" decision support systems through commonsense knowledge-based decision making and control models: a fuzzy approach. *Computers and Artificial Intelligence*, 6(2), pp 97-122. 1987

[Kahneman et al., 1982]

Kahneman D., Slovic P., & Tversky A. *Judgement under uncertainty : heuristics and biases*. Cambridge University Press., 1982

[Koenig, 1990]

Koenig G. *Management stratégique : vision, manœuvres, tactiques* (ed. Vol. chap.1 à 3.): Nathan. 1990

[Lagadec, 1988]

Lagadec P. *Etats d'urgence. défaillances technologiques et déstabilisation sociale* (ed. Seuil). 1988

[Larousse, 1992]

Larousse. *Thésaurus Larousse - des idées aux mots - des mots aux idées*. 2-03-320-148-1: Larousse. 1992

[Le_Moigne, 1998]

Le_Moigne J.-L. *Connaissance actionnable et action intelligente*. Grand Atelier MCX, au Futuroscope, Poitiers, France. 1998

- [Lévine *et al.*, 1989]
Lévine P., & Pomerol J. C. *Systèmes interactifs d'aide à la décision et systèmes experts*. Hermès. 1989
- [Lewis, 1992]
Lewis D. D. *Representation and learning in information retrieval*. PhD Thesis, Amherst, MA, USA. 1992
- [Manning *et al.*, 1999]
Manning C., & Schütze H. *Foundations of Statistical Natural Language Processing*. Cambridge, Massachusetts: The MIT Press. 1999
- [March, 1988]
March J. *Decisions and Organization*. New-York: Blackwell. 1988
- [March, 1991]
March J. *Décisions et organisations* (ed. Editions d'Organisation.). 1991
- [Marcu, 1997]
Marcu D. *The Rhetorical Parsing, Summarization, and Generation of Natural Language Texts*. PhD Thesis, University of Toronto, Toronto. 1997
- [Mavrovouniotis *et al.*, 1988]
Mavrovouniotis M., & Stephanopoulos G. Formal order-of-magnitude reasoning in process engineering. *Computer and Chemical Engineering*, 12(9/10), pp 867-880. 1988
- [Mayère, 1997]
Mayère A. *Capitalisation des connaissances et nouveau modèle industriel*. (ed. l'Harmattan). Paris: sous la direction de Monnoyer M.-C., . L'entreprise et l'outil informationnel, sous la direction de Monnoyer M.-C. 1997
- [McNee *et al.*, 2003]
McNee S. M., Lam S. K., Guetzlaff C., Konstan J. A., & Riedl J. *Confidence Displays and Training in Recommender Systems*. in *Proceedings of INTERACT '03 IFIP TC13 International Conference on Human-Computer Interaction*. pp. 176-183. 2003
- [Minel *et al.*, 2001]
Minel J.-L., Desclés J.-P., Cartier E., Crispino G., Hazez S. B., & Jackiewicz A. Résumé automatique par filtrage sémantique d'informations dans des textes. *Revue Technique et Science Informatiques*(3). 2001
- [Montero, 1989]
Montero F. J. D. J. *Weighted aggregation and single-peaked intensities*. in *Proceedings of Workshop on Aggregation and best choices of imprecise opinions*, Bruxelles, Belgium. 1989
- [Montmain *et al.*, 2002]
Montmain J., Akharraz A., & Mauris G. *Knowledge management as a support for collective decision-making and argumentation processes*. in *Proceedings of IPMU'2002*, 9th International Conference on Information processing and Management of uncertainty in Knowledge-Based System, Annecy, France. 2002
- [Montmain *et al.*, 2005]
Montmain J., Denguir-Rekik A., & Mauris G. *How deriving benefits from expert advices to make the right choice in multi-criteria decisions based on the Choquet integral?* . European workshop on the Use of Expert Judgement in Decision-Making, Aix-en-Provence, France. 2005
- [Montmain *et al.*, 2006]
Montmain J., Plantié M., & Akharraz A. *Gestion des connaissances et analyse multicritères pour un système interactif d'aide à la décision en organisation* (ed. Cépaduès). Revue des Nouvelles Technologies de l'Information. 2006

[Mukherjee et al., 2001]

Mukherjee R., Dutta P. S., & Sen S. *Movies2go – a new approach to online movie recommendations. Working notes. in Proceedings of IJCAI 2001 workshop on "Intelligent Techniques for Web Personalisation"*. pp. 10-19. 2001

[Musselin, 1990]

Musselin C. *Quelle stratégie de recherche pour les anarchies organisées. in Proceedings of séminaire CONDOR*. pp. 151-168. 1990

[Nakache et al., 2005]

Nakache D., Metais E., & Timsit J. F. *Evaluation and NLP. in Proceedings of DEXA, Database and Expert Systems Applications, Copenhagen*. pp. 626-632. 22-26 august 2005, 2005

[Neumann et al., 1944]

Neumann V., Morgenstern J., & Morgenstern O. *Theory of games and economic behavior*. Princeton N.J.: Princeton Univ. 1944

[Newell et al., 1972]

Newell A., & H.A. Simon. *Human problem solving*. Englewood Cliffs NJ USA: Prentice-Hall. 1972

[Nicolet et al., 1985]

Nicolet J. L., & Celier J. *La fiabilité humaine dans l'entreprise* (ed. Masson). Paris. 1985

[Nonaka et al., 1996]

Nonaka I., & Takeuchi H. *The Knowledge Creating Company* New York, Oxford Univ. Press. 1996

[Penalva, 2000]

Penalva J.-M. *Connaissances actionnables et intelligence collective*. Ingénierie système et NTIC (Nimestic'2000), Nîmes, France. 2000

[Penalva et al., 2002]

Penalva J.-M., & Montmain J. *Travail collectif et intelligence collective : les référentiels de connaissances. in Proceedings of IPMU'2002, Annecy, France*. 2002

[Perny et al., 2001]

Perny P., & Zucker J.-D. Preference based search and machine learning for collaborative filtering : the "film-conseil" movie recommender system. *Information, Interaction, Intelligence, 1(1)*, pp pp. 9-48. 2001

[Plantié, 2000]

Plantié M. *Segmentation dynamique de programmes de télévision basée sur l'auto-composition d'agents et le modèle AGR de Madkit* (Mémoire de DEA Informatique lgi2p_rap_2000). Montpellier: Université de Montpellier II. 5 juillet, 2000

[Plantié et al., 2002]

Plantié M., Akharraz A., & Montmain J. *De la gestion des connaissances à la décision multicritère et l'argumentation collectives, 1ère Journées d'Etudes sur les Systèmes d'Information pour l'Aide à la Décision en Ingénierie Système. in Proceedings of JESIADIS, ENSIETA, Brest, France*. 28 novembre, 2002

[Plantié et al., 2005a]

Plantié M., Dray G., Montmain J., Meimouni A., & Poncelet P. *DEFI DEFT05 : une approche par classifieur de Bayes. in Proceedings of Atelier DEFT'05, Congrès TALN 2005, Dourdan 6-10 Juin, Dourdan, France*. 6-10 Juin, 2005a

[Plantié et al., 2005b]

Plantié M., Montmain J., & Dray G. *Movies Recommender System : Automation of the Information and Evaluation Phases in a Multi-criteria Decision-Making Process. in Proceedings of DEXA 2005 : 6th International Conference on Database and Expert Systems Applications, Copenhagen, Denmark*. pp. 634-644. 22-26 august 2005, 2005b

- [Pomerol *et al.*, 1993]
Pomerol J.-C., & Barba-Romero S. *Choix multicritère dans l'entreprise*. Hermès. 1993
- [Porter, 1980]
Porter M. F. An algorithm for suffix stripping. *Program*, pp pp 130-137. 1980
- [Prax, 2000]
Prax J. Y. *Le guide du Knowledge management : concepts et pratiques du management de la connaissance*. Edition Dunod. 2000
- [Quinlan, 1993]
Quinlan J. R. *C4.5: Programs for Machine Learning*. (ed. Morgan Kaufmann). San Mateo (CA US) Morgan Kaufmann. Morgan Kaufmann series in machine learning. 1993
- [R. Yager, 1979]
R. Yager. Possibilistic decision-making. *IEEE Trans. Systems Man Cybernetics*(9), pp 177-200. 1979
- [Rasmussen *et al.*, 1981]
Rasmussen J., Pedersen O. M., Carnino A., Griffon M., G. Mancini, & Gagnolet P. *Classification system for reporting events involving human malfunctions*. EUR 7444 EN Luxembourg: Commission of the European Communities. 1981
- [Reason, 1993]
Reason J. *L'erreur humaine* (ed. PUF). Traduit_de anglais, par J.-M. Hoc, 1993
- [Rijsbergen, 1979]
Rijsbergen C. J. V. *Information Retrieval* (ed. 2e). Butterworths, London. 1979
- [Rouet *et al.*, 1995]
Rouet J.-F., & Tricot A. Recherche d'informations dans les systèmes hypertextes : des représentations de la tâche à un modèle de l'activité cognitive, *Sciences et techniques éducatives* (ed. Vol. 3/1995): Herrès. 1995
- [Roy, 1985]
Roy B. *Méthodologie multicritère d'aide à la décision*. Paris: Economica. 1985
- [Roy *et al.*, 1993]
Roy B., & Bouysou D. *Aide Multicritère à la décision : Méthodes et Cas*. Economica., 1993
- [Salton, 1971]
Salton G. The SMART Retrieval System. *Experiments in Automatic Document Processing*. 1971
- [Salton, 1981]
Salton G. A Blueprint for Automatic Indexing. *SIGIR Forum*, 16,(2). 1981
- [Salton, 1983]
Salton G. *Introduction to modern information retrieval*. MCGILL M. J., 1983
- [Salton *et al.*, 1988]
Salton G., & Buckley C. Term Weighting Approaches in Automatic Text Retrieval. *Information Processing and Management*, 24(5), pp 513-523. 1988
- [Salton *et al.*, 1975]
Salton G., Yang C. S., & Yu C. T. A theory of term importance in automatic text analysis. *Journal of the American Society for Information Science*, 26(1), pp 33-44. 1975
- [Savage, 1972]
Savage L. J. *The foundation of statistics*. New-York.: Dover. 1972

- [Schmid, 1994]
Schmid H. *Probabilistic part-of-speech tagging using decision trees. in Proceedings of International Conference on New Methods in Language Processing*, Manchester, UK. 1994
- [Sfez, 1992]
Sfez L. *Critique de la décision (1972)* (ed. 4^e). Presses de la Fondation des sciences politiques édition originale:1973). 1992
- [Shannon, 1948]
Shannon C. A mathematical theory of communication. *Bell Systems Technical Journal*(27), pp 379-423 & 623-656. 1948
- [Simon, 1947]
Simon H. A. *Administrative Behavior: A Study of Decision-Making Processes in Administrative Organizations* (ed. 4th). The Free Press. Human Problem Solving, jointly with Allen Newell. 1947
- [Simon, 1980]
Simon H. A. *Le nouveau management. La décision par les ordinateurs*. Paris: Économica. 1980
- [Simon, 1983]
Simon H. A. *Administration et processus de décision (trad. de Administrative Behavior, 1947)*. Paris.: Économica. 1983
- [Simon, 1991]
Simon H. A. Bounded rationality and Organizational Learning. *Organization Science*(2), pp 125-139. , 1991
- [Simon, 1997]
Simon H. A. *Models of bounded rationality*. Cambridge Massachusetts USA: MIT Press. 1997
- [Slowinski, 1998]
Slowinski R. *Fuzzy sets in decision analysis, operations research and statistics. The handbook of fuzzy set*. Kluwer Academic. The handbook of fuzzy set. 1998
- [Stratégor, 1988]
Stratégor. Stratégie, structure, décision, identité, politique générale d'entreprise, *STRATEGOR* (ed. Vol. chap.13 à 16.): Interéditions. 1988
- [Synapse, 2001]
Synapse. Synapse Analyser: Synapse. "Synapse Analyser." <http://synapse-fr.com>. , 2001
- [Terveen et al., 2001]
Terveen L., & Hill W. Beyond recommender systems : helping people help each other. In e. J. Carroll (Ed.), *HCI in the milenium* (ed. pp. 1-21.): Addison-Wesley. 2001
- [Thoenig, 1987]
Thoenig J.-C. *L'ère des technocrates* (ed. L'Harmattan). Paris. 1987
- [Tsuchiya, 1993]
Tsuchiya S. *Improving Knowledge Creation Ability through Organizational Learning. in Proceedings of ISMICK'93, nternational Symposium on the Management of Industrial and Corporate Knowledge*, Compiègne UTC. 1993
- [Tsuchiya, 1995]
Tsuchiya S. *Commensurability, a key concept of business re-engineering. in Proceedings of 3rd international symposium of the management of information and corporate knowledge*, institut national pour l'intelligence artificielle, Compiègne. 1995
- [Turban, 1993]
Turban E. *Decision Support and Expert Systems*. New York: Macmillan. 1993

[Vacher, 2000]

Vacher B. Connaissance de l'entreprise et de la gestion de l'information. In WEKA (Ed.), *Techniques Documentaires chap. 1* (ed. pp. 146). 2000

[Vapnik, 1982]

Vapnik. Estimation of Dependencies Based on Empirical Data (ed. Springer, pp. 54-55). 1982

[Villemeur, 1988]

Villemeur A. *Sûreté de fonctionnement des systèmes industriels, Fiabilité - Facteurs humains - Informatisation* (ed. Eyrolles). Collection de la Direction des Etudes et Recherches d'Electricité de France. 1988

[Vinck, 1997]

Vinck D. La connaissance : ses objets et ses institutions, *Connaissance et savoir-faire en entreprise, intégration et capitalisation. J.M.Fouet* (ed. pp. chapitre 3): Hermès Science Publications. 1997

[Waliser, 1977]

Waliser B. *Systèmes et modèles. Introduction critique à l'analyse de systèmes* (ed. Seuil). 1977

[Wang et al., 2003]

Wang Y., Hodges J., & Tang B. Classification of Web Documents using a Naive Bayes Method. *IEEE*, pp 560-564. 2003

[Weka Project, 2002-2005]

Weka Project U. O. W. Weka: University of waikato. 2002-2005

[Witten et al., 2005]

Witten I. H., & Frank E. *Data Mining, Practical Machine Learning Tools and Techniques*. Elsevier. data management systems. 2005

[Yager, 1988]

Yager R. On ordered weighted averaging aggregation operators in multicriteria decision-making. *IEEE Transaction Systems, Man and Cybernetics*, 18, pp 183-190. 1988

[Yang et al., 1997]

Yang Y., & Petersen J. O. *A comparative Study on Feature Selection in Text Categorisation. in Proceedings of Fourteenth International Conference on Machine Learning, ICML'97*. 1997

[Zadeh, 1983]

Zadeh L. A computational approach to fuzzy quantifiers in natural languages. *Computers and Mathematics with Applications*, 9, pp 149-184. 1983

Annexes

1. Calcul d'un modèle de classification par les arbres de décision

Un arbre de décision est une représentation graphique d'une procédure de classification. Les nœuds internes de l'arbre sont des tests sur les champs ou attributs, les feuilles sont les classes. Lorsque les tests sont binaires, le fils gauche correspond à une réponse positive au test et le fils droit à une réponse négative. Pour classer un enregistrement, il suffit de descendre dans l'arbre selon les réponses aux différents tests pour l'élément considéré. Voici un exemple d'arbre de décision illustrant notre procédure de classification.

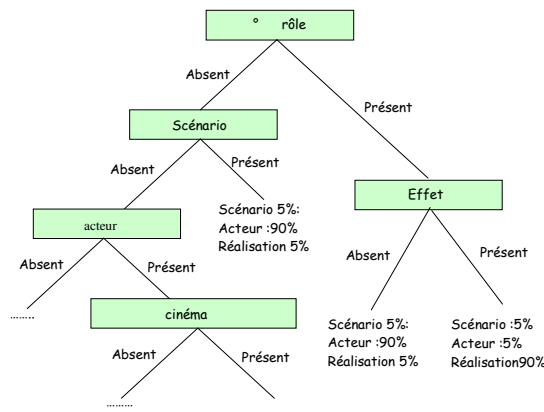


Figure VI.30 : exemple d'arbre de décision

Nous utilisons la technique d'apprentissage pour le calcul de l'arbre de décision : en entrée, on dispose d'un échantillon de m enregistrements classés par des experts et l'algorithme d'apprentissage fournit en sortie un arbre de décision permettant de classer chacun des échantillons. Chaque nœud de l'arbre correspond à un test à effectuer sur l'ensemble à classer. Pour choisir le test, on utilise des fonctions qui mesurent le "degré de mélange" des différentes classes. Pour les problèmes à deux classes, il existe principalement deux méthodes d'apprentissage pour l'élaboration d'arbres de décision :

- CART de [Crawford *et al.*, 1991] dans sa version ne considérant que des tests binaires ;
- C5, version la plus récente de l'algorithme ID3 développé par R. Quinlan [Quinlan, 1993].

Ces deux techniques bien connues sont décrites abondamment dans la littérature.

1.1. Modèles de classification et arbres décisionnels

Considérons un ensemble d'enregistrements. Chaque enregistrement a la même structure et est constitué des paires attribut/valeur. L'un de ces attributs est la classe de l'enregistrement. Le problème est de déterminer un arbre de décision qui sur la valeur des attributs autres que la classe prédise correctement la valeur de l'attribut *classe*. Généralement l'attribut *classe* prend

seulement les valeurs (vrai, faux) ou quelque chose d'équivalent. Dans tous les cas, l'une de ses valeurs est significative d'une erreur.

Les idées de base derrière l'algorithme ID3 (C4.5 et C5.0 n'en sont que des extensions ; C4.5 est une extension de ID3 qui prend en compte les valeurs manquantes, les attributs à valeurs continues...) sont les suivantes :

- Dans un arbre de décision, chaque nœud correspond à un attribut différent de la classe et chaque arc à une valeur possible de cet attribut. Une feuille de l'arbre spécifie la valeur attendue de la classe pour les enregistrements décrits par le chemin menant de la racine à la feuille [c'est la définition de l'arbre de décision] ;
- Dans un arbre de décision, à chaque nœud devrait être associé l'attribut autre que la classe le plus *informationnel* parmi les attributs non encore considérés dans le chemin depuis la racine [ceci établit ce qu'est un *bon* arbre de décision] ;
- L'entropie pour R. Quinlan ou une autre fonction est utilisée pour mesurer la quantité d'information qu'apporte un nœud [ceci correspond à ce que l'on entend par *bon*. Cette notion est introduite en Théorie de l'information par Shannon].

1.2. Pouvoir de décision et quantité d'information

S'il y a n messages équiprobables alors la probabilité d'occurrence p de chacun d'eux est $1/n$ et l'information portée par le message est $-\log(p) = \log(n)$ (dans ce qui suit tous les logarithmes sont en base 2). Ainsi, s'il y a 16 messages, alors $\log(16) = 4$ et 4 bits seront nécessaires à l'identification de ce message.

En général, si l'on se donne une distribution de probabilité $P = (p_1, p_2, \dots, p_n)$ alors l'information portée par cette distribution, également appelée entropie P est :

$$I(P) = -\sum_{i=1}^n p_i \cdot \log p_i \quad \text{VI-11}$$

Par exemple, si P est (0.5, 0.5) alors la quantité d'information $I(P)$ vaut 1 et si P est (0.67, 0.33) alors $I(P)$ vaut 0.92, si $P = (1, 0)$ alors $I(P)$ vaut 0 : on peut noter que plus la distribution est uniforme et plus la quantité d'information est grande.

Si un ensemble T d'enregistrements est partitionné en classes disjointes $C_1, C_2, \dots, C_k$ sur la base de la valeur de l'attribut classe, alors l'information nécessaire pour identifier la classe d'un élément de T est $\text{Info}(T) = I(P)$, où P est la distribution de probabilité de la partition ($C_1, C_2, \dots, C_k$) :

$$P = (|C_1|/|T|, |C_2|/|T|, \dots, |C_k|/|T|) \quad \text{VI-12}$$

Si l'on partitionne d'abord T sur la base de la valeur d'un attribut non-catégorique X en ensembles $T_1, T_2, \dots, T_n$ alors l'information nécessaire à l'identification de la classe d'un élément de T devient la moyenne pondérée des quantités d'information $I(T_i)$ nécessaires à l'identification de la classe d'un élément de T_i :

$$I(X, T) = \sum_{i=1}^n |T_i|/|T| \cdot I(T_i) \quad \text{VI-13}$$

On peut alors introduire la quantité $\text{Gain}(X, T)$ définie comme :

$$\text{Gain}(X, T) = I(T) - I(X, T) \quad \text{VI-14}$$

Cette quantité représente la différence entre l'information nécessaire à l'identification d'un élément de T et l'information nécessaire à l'identification d'un élément de T après que l'on a obtenu la valeur de l'attribut X , c'est-à-dire, le gain d'information dû à l'attribut X .

On peut alors ensuite utiliser cette information de gain d'information pour ranger les attributs et construire des arbres de décision où chaque nœud correspond à l'attribut avec le plus grand gain parmi les attributs encore non considérés dans le chemin depuis la racine.

Pour les attributs ayant des échelles continues on peut procéder de la façon suivante. Imaginons qu'un attribut C_i soit à valeurs continues. On examine les valeurs pour cet attribut sur l'ensemble d'apprentissage et on les range par ordre croissant $A_1, A_2, \dots, A_m$. Alors pour chaque valeur $A_j, j=1, 2, \dots, m$, on partitionne l'enregistrement en une population qui prend des valeurs pour C_i inférieures ou égales à A_j et une population qui prend des valeurs supérieures à A_j . Pour chacune de ces partitions on calcule le gain, et on choisit la partition qui maximise le gain d'information.

1.3. Autre mesure du meilleur candidat

La méthode CART de [Crawford *et al.*, 1991] utilise une autre fonction pour choisir le meilleur candidat à chaque étape du processus, qui mesure le "degré de mélange" des différentes classes. CART utilise la fonction de Gini : $Gini(I(p)) = 4p(1-p)$

2. Calcul d'un modèle de classification par la méthode des SVM

La méthode SVM (Support Vector Machines) [Burges *et al.*, 1995; Christianini *et al.*, 2000] est un algorithme d'apprentissage supervisé pour les problèmes de classification à 2 classes.

L'ensemble d'apprentissage est donné par un ensemble de vecteurs associé à leur classe d'appartenance : $(X_1, y_1), \dots, (X_u, y_u)$, $X_j \in \mathbb{R}^n, y_j \in \{+1, -1\}$, avec :

- y_j représente la classe d'appartenance. Dans un problème à deux classes, on peut considérer que la première classe correspond à une réponse positive (+1), c'est-à-dire y_j égal à +1, et la deuxième classe correspond à une réponse négative (-1), c'est-à-dire y_j égal à -1 ;
- $X_j \in \mathbb{R}^n$, X_j représente le vecteur du texte numéro j de l'ensemble d'apprentissage. La dimension de ce vecteur est n . X_j est un vecteur occurrence correspondant à la définition donnée ci avant.

La méthode SVM sépare les vecteurs à classe positive et les vecteurs à classe négative par un hyperplan défini par l'équation suivante : $W.X + b = 0$, $W \in \mathbb{R}^n, b \in \mathbb{R}$, où « . » représente le produit scalaire. En général, un tel hyperplan n'est pas unique. La méthode SVM détermine l'hyperplan optimal en maximisant la marge : la marge est la distance entre les vecteurs étiquetés positifs et les vecteurs étiquetés négatifs. L'ensemble d'apprentissage n'est pas nécessairement séparable linéairement, des variables d'écart ξ_j sont introduites pour tous les X_j . Ces ξ_j prennent en compte l'erreur de classification, et doivent satisfaire les inégalités suivantes :

- $W.X_j + b \geq 1 - \xi_j$;
- $W.X_j + b \leq -1 + \xi_j$.

En prenant en compte ces contraintes, on doit minimiser la fonction d'objectif suivante :

$\frac{1}{2}\|W\|^2 + C \sum_{j=1}^u \xi_j$. Le premier terme de cette fonction correspond à la taille de la marge et le second terme représente l'erreur de classification, avec u représentant le nombre de vecteurs de l'ensemble d'apprentissage.

Trouver la fonction objectif précédente revient à résoudre le problème quadratique suivant : trouver la fonction de décision : $f(X) = \text{signe}(g(X))$ dans laquelle la fonction $g(X)$ est :

$$g(X) = \sum_{i=1}^m \lambda_i y_i X_i \cdot X + b, \text{ avec :}$$

- $\text{Signe}(x)$ représente la fonction suivante :
 - Si $x > 0$ alors $\text{Signe}(x) = 1$
 - Si $x < 0$ alors $\text{Signe}(x) = -1$
 - Si $x = 0$ alors $\text{Signe}(x) = 0$
- y_j représente la classe d'appartenance,
- λ_i représente les paramètres à trouver,
- $X_i \cdot X$ représente le produit scalaire du vecteur X_i avec le vecteur X .

La fonction de décision $f(X)$ dépend uniquement des « vecteurs de support » X_i . Les vecteurs des textes de l'ensemble d'apprentissage, excepté ceux correspondant aux « vecteurs de support » n'ont aucune influence sur la fonction de décision à calculer.

Les surfaces de décision non linéaires peuvent être calculées en remplaçant le produit scalaire $X_i \cdot X$ par une fonction de noyau $K(X_i, X)$: $g(X) = \sum_{i=1}^m \lambda_i y_i K(X_i, X) + b$. VI-15

Pour cette fonction K , des fonctions polynomiales donnent de très bons résultats comme par exemple : $K(X, Y) = (X \cdot Y + 1)^d$, (voir à ce sujet [Joachims, 1998]).

N° d'ordre : 411 I

First name, Name: Michel Plantié

Title of the thesis: Automatic Knowledge Extraction pour Multicriteria Decision Making

Speciality: Computer Science

Keywords

Decision-Making, Information Decision-Making Support System, Knowledge management, Actionable Knowledge, Information Fusion, Explanation, Decisional Risk, Text-Mining, Data-Mining, Classification, Automated indexation

Summary

This thesis deals with the biased subject of cognitive automation without really coming to a decision. A complete information processing system to support each step of a collective decision-making process (DMP) is proposed. The automation of the learning stage in a DMP is emphasized: the actionable knowledge, i.e. the knowledge relevant for action, becomes a data that can be manipulated by algorithms.

The model that supports our interactive decision-making support system (IDMSS) widely relies on automated knowledge processing. Data-mining, multi criteria decision and optimization techniques are together required in our IDMSS that corresponds to a cybernetic interpretation of the model of decision by the economist Simon. The epistemic incertitude inherent to a decision-making is captured by the notion of decisional risk that analyses the discriminating factors between competing alternatives. Several attitudes in the control of decisional risk may be envisaged: the IDMSS can be used to validate, check or refute an opinion. In every case, the control of epistemic incertitude is not neutral w.r.t. the dynamics of the DMP. The instrumentation of the learning stage in the DMP leads to elaborate the actuator of a retroaction loop that enables to control the dynamics of the DMP. Our approach provides a new framework to formalize the relations between epistemic incertitude, decisional risk and decision stability concepts.

An analysis of the fundamental concepts of actionable knowledge and of automated indexation that support our models and tools is proposed. The actionable knowledge concept acquires a new interpretation in this cybernetic framework: that is the knowledge manipulated by the IDMSS actuator to control the dynamics of a DMP. A brief synthesis of the most popular learning techniques in knowledge discovery is proposed. All these techniques are instantiated for the automated extraction of actionable knowledge in a multi criteria evaluation process. The problem of a video-club manager who tries to optimize his investments w.r.t. the preferences of his customers provides a relevant application that illustrates our whole information processing system.

N° d'ordre : 411 I

Prénom Nom : Michel Plantié

Titre de la thèse : Extraction automatique de Connaissances pour la décision multicritère

Spécialité : Informatique

Mots clefs

Décision, Système d'Aide à la Décision, Gestion des Connaissances, Connaissance Actionnable, Fusion d'Informations, Explication, Argumentation, Risque décisionnel, Text-Mining, Datamining, TALN, Classification, Indexation automatique.

Résumé

Cette thèse, sans prendre parti, aborde le sujet délicat qu'est l'automatisation cognitive. Elle propose la mise en place d'une chaîne informatique complète pour supporter chacune des étapes de la décision. Elle traite en particulier de l'automatisation de la phase d'apprentissage en faisant de la connaissance actionnable—la connaissance utile à l'action—une entité informatique manipulable par des algorithmes.

Le modèle qui supporte notre système interactif d'aide à la décision de groupe (SIADG) s'appuie largement sur des traitements automatiques de la connaissance. Datamining, multicritère et optimisation sont autant de techniques qui viennent se compléter pour élaborer un artefact de décision qui s'apparente à une interprétation cybernétique du modèle décisionnel de l'économiste Simon. L'incertitude épistémique inhérente à une décision est mesurée par le risque décisionnel qui analyse les facteurs discriminants entre les alternatives. Plusieurs attitudes dans le contrôle du risque décisionnel peuvent être envisagées : le SIADG peut être utilisé pour valider, vérifier ou infirmer un point de vue. Dans tous les cas, le contrôle exercé sur l'incertitude épistémique n'est pas neutre quant à la dynamique du processus de décision. L'instrumentation de la phase d'apprentissage du processus décisionnel conduit ainsi à élaborer l'actionneur d'une boucle de rétroaction visant à asservir la dynamique de décision. Notre modèle apporte un éclairage formel des liens entre incertitude épistémique, risque décisionnel et stabilité de la décision.

Les concepts fondamentaux de connaissance actionnable (CA) et d'indexation automatique sur lesquels reposent nos modèles et outils de TALN sont analysés. La notion de connaissance actionnable trouve dans cette vision cybernétique de la décision une interprétation nouvelle : c'est la connaissance manipulée par l'actionneur du SIADG pour contrôler la dynamique décisionnelle. Une synthèse rapide des techniques d'apprentissage les plus éprouvées pour l'extraction automatique de connaissances en TALN est proposée. Toutes ces notions et techniques sont déclinées sur la problématique spécifique d'extraction automatique de CAs dans un processus d'évaluation multicritère. Enfin, l'exemple d'application d'un gérant de vidéoclub cherchant à optimiser ses investissements en fonction des préférences de sa clientèle reprend et illustre le processus informatisé dans sa globalité.