

Laboratoire d'Informatique
Fondamentale de Lille



THÈSE

Présentée à

L'UNIVERSITÉ DES SCIENCES ET TECHNOLOGIES DE LILLE

Pour obtenir le titre de

DOCTEUR EN INFORMATIQUE

par

Sébastien KONIECZNY



SUR LA LOGIQUE DU CHANGEMENT : RÉVISION ET FUSION DE BASES DE CONNAISSANCE

Thèse soutenue le 29 novembre 1999

Composition du jury :

Rapporteurs :	Anthony HUNTER	University College de Londres
	Jérôme LANG	Université de Toulouse 3
	Pierre MARQUIS	Université d'Artois
Examineurs :	Jean-Paul DELAHAYE	Université de Lille 1
	Ramón PINO PÉREZ	Université d'Artois
	Karl SCHLECHTA	Université d'Aix-Marseille 1
	Carlos UZCÁTEGUI	Université des Andes

Université des Sciences et Technologies de Lille
LIFL - U.F.R. d'I.E.E.A. - Bât. M3. 59655 Villeneuve d'Ascq CEDEX
Tél. 03.20.43.47.24 Fax. 03.20.43.65.66

Ce qui est créé par l'esprit est plus vivant que la matière.
Charles Baudelaire

Remerciements

Je ne peux exprimer ici toute ma gratitude envers mon directeur de thèse, Ramón Pino Pérez. C'est en grande partie grâce à lui que ces années de thèse ont été si intéressantes. Il m'a d'abord initié à ce domaine fascinant de la modélisation du raisonnement, il m'a laissé explorer ce dédale lorsque je trouvais un chemin intéressant, mais a été suffisamment présent pour me remettre sur la bonne voie lorsque je m'égarais. J'ai pu apprécier pendant ces années ses qualités scientifiques indiscutables mais également des qualités humaines rares. Je veux donc remercier ici le maître, le collègue, et l'ami. J'espère simplement qu'il trouvera que ce manuscrit est à la hauteur du temps incalculable qu'il m'a consacré.

Je tiens à remercier également toutes les personnes qui ont participé aux séminaires GNOM, et plus particulièrement Stéphane Janot et Hassan Bezzazi. Ces réunions ont toujours été agréables et fructueuses.

Je remercie les membres du jury de l'honneur qu'ils me font en ayant accepté de s'intéresser à ce travail. Merci à Anthony Hunter, Pierre Marquis et Jérôme Lang d'avoir accepté si chaleureusement de rapporter ce travail.

Merci à Anthony Hunter dont les questions m'ont aidé à éclaircir certains passages et m'ont suggéré de nouvelles pistes de recherche.

Je remercie Pierre Marquis pour son soutien moral et pour ses nombreuses remarques qui m'ont permis d'améliorer ce manuscrit tant sur le fond que sur la forme.

Je remercie également Jérôme Lang pour l'intérêt qu'il a porté à mon travail pendant ces années. Les conversations que nous avons eues, scientifiques ou autres, ont toujours été très enrichissantes.

Merci à Jean-Paul Delahaye d'avoir bien voulu participer à ce jury. Son ouverture d'esprit et sa curiosité pour des domaines variés sont pour moi un exemple.

Je remercie Karl Schlechta d'avoir accepté de participer au jury. Ses travaux sur la définition d'opérateurs de révision à partir de distances m'ont beaucoup inspiré.

Je remercie également Carlos Uzcátegui pour sa participation au jury. Les discussions que nous avons eues lorsque j'ai commencé ma thèse, ainsi que les séminaires qu'il a donnés alors, m'ont apporté beaucoup.

Un merci particulier à tous les membres permanents, temporaires et honoraires du bureau 318, Stéphane, Jean-Marc, Hélène, Tof, Olivier, ainsi qu'à mes deux "camarades de tranchée" Bruno et Jean-Stéphane. Je n'aurais pas pu imaginer meilleur environnement de travail. Merci à tous pour toutes ces discussions si délicieusement inutiles.

Merci à Isabelle, Claude, Johan, Eric, Laurent et les autres bouffons... Merci simplement d'être vous, et d'être un peu de moi.

Merci à ma famille, et particulièrement à mes parents, mon frère et ma soeur, et à Christophe et Béatrice, de continuer à me supporter avec autant de vaillance!

Merci enfin aux relecteurs de cette thèse, Isabelle, Stéphane, Claude, et merci à Jean-Marc dont la relecture plus qu'attentive aurait dû lui valoir le titre de quatrième rapporteur...

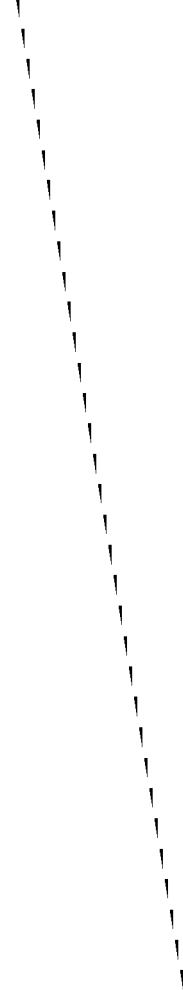


Table des matières

Introduction	11
Préliminaires	19
I Révision	21
1 Le cadre AGM	23
1.1 Expansion-Révision-Contraction	24
1.1.1 Expansion	25
1.1.2 Révision	26
1.1.3 Contraction	27
1.1.4 Identités	28
1.2 Théorèmes de représentation	29
1.2.1 Contraction par intersection partielle	30
1.2.2 Enracinements épistémiques	32
1.2.3 Contraction sûre	33
1.2.4 Assignment fidèle et système de sphères	34
1.2.5 Révision et relations d'inférence	38
1.2.6 Logique possibiliste	41
1.3 Critiques de la révision AGM	43
1.3.1 Restauration	43
1.3.2 Révision syntaxique	45
1.3.3 Itération	49
1.3.4 Semi-révision	52
1.3.5 Mise à jour	54
2 Révision itérée	57
2.1 Proposition de Darwiche et Pearl pour l'itération	58
2.1.1 AGM pour états épistémiques	58
2.1.2 Les postulats supplémentaires de Darwiche et Pearl	60
2.2 Révision naturelle	62
2.3 Révision rangée	64
2.4 Fonctions ordinales conditionnelles et transmutations	68
2.5 Entrenchment Kinematics	71
2.6 Cas limites	73

3	Principe de primauté forte	77
3.1	De la notion d'état épistémique	79
3.2	Révision à mémoire	82
3.3	Opérateurs de révision à mémoire et opérateurs AGM	88
3.4	Quelques opérateurs de révision à mémoire	89
3.4.1	Opérateur basique à mémoire	89
3.4.2	Opérateur Dalal à mémoire	92
3.4.3	Les opérateurs OTP	94
3.4.4	Exemple	97
3.5	Révision dynamique	97
3.6	Révision par états épistémiques	101
3.7	Semi-révision à mémoire	104
3.8	Remarques	104
4	Révision et chaînage avant	107
4.1	Définition des opérateurs	108
4.1.1	Flock	109
4.1.2	Mise à jour factuelle	110
4.1.3	Révision hiérarchique	111
4.1.4	Révision gangue	111
4.1.5	Révision gangue étendue	112
4.1.6	Exemples	113
4.1.7	Coder l'opérateur de révision gangue	114
4.2	Propriétés logiques	115
4.3	Révision gangue sélective	120
4.4	Conclusion	122
II	Fusion	123
5	Introduction	125
5.1	Notations	125
5.2	Autres travaux	126
5.2.1	Revesz	127
5.2.2	Liberatore et Schaerf	129
5.2.3	Lin et Mendelzon	131
6	Fusion contrainte	133
6.1	Caractérisation logique	133
6.2	Discussion sur les postulats	139
6.2.1	Equité	139
6.2.2	Majorité	140
6.2.3	Arbitrage	141
6.3	Théorèmes de représentation	142

7	Quelques opérateurs de fusion	151
7.1	Opérateurs simples	151
7.2	Distances	152
7.3	Opérateurs Max	153
7.4	Opérateurs Σ et Σ^n	155
7.5	Opérateurs GMax	157
7.6	Exemples	159
7.7	Opérateurs d'arbitrage et opérateurs majoritaires	162
7.7.1	Distance drastique	162
7.7.2	Etude graphique	162
8	Fusion contrainte et autres travaux	167
8.1	Fusion contrainte et révision AGM	167
8.2	Fusion contrainte et fusion pure	171
8.2.1	Postulats	171
8.2.2	Caractérisations sémantiques	172
8.3	Fusion contrainte et révision commutative	174
8.4	Fusion contrainte et théorie du choix social	180
8.4.1	Le théorème d'impossibilité d'Arrow	180
8.4.2	Impossibilité de la fusion contrainte?	183
9	Fusion syntaxique	187
9.1	Opérateurs de combinaison	188
9.2	Fonctions de sélection	192
9.2.1	Opérateur majoritaire drastique	193
9.2.2	Les opérateurs cardinaux	194
9.3	Remarques	198
	Conclusion	199
	Index des propriétés logiques	205
	Index	211

Introduction

L'intelligence artificielle est l'étude des concepts qui permettent de rendre les machines intelligentes. Vue la difficulté qu'ont les philosophes à définir cette notion d'intelligence, on ne peut pas s'attendre à ce que le domaine de l'intelligence artificielle soit nettement délimité. Mais, bien que l'on puisse discuter sur les limites exactes de cette frontière, il est indéniable que la capacité à mener un raisonnement est une des multiples expressions de ce que l'on nomme intelligence. De ce fait, la mécanisation de ce raisonnement a une relation très intime avec l'intelligence artificielle. Cette thèse est une contribution à l'intelligence artificielle et plus précisément aux modèles du raisonnement dynamique.

La logique, depuis sa formalisation mathématique à la fin du XIX^e siècle, est l'outil privilégié pour modéliser le raisonnement valide. Elle offre une première réponse à la mécanisation du raisonnement, au travers de la démonstration automatique. Mais, bien que la logique classique soit bien adaptée au raisonnement valide, le problème est qu'elle ne peut pas être utilisée pour des raisonnements simplement plausibles. Or, la plupart des raisonnements que nous effectuons ne sont pas des raisonnements valides. Lorsque nous raisonnons, nous utilisons des informations incomplètes et imparfaites. Si l'on m'apprend que *Titi est un oiseau*, j'en déduirai que *Titi vole* alors que je n'ai aucune certitude à ce sujet. De même, si j'apprends ensuite que *Titi est un manchot*, je remplacerai dans mes connaissances *Titi vole* par *Titi ne vole pas* sans que cela ne me pose le moindre problème. Mais quelles difficultés cela pose en logique ! Il suffit de compter le nombre de propositions qui ont été faites pour tenter de faire atterrir *Titi* pour s'en rendre compte. L'inconvénient de la logique classique pour résoudre ce genre de problèmes réside en sa monotonie. C'est-à-dire qu'augmenter le nombre d'hypothèses, augmente le nombre de conclusions. Il n'est donc pas possible de remettre en cause une conclusion en augmentant le nombre de prémisses. Ainsi, lorsque l'on déduit que *Titi vole*, il ne sera pas possible d'enlever ce fait quelle que soit la nouvelle information.

Lorsque l'on veut mener un raisonnement à propos d'un problème dont on ne connaît pas tous les paramètres, comme c'est généralement le cas, et où certaines de nos connaissances peuvent être remises en question parce que celles-ci étaient inexactes ou parce que le monde a changé, la logique classique ne suffit pas. Pour résoudre ces problèmes de *raisonnement de sens commun*, plusieurs méthodes ont été proposées. Elles ont toutes en commun le fait de ne pas être monotones, et ont donc été nommées *logiques non monotones*. Mais certaines méthodes avaient très peu de contenu logique. Les travaux de Gabbay et de Makinson [Gab85, Mak89, Mak94] ont été une avancée importante de ce point de vue. Ces derniers ont proposé non pas d'étudier des méthodes particulières mais de considérer ces méthodes comme des relations d'inférence et d'étudier les propriétés logiques que pouvaient satisfaire les relations de conséquence associées. Cela a permis de comparer plus finement les différentes

méthodes proposées et de définir des familles de méthodes ayant des propriétés particulières.

Cette thèse aborde les domaines des modèles de raisonnement et des logiques non monotones. Le problème abordé n'est pas directement celui des logiques non monotones mais on s'intéresse à la notion duale d'opérateurs de révision de la connaissance. Comme le souligne Peter Gärdenfors [Gär90]:

“Belief revision and nonmonotonic logic are two sides of the same coin”

Le problème posé par la révision de la connaissance est de déterminer l'état d'une base de connaissance après qu'une nouvelle information y ait été ajoutée, cette nouvelle information étant bien souvent en conflit avec certaines connaissances de la base. Un problème similaire se pose lorsque l'on dispose de plusieurs sources d'information et que l'on désire exploiter ces informations potentiellement contradictoires. On parle alors de fusion de bases de connaissance. Notre contribution est donc une étude des opérateurs de révision et de fusion de bases de connaissance.

Révision de la connaissance

Beaucoup de problèmes en intelligence artificielle reposent sur l'exploitation d'une base de connaissance. Cela est particulièrement vrai pour les systèmes experts, pour le diagnostic, pour la planification... Or, lorsque l'on travaille avec un système à base de connaissance, il est nécessaire de pouvoir faire évoluer cette base. Pour les systèmes experts, une nouvelle connaissance doit pouvoir être incorporée pour augmenter l'expertise du système. Pour le diagnostic, une nouvelle règle concernant le phénomène étudié doit pouvoir être ajoutée. Pour la planification, si l'on veut réaliser un robot autonome, il est nécessaire que celui-ci puisse faire face à une situation éventuellement différente de celle qui était prévue initialement. Ainsi la dynamique de la connaissance est un point crucial dans de nombreux domaines de l'intelligence artificielle et c'est donc un domaine clef pour la conception de machines “intelligentes”. Le cadre traitant de cette dynamique de la connaissance est nommé *théorie de la révision*. Le problème posé par la théorie de la révision est simple : étant données une base de connaissance et une nouvelle information, quelles modifications dois-je effectuer à la base de connaissance pour incorporer la nouvelle information ?

Dans le domaine des bases de données, le problème de la révision est une question primordiale. Illustrons cela sur l'exemple suivant donné par [FUV83] :

“Considérons par exemple une base de données relationnelle avec une relation ternaire FOURNIT, où un tuple $\langle a, b, c \rangle$ signifie que le fournisseur a fournit l'élément b au projet c . Supposons à présent que la relation contient le tuple $\langle \text{Hughes}, \text{tôles}, \text{Navette Spatiale} \rangle$, et que l'utilisateur demande d'effacer ce tuple. Une approche simpliste serait d'effectuer cette requête et d'effacer simplement le tuple de la relation. Toutefois, bien qu'il soit vrai que Hughes ne fournit plus de tôles pour le projet de la navette spatiale, il n'est pas évident de savoir comment manier trois autres faits qui étaient impliqués par le tuple, à savoir que Hughes fournit des tôles, que Hughes fournit des éléments du projet de la navette spatiale, et que le projet de la navette spatiale utilise des tôles.”

Ainsi, de nombreux opérateurs ont été proposés dans ce cadre [FUV83, FKUV86, Bor85, Dal88a, Win88].

Les philosophes s'intéressent également à la révision de la connaissance puisque cela recouvre la question des théories scientifiques dans le cas de ce que Bachelard appelle une rupture épistémologique. La question devient alors : "que se passe-t-il lorsqu'une découverte remet en cause une théorie existante?". Un autre problème que se posent les philosophes est simplement de modéliser l'effet d'une nouvelle information sur les connaissances d'un individu. Sous cet angle, la théorie de la révision tente de modéliser le raisonnement humain.

Ce problème de la révision de la connaissance a été abordé de différentes manières, suivant que l'on étudie le problème sous un aspect quantitatif ou qualitatif, probabiliste ou logique, donnant naissance à des cadres de travail très différents. On peut par exemple citer le cadre des probabilités qualitatives [GP96, Gef92, Spo87, Spo90], les fonctions de croyance [SK94], la théorie des possibilités, etc...

Dans le cadre logique, c'est le cadre AGM [AGM85, Gär88] (d'après Carlos Alchourrón, Peter Gärdenfors et David Makinson) qui s'est imposé. Ces trois auteurs ont proposé une caractérisation logique des opérateurs de révision, c'est-à-dire un ensemble de propriétés qu'un opérateur de révision "sensé" doit satisfaire. Il s'avère que cette caractérisation logique correspond à des méthodes de révision assez intuitives. C'est ce que l'on appelle les théorèmes de représentation. En effet, pour chacune de ces méthodes de révision, il a été montré qu'un opérateur était défini à partir de cette méthode si et seulement si il vérifiait les propriétés AGM. Ce sont ces théorèmes qui donnent une consistance au cadre AGM et en constituent, en quelque sorte, les fondations. De plus, il a été montré que le cadre AGM correspondait à d'autres domaines comme les logiques non monotones et la logique possibiliste.

D'un autre côté, la caractérisation logique n'est qu'un ensemble de propriétés logiques et il est toujours possible, pour une application particulière, de critiquer tel ou tel postulat. Les fondations du cadre AGM, constituées des théorèmes de représentation et des correspondances aux domaines connexes, ne font pas craindre l'effondrement de l'édifice sous les critiques. Il est, en effet, possible de critiquer le cadre AGM, mais ces critiques ne remettent pas ce cadre en question. Elles suggèrent plutôt des pistes vers des adaptations ou des généralisations de ce cadre.

Dans cette thèse nous proposons d'étudier trois généralisations de ce cadre.

Tout d'abord le cadre AGM ne contraint pas suffisamment la révision pour capturer des propriétés satisfaisantes pour l'itération du processus de révision. Or, lorsque l'on désire étudier le comportement d'un système à base de connaissance, il est nécessaire de pouvoir itérer ce processus. La caractérisation AGM ne s'intéresse qu'à une seule itération de ce processus. Or, pour avoir de bonnes propriétés pour l'itération, il est nécessaire de contraindre les révisions successives aux itérations précédentes, c'est-à-dire avoir une stratégie de révision. La caractérisation AGM suppose l'existence d'une stratégie de révision extérieure aux opérateurs, le problème de la révision itérée est de coder et de manipuler cette stratégie de révision dans les opérateurs.

Un autre problème posé par le cadre AGM est que celui-ci travaille à partir de bases closes

déductivement, appelées également théories. Cette hypothèse rend les méthodes de révision usuelles difficilement implémentables, car le fait de devoir manipuler toutes les conséquences logiques d'une base de connaissance pose, entre autres, d'importants problèmes de représentation. Plusieurs solutions existent pour réduire la complexité de ce calcul. On peut par exemple travailler avec des bases non closes déductivement, utiliser des méthodes incomplètes ou affaiblir la logique. Nous étudions des opérateurs travaillant sur des bases non closes déductivement et basés sur une logique induite par le chaînage avant.

Une dernière généralisation du cadre AGM est de supposer que l'information à réviser est en fait un ensemble de bases de connaissance. Cela permet de définir la révision AGM comme un cas particulier de fusion de connaissances. Le problème de la fusion de bases de connaissance constitue la deuxième partie de ce mémoire et nous motivons donc son intérêt indépendamment du problème de la révision.

Fusion de bases de connaissance

La tendance actuelle, tant en intelligence artificielle qu'en informatique en général, va à la distribution de l'information. Cela permet de construire des systèmes complexes à partir d'éléments spécialisés : les systèmes à bases de connaissance sont à présent des systèmes multi-agents, un problème majeur en bases de données est l'intégration de bases de données hétérogènes, et le world wide web contribue également à la distribution de données à l'échelle planétaire. Or une fois réglés les problèmes de négociation et de langage pour les systèmes multi-agents, les problèmes de la détermination d'un schéma global pour les bases de données et les problèmes de langages sur le web, il reste une question commune et essentielle à traiter : que faire lorsque les sources d'information fournissent des informations contradictoires ?

Pour illustrer le problème, considérons l'exemple suivant : on désire construire un système expert. Pour ce faire, on consulte plusieurs experts humains et on code chaque expertise dans une base de connaissance. Pour réaliser le système expert, le but est alors de définir ce qu'est la connaissance du groupe d'experts. La question revient alors à fusionner l'ensemble des bases de connaissance. Ce processus n'est pas trivial, puisque opérer simplement la réunion de toutes les bases de connaissance mènera souvent à des inconsistances. Il est donc nécessaire de définir des opérateurs capables de manier ces bases de connaissance de manière adéquate.

Il est possible de tenter d'éviter le problème par une phase de négociation dans les systèmes multi-agents ou en supposant l'existence d'un oracle, comme cela est fréquemment fait en bases de données. Mais il y aura toujours des cas où la négociation échouera, ou où l'oracle ne pourra pas trancher. Le problème de la fusion de bases de connaissance est donc un problème important dont les applications sont nombreuses dans plusieurs domaines majeurs.

Nous proposons une caractérisation logique de ces opérateurs de fusion, inspirée de la caractérisation AGM pour la révision. Nous distinguons en particulier deux classes d'opérateurs de fusion : les opérateurs majoritaires et les opérateurs d'arbitrage. Les premiers effectuent la fusion en donnant la priorité à la majorité. Les seconds ont un comportement plus égalitaire, tentant de satisfaire toutes les bases. Nous donnons également un théorème de représentation pour ces opérateurs.

Ce problème de la fusion de bases de connaissance est très proche de la théorie du vote, de la théorie du choix social et de la théorie de la décision. Ces différents domaines travaillent sur des objets différents, mais les méthodes et les buts sont similaires.

Plan de la thèse

Cette thèse est composée de deux parties. La première partie est consacrée au cadre AGM de la théorie de la révision. Dans la seconde, nous étudions les opérateurs de fusion de bases de connaissance.

Le chapitre 1 tente de faire un résumé succinct mais complet du cadre AGM. Pour le lecteur voulant se familiariser avec le domaine, nous espérons que ce manuscrit pourra servir de point de départ. Nous avons voulu en faire un état de l'art du domaine. Ce survol peut être complété par la lecture du livre de Peter Gärdenfors, *Knowledge in Flux* [Gär88] qui détaille et motive la caractérisation logique des opérateurs de changement de la connaissance donnée à la section 1.1 et aborde plusieurs théorèmes de représentation de la section 1.2. S'appuyer sur ces deux écrits doit permettre de se faire une idée juste du cadre AGM de la théorie de la révision. Dans une première section, nous présentons la caractérisation logique des opérateurs de changement proposée par Alchourrón, Gärdenfors et Makinson. Nous posons ensuite les fondations du cadre AGM en énumérant les différents théorèmes de représentation proposés pour les opérateurs de révision AGM. Ces théorèmes montrent que les opérateurs vérifiant la caractérisation logique proposée par Alchourrón, Gärdenfors et Makinson correspondent à des méthodes de révision naturelles. Dans cette section nous soulignons également les rapports étroits entre le cadre AGM et les domaines connexes que sont les logiques non monotones et la logique possibiliste. C'est cet ensemble de théorèmes de représentation et de correspondances avec d'autres domaines qui apporte une consistance à la caractérisation logique. Les critiques du cadre AGM, données dans la section 1.3, ne le remettent donc pas en cause. Mais elles suggèrent des adaptations ou des généralisations pour traiter d'autres types d'opérateurs de changement de la connaissance.

Le chapitre 2 est une présentation exhaustive des principales propositions concernant la révision itérée. Nous avons tenté de faire apparaître les inconvénients et avantages de chacune et nous avons souligné les connexions entre les différentes propositions. La section 2.1 présente la proposition de Darwiche et Pearl pour la révision itérée. Ils ont proposé l'utilisation d'une information plus complexe qu'une base de connaissance, nommée état épistémique, pour coder l'état cognitif d'un agent. Cela permet de coder les croyances de l'agent, *i.e.* sa stratégie de révision. La section 2.2 décrit l'opérateur de révision naturelle proposé par Craig Boutilier comme un cas particulier d'opérateur de Darwiche et Pearl. Les opérateurs de révision rangée, proposés par Daniel Lehmann, sont présentés section 2.3. Ces opérateurs sont basés sur la notion de séquences de révisions. Les postulats proposés par Lehmann contraignent donc les séquences, plutôt qu'une seule itération comme c'est le cas pour les postulats AGM. Nous présentons ensuite, section 2.4, le cadre des fonctions ordinales conditionnelles (OCF) proposé par Wolfgang Spohn. Ce cadre n'est pas purement quantitatif car la nouvelle information est accompagnée d'un degré de confiance. Spohn propose alors une famille d'opérateurs, nommés conditionnalisations, comme opérateurs de changement. Mary-Anne Williams a étudié fi-

nement les opérateurs travaillant sur les OCF, elle nomme ces opérateurs des transmutations. Elle propose en particulier une famille d'opérateurs d'ajustement et montre que les opérateurs d'ajustement et de conditionnalisation sont interdéfinissables. Les opérateurs travaillant sur ces OCF ont de bonnes propriétés pour l'itération et peuvent suggérer la définition d'opérateurs dans un cadre purement qualitatif. Enfin, nous résumons à la section 2.5 la proposition de Nayak *et al.*. La dernière section (section 2.6) compare ces différentes propositions en ce qui concerne leur cas limite. En effet, après un certain nombre d'itérations, la plupart de ces opérateurs acquièrent le même comportement : soit les préférences (croyances) de l'agent deviennent très strictes, ce comportement pouvant être interprété comme un processus d'apprentissage, soit l'agent perd toute préférence et ce comportement semble dénoter un processus d'oubli de la part de l'agent.

Les deux derniers chapitres de cette partie concernent notre contribution au problème de la révision.

Nous présentons, dans le chapitre 3, notre proposition pour la révision itérée. Nous soulignons tout d'abord la nécessité d'avoir une définition précise des états épistémiques étudiés dans les différents travaux. Nous proposons donc, section 3.1, une syntaxe et les sémantiques correspondantes pour les états épistémiques utilisés dans les travaux présentés au chapitre précédent et dans notre approche. Nous présentons ensuite à la section 3.2 les opérateurs de révision à mémoire. Ces opérateurs, grâce à un renforcement du principe de primauté de la nouvelle information, permettent d'apporter une réponse au problème de la révision itérée. Nous montrons section 3.3 que ces opérateurs nécessitent la même quantité d'information que les opérateurs de révision AGM classiques. Nous donnons également, comme exemple d'opérateurs à la section 3.4, deux généralisations de l'opérateur de révision sur des Présentations Ordonnées de Théorie (OTP) proposé par Mark Ryan. Alors que l'opérateur de Ryan ne satisfaisait pas les propriétés de base des opérateurs de révision, les deux généralisations satisfont ces propriétés et celles des opérateurs de révision à mémoire. Nous présentons ensuite trois généralisations des opérateurs de révision à mémoire. Section 3.5, nous définissons des états épistémiques plus complexes que ceux utilisés dans les différents travaux sur la révision itérée, pour permettre de rendre les opérateurs de révision à mémoire plus dynamiques. Une autre généralisation possible est de considérer que la nouvelle information n'est pas une simple formule mais un véritable état épistémique, ce qui est fait section 3.6. Cela permet, en particulier, de réviser la connaissance par des informations conditionnelles. Enfin, section 3.7, nous montrons comment utiliser les opérateurs de révision à mémoire lorsque l'on a davantage confiance en la connaissance actuelle qu'en la nouvelle information.

Le chapitre 4 aborde le problème de la révision syntaxique. Nous proposons et étudions des opérateurs de révision syntaxiques basés sur une logique induite par le chaînage avant. La plupart de ces opérateurs sont basés sur une hiérarchie entre formules directement induite par la base de connaissance. Cette construction est inspirée de la clôture rationnelle de bases conditionnelles proposée par Kraus, Lehmann et Magidor. La différence entre les opérateurs réside en la manière dont cette hiérarchie est utilisée.

Dans une seconde partie nous abordons le problème de la fusion de bases de connaissance.

Le chapitre 5 est une introduction à la seconde partie. Il donne la problématique posée

par la fusion de bases de connaissance dans une première section. La section 5.1 rassemble les définitions et les notations propres à cette partie. On y définit notamment la notion d'ensemble de connaissance. La section 5.2 présente les autres propositions existantes, c'est-à-dire les opérateurs d'adéquation sémantique proposé par Peter Revesz, les opérateurs de révision commutative proposés par Liberatore et Schaerf et les opérateurs de fusion majoritaire de théories proposés par Lin et Mendelzon.

Le chapitre 6 expose notre proposition pour la fusion de bases de connaissance. Nous proposons section 6.1, une caractérisation logique des opérateurs de fusion contrainte. Nous distinguons, en particulier, trois sous-classes d'opérateurs de fusion particuliers : les opérateurs de quasi-fusion, les opérateurs de fusion majoritaire et les opérateurs d'arbitrage. Nous donnons ensuite, à la section 6.3 les théorèmes de représentation pour les différents types d'opérateurs définis. Ces théorèmes montrent que les opérateurs de fusion contrainte correspondent à des familles de pré-ordres sur les interprétations.

Nous donnons au chapitre 7 des exemples d'opérateurs de fusion. Nous montrons tout d'abord que la caractérisation logique est suffisamment contraignante pour éliminer des opérateurs trop "simples". Nous définissons ensuite trois familles d'opérateurs de fusion définis à partir d'une distance. La première famille, que nous nommons opérateurs *Max*, donne des opérateurs de quasi-fusion. La seconde famille, les opérateurs Σ , est composée d'opérateurs de fusion majoritaires et la dernière famille, les opérateurs *GMax*, définit des opérateurs d'arbitrage. Dans une dernière section, nous montrons qu'il est possible pour un opérateur d'être à la fois majoritaire et d'arbitrage.

Le chapitre 8 compare notre proposition à divers autres travaux. Nous montrons que les opérateurs de révision AGM sont un cas particulier des opérateurs de fusion contrainte, et nous montrons comment définir un opérateur de fusion contrainte à partir d'un opérateur de révision. Section 8.2, nous étudions la définition d'opérateurs de révision lorsqu'il n'y a pas de contraintes d'intégrité. Nous montrons ensuite, section 8.3, que les opérateurs de révision commutative peuvent être définis à partir des opérateurs de fusion contrainte. Enfin, section 8.4, nous abordons les rapports entre fusion de base de connaissance et théorie du choix social. Cette théorie a comme objet d'étude l'agrégation de préférences individuelles en une préférence sociale. K. J. Arrow a prouvé un théorème d'impossibilité montrant qu'il n'existe pas de "bonne" fonction d'agrégation. Nous étudions comment les opérateurs de fusion échappent à ce théorème d'impossibilité.

Le chapitre 9 s'intéresse à ce que nous avons nommé la fusion syntaxique. Ces opérateurs particuliers, souvent appelés opérateurs de combinaison dans la littérature, effectuent la fusion de bases de connaissance de la façon suivante : tout d'abord ils réalisent l'union de toutes les bases de connaissance à fusionner, puis ils sélectionnent certains sous-ensembles maximaux consistants de formules de cette union comme résultat de la fusion. Nous étudions les propriétés logiques de ces opérateurs, particulièrement ceux proposés par Baral *et al.*, et nous montrons qu'ils ne sont pas satisfaisants pour fusionner des bases de connaissances. Nous proposons ensuite l'utilisation de fonctions de sélection, inspirées de celles utilisées pour la révision, pour améliorer les propriétés de ces opérateurs.

Préliminaires

Ensembles, relations et pré-ordres

Soit E un ensemble, on note $\mathcal{P}(E)$ l'ensemble des sous-ensembles de E . Une relation binaire sur E est un sous-ensemble de $E \times E$, i.e. un ensemble de couples (x,y) avec $x,y \in E$.

Une relation binaire $\mathcal{R}$, définie sur $E \times E$, est dite :

- *réflexive* si pour tout x de E , $x\mathcal{R}x$.
- *transitive* si pour tout x,y,z de E , si $x\mathcal{R}y$ et $y\mathcal{R}z$, alors $x\mathcal{R}z$.
- *totale* si pour tout x,y de E , on a $x\mathcal{R}y$ ou $y\mathcal{R}x$.
- *symétrique* si pour tout x,y de E , si $x\mathcal{R}y$ alors $y\mathcal{R}x$.
- *anti-symétrique* si pour tout x,y de E , si $x\mathcal{R}y$ et $y\mathcal{R}x$ alors $x = y$.
- *modulaire* si pour tout x,y,z de E , si $x\mathcal{R}y$, $y\mathcal{R}x$ et $z\mathcal{R}x$, alors $z\mathcal{R}y$.
- *acyclique* si pour tout $x_1, \dots, x_n$ de E , on n'a pas $x_1 \mathcal{R} x_2 \mathcal{R} \dots \mathcal{R} x_n \mathcal{R} x_1$.

Une relation qui n'est pas totale est dite *partielle*. Un *pré-ordre* est une relation réflexive et transitive sur $E \times E$. Une *relation d'équivalence* est une relation réflexive, transitive et symétrique sur $E \times E$. Un *ordre* est une relation réflexive, anti-symétrique et transitive sur $E \times E$. Un *ordre strict* est une relation irreflexive et transitive sur $E \times E$. Soit un pré-ordre $\leq$ défini sur $E \times E$, on définit l'ordre strict $<$ associé, comme $x < y$ si $x \leq y$ et $y \not\leq x$. On définit également la relation d'équivalence $\simeq$ induite par $\leq$, comme $x \simeq y$ si $x \leq y$ et $y \leq x$.

Soient un pré-ordre $\leq$ défini sur $E \times E$ et E' un sous-ensemble de E , on note $\min(E', \leq)$ l'ensemble des éléments minimaux de E' pour $\leq$, i.e. $\min(E', \leq) = \{x \in E' : \nexists y \in E' y < x\}$.

Un pré-ordre défini sur E est *bien fondé* si chaque sous-ensemble de E admet un élément minimal.

Logique

Soit $\mathcal{L}$ un langage comprenant un ensemble d'atomes $\mathcal{A} = \{A, B, C, \dots\}$ et les connecteurs usuels $\neg$ (négation), $\wedge$ (conjonction), $\vee$ (disjonction), $\rightarrow$ (implication) et $\leftrightarrow$ (équivalence). $\perp$ dénote la contradiction et $\top$ la tautologie.

Nous utiliserons une opération de conséquence au sens de Tarski [Tar56] :

Définition 1 Une opération de conséquence sur un langage $\mathcal{L}$ est une fonction $Cn : \mathcal{P}(\mathcal{L}) \rightarrow \mathcal{P}(\mathcal{L})$ vérifiant les conditions :

1. $A \subseteq Cn(A)$ (inclusion)

- 2. Si $A \subseteq B$, alors $Cn(A) \subseteq Cn(B)$ (monotonie)
- 3. $Cn(A) = Cn(Cn(A))$ (idempotence)

Dans la suite du manuscrit on supposera que la relation de conséquence Cn vérifie également les propriétés suivantes :

- 1. Cn contient les conséquences logiques classiques (supraclassicalité)
- 2. Si $\alpha \in Cn(A)$ alors il existe A' un sous-ensemble fini de A tel que $\alpha \in Cn(A')$ (compacité)
- 3. $\beta \in Cn(A \cup \{\alpha\})$ si et seulement si $\alpha \rightarrow \beta \in Cn(A)$ (déduction)

On définit alors $A \vdash \alpha$ comme une notation pour $\alpha \in Cn(A)$.

Définition 2 Une base de connaissance K est un ensemble de propositions de $\mathcal{L}$. Une base de connaissance est dite close déductivement si $K = Cn(K)$.

Lorsque cela ne sera pas précisé on supposera que les bases de connaissance sont closes déductivement. Dans ce cas les bases de connaissance sont également appelées *théories*. On notera $K_\perp$ la base de connaissance triviale, i.e. celle contenant toutes les formules et $K_\top$ la base de connaissance qui ne contient que les tautologies. $\mathcal{K}_\mathcal{L}$ représente l'ensemble des bases de connaissance définies sur $\mathcal{L}$. Lorsqu'il n'y a pas d'ambiguïté on note simplement $\mathcal{K}$.

Lorsque l'on travaille en logique propositionnelle finie, une base de connaissance K est équivalente à la formule φ qui est la conjonction des formules de K .

On appelle *interprétation* une fonction de $\mathcal{A}$ dans $\{0,1\}$. On notera $\mathcal{W}$ l'ensemble des interprétations de $\mathcal{L}$.

Un *modèle* I d'une formule propositionnelle φ est une interprétation qui rend φ Vrai au sens usuel, ce que l'on note $I \models \varphi$. Un *contre-modèle* d'une formule est une interprétation qui n'est pas modèle de cette formule. On note $Mod(\varphi)$ l'ensemble des modèles de φ , i.e. $Mod(\varphi) = \{I \in \mathcal{W} : I \models \varphi\}$. Une formule est *consistante* si elle admet au moins un modèle. Soit M un ensemble d'interprétations, on note $\varphi_{\{M\}}$ la formule (à équivalence logique près) qui a M comme ensemble de modèles.

Première partie

Révision

Chapitre 1

Le cadre AGM

Pour atteindre la connaissance, ajoute des choses chaque jour.
Pour atteindre la sagesse, retire des choses chaque jour.
(Lao Tzu, Tao-te Ching, ch. 48)

La question soulevée par la révision de la connaissance est simple : étant données une connaissance du monde et une nouvelle information, comment incorporer cette nouvelle information dans nos connaissances tout en restant cohérent ? Si la question est simple, la réponse l'est moins...

L'étude de la dynamique des connaissances est nécessaire pour pouvoir utiliser et gérer des connaissances qui par nature sont incertaines et/ou incomplètes.

Le premier système utilisant de telles connaissances que l'on peut citer n'est autre que l'être humain ! Dans la vie de tous les jours, nous utilisons souvent la révision pour "corriger" notre connaissance du monde. Cela est dû au fait qu'en l'absence d'informations précises, nous nous reposons sur des hypothèses qui sont souvent contrariées par des évidences issues de notre observation du monde réel. En fait, la révision est un mécanisme à part entière de nos processus cognitifs.

En informatique ce problème est également de première importance dans tous les domaines où il est nécessaire de raisonner en présence d'incertitudes et, principalement, dans les domaines des bases de données et de l'intelligence artificielle.

En intelligence artificielle, le problème de la révision de la connaissance est fortement lié à celui de l'inférence non monotone. C'est un point clef pour définir des systèmes capables de raisonner en présence d'informations incomplètes et de mener des raisonnements de sens commun.

Dans le cadre logique c'est le paradigme AGM qui s'est imposé. Au lieu de proposer des opérateurs particuliers comme cela a été le cas dans de nombreux travaux [FUV83, FKUV86, Bor85, Dal88a, Win88], Carlos Alchourrón, Peter Gärdenfors et David Makinson ont proposé un ensemble de propriétés logiques que les opérateurs de révision "sensés" doivent satisfaire. Ils ont également montré que les opérateurs satisfaisant ces propriétés pouvaient être construits à partir de méthodes de révision très naturelles. D'ailleurs, depuis, de nombreux théorèmes de représentation ont montré l'équivalence, ou tout au moins les rapports étroits, entre cette caractérisation logique et des méthodes de révision, ou d'autres domaines comme les relations

d'inférence non monotones par exemple. C'est cette accumulation de résultats qui a renforcé la caractérisation logique en un véritable cadre de travail.

Bien sûr cette caractérisation n'est pas parfaite et ne peut être appliquée à tous les opérateurs de changement. De nombreuses critiques ont été adressées à l'encontre du cadre AGM. Ces critiques ne remettent pas en compte le cadre en lui-même, sa robustesse étant prouvée par les théorèmes de représentation. Mais elles suggèrent des adaptations ou des généralisations de ce cadre.

1.1 Expansion-Révision-Contraction

Nous allons présenter dans cette section le cadre AGM (Alchourrón, Gärdenfors, Makinson) [AGM85, Gär88] de la révision de la connaissance. Après avoir donné quelques définitions, nous nous intéresserons aux trois types d'opérateurs de changement de la connaissance, à savoir les opérateurs d'expansion, de révision et de contraction. Nous donnerons la caractérisation logique de chacune de ces familles au travers de propriétés logiques. Nous verrons enfin les relations étroites entre ces différents types d'opérateurs.

Les postulats AGM sont donnés sans grande exigence sur la nature des bases de connaissance, on suppose simplement que les bases de connaissance sont des ensembles de formules exprimées dans un certain langage et clos pour une relation de conséquence.

Etant données une base de connaissance K codant les connaissances d'un agent et une information A , il y a trois attitudes possibles pour K vis à vis de A :

- $A \in K$, i.e. l'information est dans la base : on dit que A est *acceptée*.
- $\neg A \in K$, i.e. la négation de l'information est dans la base : on dit que A est *refusée*.
- $A \notin K$ et $\neg A \notin K$, i.e. ni la nouvelle information, ni sa négation ne sont dans la base : on dit que A est *indéterminée* (ou *contingente*).

Gärdenfors [Gär88] définit alors les opérateurs de changement comme des changements d'attitude envers une information. Il y a donc 6 changements d'attitude possibles mais, pour des raisons de symétrie, ils se regroupent en trois catégories :

- lorsque l'on passe de A est indéterminée à A est acceptée ou A est refusée (i.e. $\neg A$ est acceptée), on effectue une *expansion*. Cela consiste à ajouter une information à la base de connaissance sans retirer aucune autre information.
- lorsque l'on passe de A est acceptée ou A est refusée à A est indéterminée, on effectue une *contraction*, c'est-à-dire que l'on "supprime" la connaissance à propos de A .
- enfin lorsque l'on passe de A est acceptée à A est refusée, ou symétriquement de A est refusée à A est acceptée, on effectue une *révision* : la croyance en A est totalement remise en question.

Alchourrón, Gärdenfors et Makinson ont proposé une série de postulats qui caractérisent ces trois types d'opérateurs de changement de la connaissance. Ces postulats ont pour but de capturer les propriétés de rationalité que l'on peut attendre de ces changements et sont connus sous le nom de postulats AGM.

1.1.1 Expansion

Typiquement, l'expansion est l'opération qui permet d'ajouter une information à une base de connaissance lorsque celles-ci sont compatibles. Si K est la base de connaissance de départ, l'expansion de K par A est notée $K + A$. Un opérateur d'*expansion* est une fonction $+$ de $\mathcal{K} \times \mathcal{L}$ vers $\mathcal{K}$ qui vérifie les propriétés suivantes :

- | | |
|--|--------------|
| (K+1) $K + A$ est une théorie | (clôture) |
| (K+2) $A \in K + A$ | (succès) |
| (K+3) $K \subseteq K + A$ | (inclusion) |
| (K+4) Si $A \in K$, alors $K + A = K$ | (vacuité) |
| (K+5) Si $K \subseteq H$, alors $K + A \subseteq H + A$ | (monotonie) |
| (K+6) $K + A$ est la plus petite base de connaissance satisfaisant (K+1)-(K+5) | (minimalité) |

L'explication intuitive de ces axiomes est la suivante : (K+1) assure que le résultat de l'expansion est bien une théorie. (K+2) dit que la nouvelle information doit être vraie dans la nouvelle base de connaissance. La motivation du nom expansion peut être expliquée par (K+3) qui certifie que l'on garde toutes les informations de l'ancienne base. (K+4) dit que si la nouvelle information appartient déjà à la base de connaissance alors il n'y a rien à faire pour l'accepter. Le postulat (K+5) exprime la monotonie de l'expansion. Et le dernier postulat (K+6) exprime la minimalité du changement, c'est-à-dire qu'il s'assure que la nouvelle base de connaissance ne contient pas de connaissance non justifiée par l'ajout de la nouvelle information.

Trois conséquences intéressantes de ces postulats sont :

- (1) $K + A = K + B$ ssi $B \in K + A$ et $A \in K + B$

Cette propriété exprime une caractérisation de l'équivalence entre deux résultats d'expansion. Si une information A est dans le résultat de l'expansion de la base de connaissance par B et symétriquement si B est une conséquence de l'expansion par A , alors les deux expansions donnent des bases équivalentes.

- (2) $(K + A) + B = (K + B) + A$

La propriété (2) exprime la commutativité de l'expansion, c'est-à-dire que l'ordre dans lequel on incorpore les nouvelles informations n'influe pas sur le résultat final.

- (3) Si $\neg A \in K$, alors $K + A = K_{\perp}$

La propriété (3) dit que si l'on effectue une expansion d'une base avec une information qui n'est pas cohérente avec elle, alors le résultat est la base triviale. Une autre remarque intéressante est que si l'on atteint la base triviale, il n'y a aucun moyen d'en sortir en utilisant l'expansion.

Il n'y a qu'un opérateur satisfaisant ces propriétés. En effet, un opérateur qui satisfait les postulats de l'expansion donne comme résultat une base de connaissance qui est l'ensemble des conséquences de la conjonction de la base de connaissance et de la nouvelle information :

Théorème 1 La fonction d'expansion $+$ satisfait les postulats (K+1)-(K+6) ssi $K + A = Cn(K \cup A)$.

1.1.2 Révision

Lorsque la nouvelle information contredit la base de connaissance, on ne peut pas utiliser l'expansion pour incorporer celle-ci, sous peine de rendre la base de connaissance inconsistante. Il faut alors abandonner certaines connaissances pour maintenir la consistance de la base. Il est donc évident que ce type de changement ne sera pas monotone (au sens de (K+5)).

Les principales propriétés de rationalité que l'on peut attendre d'un opérateur de révision sont les suivantes :

- La première est évidemment que la nouvelle base de connaissance soit cohérente.
- La seconde est la primauté de la nouvelle information (*primacy of update*), c'est-à-dire que la nouvelle information doit être vraie dans la nouvelle base de connaissance.
- La troisième est la minimalité du changement, c'est-à-dire que la révision doit conserver le maximum de connaissances de l'ancienne base. Ce principe de conservation a été annoncé de la manière suivante dans [Har86] p. 46 :

“When changing beliefs in response to new evidence, you should continue to believe as many of the old beliefs as possible.”

Alchourrón, Gärdenfors et Makinson ont proposé une série de postulats qui tentent de capturer logiquement ces propriétés.

Un opérateur de *révision* $*$ est une fonction de $\mathcal{K} \times \mathcal{L}$ vers $\mathcal{K}$ qui, à une base de connaissance K et une nouvelle information A , associe une nouvelle base de connaissance $K * A$ qui vérifie les propriétés suivantes :

- | | |
|---|-------------------------|
| (K*1) $K * A$ est une théorie | (clôture) |
| (K*2) $A \in K * A$ | (succès) |
| (K*3) $K * A \subseteq K + A$ | (inclusion) |
| (K*4) Si $\neg A \notin K$, alors $K + A \subseteq K * A$ | (vacuité) |
| (K*5) $K * A = K_{\perp}$ ssi $\vdash \neg A$ | (consistance) |
| (K*6) Si $A \leftrightarrow B$, alors $K * A = K * B$ | (extensionnalité) |
| (K*7) $K * (A \wedge B) \subseteq (K * A) + B$ | (inclusion conjonctive) |
| (K*8) Si $\neg B \notin K * A$, alors $(K * A) + B \subseteq K * (A \wedge B)$ | (vacuité conjonctive) |

L'interprétation de ces postulats est la suivante: (K*1) s'assure que le résultat de la révision est bien une théorie. (K*2) dit que la nouvelle information est vraie dans la nouvelle base de connaissance. Le postulat (K*3) implique que la révision par la nouvelle information ne peut pas ajouter de connaissance qui ne soit une conséquence de la nouvelle information et de la base de connaissance. Et les postulats (K*3) et (K*4) ensemble signifient que, lorsque la nouvelle information n'est pas contradictoire avec l'ancienne base de connaissance, alors la révision de la base de connaissance se résume à l'expansion de cette base. (K*5) exprime le fait que la seule façon d'arriver à une base inconsistante par une révision est de réviser par une information contradictoire. (K*6) dit que le résultat de la révision ne dépend pas de la syntaxe de la nouvelle information.

Ces six postulats sont les postulats de base pour les opérateurs de révision, les deux postulats (K*7) et (K*8) ont été appelés postulats supplémentaires par Gärdenfors et expriment le bon comportement des opérateurs de révision en terme de minimalité de changement. Ils

assurent que la révision par une conjonction de deux informations revient à une révision par la première information et une expansion par la seconde dès que cela est possible (i.e. dès que la seconde information ne contredit aucune connaissance issue de la première révision).

Katsuno et Mendelzon [KM91b] ont proposé une formulation équivalente des postulats AGM lorsque les bases de connaissance sont exprimées dans un langage propositionnel fini.

Soient deux formules propositionnelles φ et μ , φ dénotant la base de connaissance et μ la nouvelle information. $\varphi \circ \mu$ dénote la formule résultat de la révision de φ par μ . L'opérateur $\circ$ est un opérateur de révision s'il vérifie les postulats suivants :

- (R1) $\varphi \circ \mu \vdash \mu$
- (R2) Si $\varphi \wedge \mu$ est consistant alors $\varphi \circ \mu \leftrightarrow \varphi \wedge \mu$
- (R3) Si μ est consistant alors $\varphi \circ \mu$ est consistant
- (R4) Si $\varphi_1 \leftrightarrow \varphi_2$ et $\mu_1 \leftrightarrow \mu_2$ alors $\varphi_1 \circ \mu_1 \leftrightarrow \varphi_2 \circ \mu_2$
- (R5) $(\varphi \circ \mu) \wedge \phi \vdash \varphi \circ (\mu \wedge \phi)$
- (R6) Si $(\varphi \circ \mu) \wedge \phi$ est consistant alors $\varphi \circ (\mu \wedge \phi) \vdash (\varphi \circ \mu) \wedge \phi$

Soit un opérateur de révision $*$ sur des théories et $\circ$ un opérateur de révision sur des bases de connaissance propositionnelles. Si $Cn(\varphi) = K$, on dit que l'opérateur $*$ correspond à l'opérateur $\circ$ si $K * A = Cn(\varphi \circ A)$.

Théorème 2 *Soit un opérateur de révision $*$ et son opérateur $\circ$ correspondant. Alors $*$ satisfait les postulats $(K * 1) - (K * 8)$ si et seulement si $\circ$ vérifie les postulats (R1) – (R6).*

1.1.3 Contraction

On nomme contraction le type de changement intervenant lorsque l'on rétracte une information d'une base de connaissance mais qu'aucune nouvelle information n'est ajoutée. Typiquement, lors d'une contraction, on passe donc de l'attitude "A est acceptée" ou "A est refusée" à "A est indéterminée". On peut donner un exemple d'utilisation de la contraction lorsque l'on mène des raisonnements hypothétiques. On utilise l'expansion pour se rendre compte des conséquences qu'aurait tel ou tel fait sur nos connaissances. Puis on utilise la contraction pour retrouver la base de connaissance initiale.

Dans ce cas nous sommes confrontés aux mêmes exigences de minimalité de changement que pour la révision, c'est-à-dire que, lorsque l'on contracte par une information, on ne veut supprimer de la base de connaissance que ce qui est nécessaire pour ne plus impliquer l'information.

Un opérateur de *contraction* $\div$ est une fonction de $\mathcal{K} \times \mathcal{L}$ vers $\mathcal{K}$, qui rejette l'information A de la base de connaissance K et a comme résultat la base de connaissance $K \div A$ qui vérifie les propriétés suivantes :

- | | |
|--|-------------|
| (K÷1) $K \div A$ est une théorie | (clôture) |
| (K÷2) $K \div A \subseteq K$ | (inclusion) |
| (K÷3) Si $A \notin K$, alors $K \div A = K$ | (vacuité) |
| (K÷4) Si $\not\models A$, alors $A \notin K \div A$ | (succès) |

- (K÷5) Si $A \in K$, alors $K \subseteq (K \div A) + A$ (restauration)
 (K÷6) Si $A \leftrightarrow B$, alors $K \div A = K \div B$ (préservation)
 (K÷7) $(K \div A) \cap (K \div B) \subseteq K \div (A \wedge B)$ (intersection)
 (K÷8) Si $A \notin K \div (A \wedge B)$, alors $K \div (A \wedge B) \subseteq K \div A$ (conjonction)

(K÷1) assure que le résultat de la contraction est bien une théorie. (K÷2) garantit que lors de la contraction aucune nouvelle information n'est ajoutée à la base de connaissance. (K÷3) dit que si l'information A n'est pas acceptée par K , il n'y a rien à faire pour retirer A de K . Le postulat (K÷4) assure le succès de la contraction, c'est-à-dire que si A n'est pas une tautologie, alors la contraction réussira. Le postulat (K÷5) assure que la contraction de K par A suivie de l'expansion par A redonne la théorie K comme résultat (l'inclusion inverse de (K÷5) étant une conséquence de (K÷1)-(K÷4)). (K÷6) dit que le résultat de la contraction ne dépend pas de la syntaxe de l'information.

Ces six postulats sont les postulats de base pour les opérateurs de contraction. (K÷7) et (K÷8) sont appelés postulats supplémentaires pour la contraction. (K÷7) dit que si une information est à la fois dans la contraction par A et dans la contraction par B alors elle doit être dans la contraction par la conjonction $A \wedge B$. (K÷8) exprime la minimalité du changement pour la conjonction.

1.1.4 Identités

L'expansion permet d'ajouter une nouvelle information à la base de connaissance, la contraction permet de retirer une information de la base, et la révision permet de modifier une information. Il semblerait alors possible d'exprimer une révision comme étant une contraction suivie d'une expansion [Gär88] :

$$K * A = (K \div \neg A) + A \quad (\text{Identité de Levi})$$

On a donc une définition de la révision en fonction de la contraction et de l'expansion. Le théorème suivant montre que les opérateurs définis grâce à l'identité de Levi sont bien des opérateurs de révision.

Théorème 3 *Si l'opérateur de contraction $\div$ satisfait (K÷1)-(K÷4) et (K÷6) et l'opérateur d'expansion $+$ satisfait (K+1)-(K+6), alors l'opérateur de révision $*$ défini par l'identité de Levi satisfait (K*1)-(K*6). De plus, si (K÷7) est satisfait, alors (K*7) est satisfait pour la révision ainsi définie, et si (K÷8) est satisfait, alors (K*8) est satisfait pour la révision ainsi définie.*

Remarque 1 *Il est important de noter que (K÷5) n'intervient pas dans ce résultat. En effet ce postulat, restauration, est le plus critiqué en ce qui concerne la contraction (cf section 1.3.1). Mais les réserves émises sur ce postulat n'affectent donc en rien les opérateurs de révision.*

De la même manière, on a la correspondance inverse, c'est-à-dire que si on dispose d'un opérateur de révision, il est possible de définir un opérateur de contraction. Cette définition est connue sous le nom d'identité de Harper.

$$K \div A = K \cap (K * \neg A) \quad (\text{Identité de Harper})$$

L'idée assez naturelle de cette identité est que les informations présentes dans la contraction de K par A sont celles n'ayant rien à voir avec la véracité de A , i.e. celles présentes à la fois dans K et dans la révision de K par $\neg A$.

Théorème 4 *Si l'opérateur de révision $*$ satisfait (K^*1) - (K^*6) , alors l'opérateur de contraction $\div$ définie par l'identité de Harper satisfait $(K\div 1)$ - $(K\div 6)$. De plus, si (K^*7) est satisfait, alors $(K\div 7)$ est satisfait pour la contraction ainsi définie, et si (K^*8) est satisfait, alors $(K\div 8)$ est satisfait pour la contraction ainsi définie.*

Le point intéressant à remarquer est que, bien que les postulats caractérisant ces changements soient indépendants (on ne se réfère pas à la contraction dans les postulats de la révision et vice-versa), il existe une véritable correspondance entre révision et contraction.

Il est donc possible de définir un opérateur de révision à partir d'un opérateur de contraction et vice-versa. On peut donc étudier indifféremment les propriétés de l'un ou de l'autre.

1.2 Théorèmes de représentation

La caractérisation de la révision (et de la contraction) donnée à la section précédente n'est qu'un ensemble de propriétés logiques qui tentent de capturer les conditions de rationalité que l'on peut attendre de ces opérateurs. En tant que telle, cette caractérisation n'est donc qu'un ensemble de postulats et sa justification peut donc sembler un peu faible. Pourquoi cet ensemble de postulats se justifierait-il plus qu'un autre ?

Il se trouve que les opérateurs obéissant à la caractérisation donnée par C. Alchourrón, P. Gärdenfors et D. Makinson correspondent à des méthodes de révision (contraction) naturelles. Ces méthodes sont basées sur des "ordres de préférences" sur les croyances, dénotant qu'aux yeux de l'agent une proposition est plus crédible qu'une autre. Et ce sont les croyances les plus crédibles qui sont sélectionnées comme nouvelle base de connaissance. La différence entre les opérateurs réside dans la "métrique" choisie pour construire l'ordre de préférences.

Ces résultats d'équivalence entre opérateurs de révision AGM et ces méthodes de révision, donnent des définitions *constructives* des opérateurs de révision et c'est pour cela qu'on les nomme théorèmes de représentation.

Les trois premières méthodes présentées, les *fonctions de contraction par intersection partielle*, les *enracinements épistémiques* et la *contraction sûre*, reposent sur l'existence d'un pré-ordre entre formules qui marque la crédibilité relative de celles-ci et sert à déterminer les informations les moins importantes de la base de connaissance qui pourront être éliminées pour maintenir la cohérence lors de la révision.

Ensuite une méthode de révision basée sur un pré-ordre sur les mondes possibles (*Assignment Fidèle* ou *système de sphères*) exprime la révision comme une sélection des modèles minimaux de la nouvelle information suivant cette mesure de confiance sur les mondes.

Enfin, les deux derniers résultats ne sont pas des définitions d'opérateurs de révision mais ils montrent le rapport étroit entre le domaine de la révision AGM et d'autres domaines "indépendants".

Premièrement il existe une correspondance entre opérateurs de révision et certaines relations d'inférence non monotones. Cette correspondance a fait dire à Gärdenfors que relations d'inférence et opérateurs de révision n'étaient que les deux faces d'une même pièce de monnaie.

Ensuite on peut montrer également le lien étroit entre opérateurs de révision et la logique possibiliste. Cette logique permet de raisonner qualitativement en présence d'incertitude et/ou d'imprécision.

C'est cet ensemble de correspondances, illustré figure 1.1, entre révision AGM, diverses méthodes de révision et d'autres domaines du raisonnement en présence d'incertitude qui montre la robustesse de cette caractérisation logique. Puisque remettre en cause la caractérisation proposée par Alchourrón, Gärdenfors et Makinson revient alors à remettre en cause la rationalité de ces différentes méthodes et le bien fondé des domaines connexes.

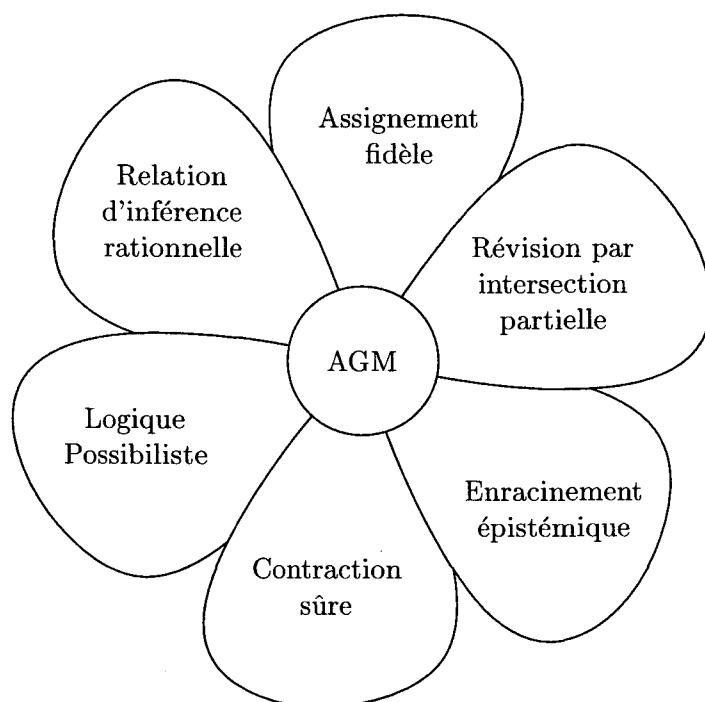


FIG. 1.1 – *Théorèmes de représentation*

1.2.1 Contraction par intersection partielle

Nous allons définir dans cette section des opérateurs de contraction à partir d'ensembles maximaux consistants de formules.

On définit d'abord l'ensemble $K \perp A$, qui est l'ensemble des sous-théories maximales de K qui n'impliquent pas A .

Définition 3 Soient une théorie K et une proposition A . $K \perp A$ est l'ensemble de tous les K'

qui vérifient :

1. $K' \subseteq K$
2. $K' \not\models A$
3. Si $K' \subsetneq K'' \subseteq K$ alors $K'' \vdash A$

Une façon naturelle de définir la contraction est alors de considérer l'intersection de toutes les sous-théories maximales de K qui n'impliquent pas A .

Définition 4 La fonction de contraction par intersection totale (*full meet contraction*) $\div_f$ est définie comme

$$K \div_f A = \begin{cases} Cn(\bigcap(K \perp A)) & \text{si } K \perp A \text{ n'est pas vide, et} \\ K \div_f A = K & \text{sinon} \end{cases}$$

Le problème est que cette contraction est trop drastique : lorsque l'on effectue une contraction par A , le résultat est l'ensemble des formules de K qui sont conséquences logiques de $\neg A$. Cela pose un problème pour la révision puisque l'on a le résultat suivant [AM82] :

Théorème 5 Si une fonction de révision $*$ est définie à partir d'une fonction de contraction par intersection totale au moyen de l'identité de Levi, alors pour chaque proposition A telle que $\neg A \in K$, on a $K * A = Cn(A)$.

Ce résultat indique donc que, pour chaque révision sévère (i.e. par une information inconsistante avec la base de connaissance) que l'on effectuera, cela "effacera" l'ancienne connaissance. Ce résultat n'est donc pas intuitivement suffisant pour l'opérateur de révision correspondant par l'identité de Levi. La fonction de contraction par intersection totale doit plutôt être considérée comme une limite inférieure pour les fonctions de contraction. C'est-à-dire que chaque opérateur de contraction raisonnable doit contenir le résultat de cette dernière.

Une solution pour obtenir des opérateurs moins drastiques est de ne pas considérer l'ensemble de toutes les sous-théories maximales, mais simplement certaines d'entre elles. Une fonction de sélection choisit alors certaines de ces sous-théories. Intuitivement, il faut interpréter cette sélection comme le choix des "meilleures" sous-théories pour un certain critère. Ceci motive les définitions suivantes :

Définition 5 Soit une théorie K , une fonction de sélection γ est une fonction qui associe à chaque proposition A l'ensemble $\gamma(K \perp A)$, qui est un sous-ensemble non vide de $K \perp A$ si celui-ci n'est pas vide et $\gamma(K \perp A) = \{K\}$ sinon.

Définition 6 Une fonction de contraction par intersection partielle (*partial meet contraction*) $\div$ est définie comme

$$K \div A = Cn\left(\bigcap \gamma(K \perp A)\right)$$

On peut noter que la fonction de contraction par intersection totale est un cas particulier de fonction de contraction par intersection partielle (lorsque la fonction de sélection garde $K \perp A$ en entier). L'autre cas limite, lorsque la fonction de sélection ne garde qu'un seul élément de l'ensemble $K \perp A$, est appelé fonction de contraction à *choix maximal* (*maxichoice*) [AM82].

Les résultats suivants montrent la correspondance entre opérateurs de contraction par intersection partielle et la caractérisation AGM [AGM85].

Théorème 6 *Soit un opérateur $\div$, $\div$ est une fonction de contraction par intersection partielle si et seulement si $\div$ satisfait les postulats $(K \div 1)$ - $(K \div 6)$.*

Définition 7 *Une fonction de sélection γ est relationnelle si et seulement si pour tout K il existe une relation $\leq$ sur K^2 telle que*

$$\gamma(K \perp A) = \{K' \in K \perp A \mid K' \leq K'', \forall K'' \in K \perp A\}$$

Si $\leq$ est une relation transitive alors γ est dite relationnelle transitive (transitively relational).

Une fonction de contraction par intersection partielle est relationnelle (respectivement relationnelle transitive) si et seulement si elle est définie à partir d'une fonction de sélection qui est relationnelle (respectivement relationnelle transitive).

Théorème 7 *Soit un opérateur $\div$, $\div$ est une fonction de contraction par intersection partielle relationnelle transitive si et seulement si $\div$ satisfait les postulats $(K \div 1)$ - $(K \div 8)$.*

Un résultat intéressant est que dans ce cas la fonction est totale (Une fonction de sélection est totale si elle est définie à partir d'une relation totale).

Théorème 8 *Soit une fonction de contraction par intersection partielle $\div$, $\div$ est relationnelle transitive si et seulement si elle est transitive et totale.*

Le fait qu'un opérateur de contraction AGM (et donc par dualité un opérateur de révision AGM) corresponde à une famille de pré-ordres totaux, résultat qui transparaît au travers des théorèmes ci-dessus, est le point clef qui permet la définition des théorèmes de représentation présentés dans cette section. En effet, ces théorèmes montrent que les opérateurs vérifiant les postulats AGM correspondent à des méthodes de changement basés sur les préférences (les pré-ordres totaux) de l'agent.

1.2.2 Enracinements épistémiques

La principale préoccupation, lorsque l'on effectue une révision (ou une contraction), est de sélectionner les formules que l'on peut enlever de la base de connaissance. Un principe évident de rationalité nous demande de retirer les formules les moins importantes. Il faut donc construire une sorte d'ordre entre les formules suivant leur importance vis à vis de la base. On appelle un tel ordre enracinement épistémique (epistemic entrenchment). En fait, on peut voir la base de connaissance comme un ensemble de "couches", plus une formule appartient à une couche profonde, plus elle est importante. Lorsque l'on a un changement à effectuer à l'intérieur de la base, ce sont les formules appartenant aux couches les plus superficielles que l'on enlève de la base.

Soient deux formules A et B , la notation $A \leq B$ signifie " B est au moins aussi enraciné que A ". $\leq$ est un *enracinement épistémique* s'il satisfait les propriétés suivantes [Gär88]:

- | | |
|--|----------------|
| (EE1) Si $A \leq B$ et $B \leq C$, alors $A \leq C$ | (transitivité) |
| (EE2) Si $A \vdash B$, alors $A \leq B$ | (domination) |

- (EE3) $A \leq A \wedge B$ ou $B \leq A \wedge B$ (conjonction)
 (EE4) Si $K \neq K_\perp$, $A \notin K$ ssi $\forall B \ A \leq B$ (minimalité)
 (EE5) Si $B \leq A \ \forall B$, alors $\vdash A$ (maximalité)

(EE1) demande que la relation d'enracinement épistémique soit transitive. (EE2) dit que si une formule est logiquement plus forte qu'une autre alors elle est moins retranchée. (EE3) assure que pour abandonner la formule $A \wedge B$, il suffit d'abandonner A ou B . (EE4) exprime que les minimaux pour la relation sont les formules qui ne sont pas dans K . Inversement, les formules les plus retranchées sont les tautologies, ce que dit (EE5).

Un enracinement épistémique correspond à un opérateur de contraction [Gär88]:

Théorème 9 Une fonction de contraction $\div$ satisfait $(K \div 1) - (K \div 8)$ si et seulement si il existe $\leq$ satisfaisant (EE1)-(EE5), où $B \leq A$ ssi $B \notin K \div A \wedge B$.

Ce qui, grâce à l'identité de Levi, donne une équivalence entre enracinements épistémiques et opérateurs de révision. De manière plus constructive, on peut définir un opérateur de révision à partir d'un enracinement épistémique de la façon suivante:

Théorème 10 Soit un enracinement épistémique $\leq$. L'opérateur $*$ défini par

$$B \in K * A \text{ ssi soit } (A \rightarrow \neg B) < (A \rightarrow B), \text{ soit } \vdash \neg A$$

est un opérateur de révision AGM.

1.2.3 Contraction sûre

La méthode de contraction sûre a été proposée par Alchourrón et Makinson [AM85], et elle est basée sur la notion de sécurité des connaissances. Intuitivement, une connaissance est en sécurité lors d'une contraction si elle ne peut être blâmée d'impliquer l'information par laquelle on effectue la contraction.

On définit tout d'abord ce qu'est une sous-théorie minimale de K impliquant A .

Définition 8 Un ensemble K' est une sous-théorie minimale de K impliquant A si et seulement si

- i. $K' \subseteq K$
- ii. $K' \vdash A$
- iii. Si $K'' \subsetneq K'$, alors $K'' \not\vdash A$

Etant donnée une théorie K , on définit une hiérarchie sur les éléments de K à l'aide d'une relation acyclique:

Définition 9 Etant donnée une relation acyclique $<$ sur une théorie K , un élément B est en sécurité (safe) vis à vis de A si et seulement si B n'est pas un élément minimal (pour $<$) d'une sous-théorie minimale de K qui implique A .

Un élément B de K est en sécurité vis à vis de A s'il ne peut être blâmé d'impliquer A , c'est-à-dire qu'étant donnée une sous-théorie minimale de K impliquant A , soit cet élément n'appartient pas à cette sous-théorie, soit il y a un élément C de cette sous-théorie qui est plus "blâmable", i.e. $C < B$.

L'ensemble des éléments de K qui sont en sécurité vis à vis de A sont notés K/A .

Définition 10 La fonction de contraction sûre (safe contraction) d'une théorie K , étant donnée une hiérarchie $<$, est l'ensemble des conséquences de K/A , i.e. $K \div A = Cn(K/A)$.

Théorème 11 Une fonction de contraction sûre satisfait les postulats $(K \div 1)$ – $(K \div 6)$.

Nous n'avons pour l'instant pas exigé grand chose de la relation $<$. On peut se demander s'il est possible de satisfaire $(K \div 7)$ et $(K \div 8)$ en augmentant les conditions sur la relation $<$.

Nous allons définir quelques propriétés sur la relation $<$. Les deux premières disent en quelque sorte que la relation est "compatible" avec la relation de conséquence. La dernière est une version faible de la propriété de totalité.

Définition 11 On dit que $<$ continue $\vdash$ vers le haut (continues up) sur la théorie K si et seulement si pour tout $A, B, C \in K$, si $A < B$ et $B \vdash C$, alors $A < C$.

On dit que $<$ continue $\vdash$ vers le bas (continues down) sur la théorie K si et seulement si pour tout $A, B, C \in K$, si $A \vdash B$ et $B < C$, alors $A < C$.

On dit que $<$ est virtuellement totale (virtually connected) sur la théorie K si et seulement si pour tout $A, B, C \in K$, si $A < B$, alors $A < C$ ou $C < B$.

Théorème 12 Soit une théorie K . Une fonction de contraction sûre définie à partir d'une hiérarchie $<$ qui continue $\vdash$ vers le haut (ou vers le bas) sur K satisfait $(K \div 7)$.

Théorème 13 Soit une théorie K . Une fonction de contraction sûre définie à partir d'une hiérarchie $<$ qui est virtuellement totale et qui continue $\vdash$ vers le haut et vers le bas sur K satisfait $(K \div 8)$.

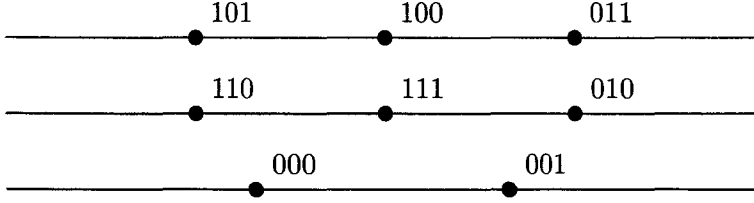
Alchourrón et Makinson [AM86] ont donné un théorème de représentation pour ces opérateurs dans le cas fini.

Théorème 14 Soit une théorie K . $\div$ est une fonction de contraction sûre définie à partir d'une hiérarchie $<$ qui continue $\vdash$ vers le haut et vers le bas sur K si et seulement si $\div$ est une fonction de contraction par intersection partielle relationnelle sur K .

Théorème 15 Soit une théorie K . $\div$ est une fonction de contraction sûre définie à partir d'une hiérarchie $<$ qui est virtuellement totale et qui continue $\vdash$ vers le haut et vers le bas sur K si et seulement si $\div$ est une fonction de contraction par intersection partielle relationnelle transitive sur K .

1.2.4 Assignement fidèle et système de sphères

Dans les sections précédentes, les relations de préférences associées aux opérateurs de révision à travers les théorèmes de représentation étaient des pré-ordres sur les formules (ou ensembles de formules). Katsuno et Mendelzon [KM91b] ont montré que les opérateurs de révision correspondaient également à des pré-ordres sur les interprétations. Cette idée peut

FIG. 1.2 – Exemple de pré-ordre associé à φ

être attribuée à Grove [Gro88] qui, avec son système de sphères, avait donné une sémantique similaire aux opérateurs de révision.

La proposition de Katsuno et Mendelzon a été faite dans le cadre propositionnel fini, mais peut être étendue à toutes bases de connaissance, puisqu'elle peut être vue comme un cas particulier du système de Sphères de Grove (cf [Gro88]).

Définition 12 *Un assignement fidèle est une fonction qui associe à chaque base de connaissance φ un pré-ordre $\leq_\varphi$ sur les interprétations tel que :*

1. Si $I \models \varphi$ et $J \models \varphi$, alors $I \simeq_\varphi J$
2. Si $I \models \varphi$ et $J \not\models \varphi$, alors $I <_\varphi J$
3. Si $\varphi_1 \leftrightarrow \varphi_2$, alors $\leq_{\varphi_1} = \leq_{\varphi_2}$

C'est-à-dire que les modèles de φ sont tous équivalents pour l'ordre associé et sont strictement préférés aux contre-modèles.

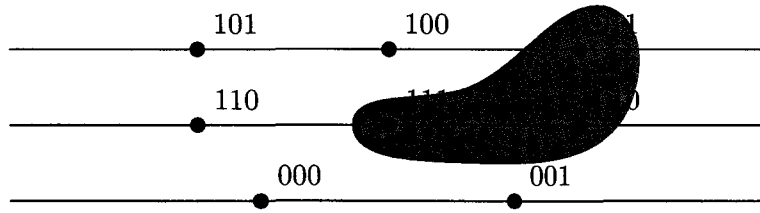
Intuitivement, l'ordre associé à une base de connaissance représente l'état épistémique de l'agent, c'est-à-dire ses connaissances (les modèles de la base) et ses croyances (préférences). Plus une interprétation est petite pour l'ordre, plus elle est préférée par l'agent, c'est-à-dire plus l'agent la considère crédible.

Lorsque l'on effectue une révision, ce sont alors les interprétations de la nouvelle information les plus crédibles pour l'état épistémique courant qui forment la nouvelle base de connaissance. Ceci est décrit formellement dans le théorème de représentation suivant :

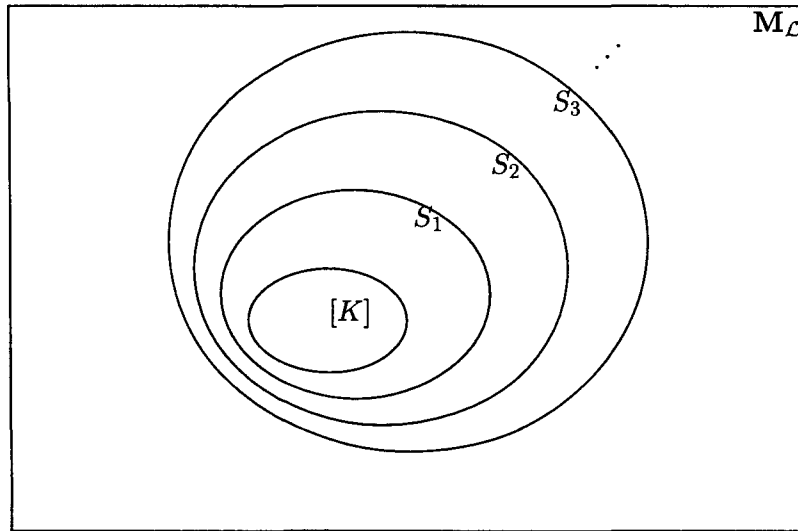
Théorème 16 *Un opérateur de révision $\circ$ satisfait les postulats (R1)-(R6) si et seulement si il existe un assignement fidèle qui associe à chaque base de connaissance φ un pré-ordre total $\leq_\varphi$ tel que $\text{Mod}(\varphi \circ \mu) = \min(\text{Mod}(\mu), \leq_\varphi)$.*

Supposons que l'on travaille avec un langage composé de trois variables propositionnelles a, b, c considérées dans cet ordre pour les valuations. On notera 100 l'interprétation (1,0,0), c'est-à-dire celle qui rend a à Vrai et b et c à Faux. Soient I et J deux interprétations. L'interprétation I est plus petite que J pour le pré-ordre associé à la base de connaissance φ (i.e. $I <_\varphi J$) si elle apparaît à un niveau inférieur. I et J sont équivalentes pour φ ($I \simeq_\varphi J$) si elles apparaissent à un même niveau. Considérons par exemple le pré-ordre de la figure 1.2.4. Ce pré-ordre est associé à la base de connaissance qui a comme modèles (0,0,0) et (0,0,1), c'est-à-dire à la formule (à équivalence logique près) $\varphi = \neg a \wedge \neg b$.

Ce pré-ordre représente les préférences de l'agent, puisque ici (0,1,0) par exemple apparaît à un niveau inférieur que (0,1,1), ce qui peut-être interprété comme le fait que l'agent trouve (0,1,0) plus crédible que (0,1,1) au vu de ses connaissances actuelles.

FIG. 1.3 – Révision de φ par μ

Supposons que l'on révisé φ par $\mu = \varphi_{\{(0,1,1),(0,1,0),(1,1,1)\}}$, comme représenté dans la figure 1.2.4, le résultat de la révision est, d'après le théorème de représentation, la formule (à équivalence logique près) dont l'ensemble des modèles est l'ensemble des modèles minimaux de la nouvelle information pour le pré-ordre associé à φ , donc ici $\varphi \circ \varphi_{\{(0,1,1),(0,1,0),(1,1,1)\}} = \varphi_{\{(0,1,0),(1,1,1)\}}$.

FIG. 1.4 – Système de sphères centré sur $[K]$

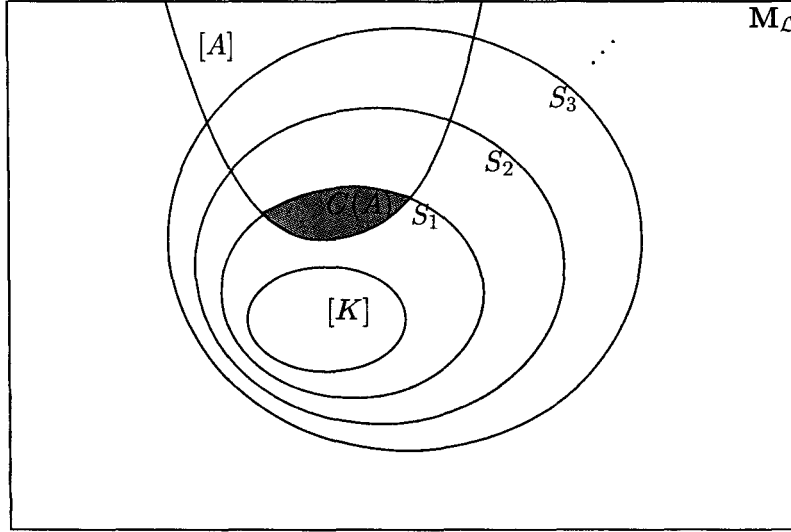
Grove a donné une sémantique similaire aux opérateurs de révision [Gro88] en termes de système de sphères sur les mondes possibles.

On appelle *monde possible* un sous-ensemble maximal consistant du langage (ce qui, dans le cas fini, s'identifie donc aux interprétations) et on note $\mathbf{M}_{\mathcal{L}}$ l'ensemble des mondes possibles du langage $\mathcal{L}$.

Une base de connaissance peut alors être représentée par l'ensemble des mondes possibles $[K] \subseteq \mathbf{M}_{\mathcal{L}}$ qui contiennent toutes les formules de K . Plus formellement :

Définition 13 Soit une base de connaissance K . Si $K = K_{\perp}$ alors $[K] = \emptyset$, sinon

$$[K] = \{M \in \mathbf{M}_{\mathcal{L}} : K \subseteq M\}$$

FIG. 1.5 – Système de sphères centré sur $[K]$: Révision par A

Ceci donne une véritable correspondance entre bases de connaissance et sous-ensembles de M_L puisque l'on peut également, à tout sous-ensemble S de M_L associer une base K_S composée de l'ensemble des formules présentes dans tous les mondes possibles de S : $K_S = \bigcap \{M : M \in S\}$.

Grove s'intéresse alors à un système de sphères centré sur $[K]$. Cette construction est inspirée de la sémantique proposée par Lewis pour les conditionnels [Lew73].

Définition 14 *Un système de sphères centré sur $[K]$ est une collection de sous-ensembles S de M_L qui vérifient les conditions suivantes :*

- (S1) *Si $S, S' \in S$, alors $S \subseteq S'$ ou $S' \subseteq S$*
- (S2) *$[K] \in S$*
- (S3) *Si $S \in S$, alors $[K] \subseteq S$*
- (S4) *$M_L \in S$*
- (S5) *Si A est une formule et si $[A]$ intersecte une sphère de S , alors il existe une sphère minimale qui intersecte $[A]$*

La condition (S1) dit que S est totalement ordonné par $\subseteq$. Les conditions (S2) et (S3) disent que $[K]$ est le plus petit élément de S (pour $\subseteq$). (S4) assure que le plus grand élément de S est l'ensemble de tous les mondes possibles. (S5) dit que si une sphère $[A]$ intersecte une sphère de S (ce qui arrive dès que $[A] \neq \emptyset$) alors il existe une sphère, notée S_A , intersectant $[A]$ et plus petite que toutes les sphères intersectant $[A]$. On note alors $C(A) = [A] \cap S_A$ les mondes possibles de cette sphère minimale satisfaisant A , c'est-à-dire les mondes les plus proches de $[K]$. Si $A = \perp$, alors $[A] = \emptyset$ et $C(A) = M_L$.

L'ensemble des mondes possibles $C(A)$ est alors le résultat de la révision de K par A comme le montre le théorème suivant :

Théorème 17 *Soit une base de connaissance K . Il existe un système de sphères S centré sur $[K]$ tel que pour toute formule A , $K * A = K_{C(A)}$ si et seulement si $*$ est un opérateur de révision satisfaisant $(K * 1) - (K * 8)$.*

Ce théorème de représentation peut donc être vu comme une généralisation du théorème de Katsuno et Mendelzón.

1.2.5 Révision et relations d'inférence

Nous allons voir dans cette section que le point commun entre opérateurs de révision et relations d'inférence non monotones ne se limite pas seulement au partage de cette propriété de non monotonie, mais qu'il existe une véritable correspondance entre ces opérateurs et les (bonnes) relations d'inférence. Comme l'a fait remarquer Gärdenfors [Gär90] :

“Belief revision and nonmonotonic logic are two sides of the same coin.”

Etudier les relations d'inférence permet donc de trouver des résultats sur les opérateurs de révision et vice-versa.

Relations d'inférence

La logique classique n'est pas adéquate pour modéliser le raisonnement de sens commun. Le problème réside en sa monotonie. De nombreuses logiques, affaiblissements ou surcharges de la logique classique, ont été proposées pour pallier ce “défaut”. Ces logiques ont donc été regroupées sous le nom de logiques non monotones. Mais le fait de ne pas satisfaire la propriété de monotonie ne suffit pas à caractériser une logique, et les logiques existantes sont loin d'avoir toutes le même comportement. On a donc étudié les propriétés que l'on pouvait attendre de ces différentes logiques, permettant ainsi d'avoir une catégorisation plus fine. Ces propriétés sont les suivantes [KLM90, LM92, Mak94] :

Définition 15 Une relation $\vdash$ est dite *préférentielle* si elle satisfait les six propriétés suivantes :

$$\begin{array}{ll}
 \text{REF} & \frac{}{\alpha \vdash \alpha} \\
 \text{LLE} & \frac{\alpha \vdash \beta \quad \vdash \alpha \leftrightarrow \gamma}{\gamma \vdash \beta} \\
 \text{RW} & \frac{\alpha \vdash \beta \quad \vdash \beta \rightarrow \gamma}{\alpha \vdash \gamma} \\
 \text{AND} & \frac{\alpha \vdash \beta \quad \alpha \vdash \gamma}{\alpha \vdash \beta \wedge \gamma} \\
 \text{OR} & \frac{\alpha \vdash \gamma \quad \beta \vdash \gamma}{\alpha \vee \beta \vdash \gamma} \\
 \text{CM} & \frac{\alpha \vdash \beta \quad \alpha \vdash \gamma}{\alpha \wedge \gamma \vdash \beta}
 \end{array}$$

Les règles ci-dessus sont la réflexivité (REF), l'équivalence logique à gauche (LLE), l'affaiblissement à droite (RW), le et (AND), le ou (OR), et la monotonie prudente (CM). Ces six règles sont également connues sous le nom de système P.

Définition 16 Une relation $\vdash$ est dite *rationnelle* si elle est préférentielle et satisfait la propriété suivante (monotonie rationnelle) :

$$\text{RM} \quad \frac{\alpha \vdash \beta \quad \alpha \not\vdash \neg \gamma}{\alpha \wedge \gamma \vdash \beta}$$

Il existe également des propriétés plus fortes que celle de monotonie rationnelle et moins forte que celle de monotonie (voir [BMP97, BP96]).

Les conditions posées sur les relations ne sont pas très fortes, mais elles sont déjà très structurantes puisqu'elles permettent de déduire les propriétés suivantes [KLM90, LM92, Mak94] :

Théorème 18 *Si $\vdash$ est une relation préférentielle alors $\vdash$ satisfait les propriétés suivantes :*

$$S \quad \frac{\alpha \wedge \beta \vdash \gamma}{\alpha \vdash \beta \rightarrow \gamma} \quad \text{CUT} \quad \frac{\alpha \wedge \beta \vdash \gamma \quad \alpha \vdash \beta}{\alpha \vdash \gamma}$$

Si $\vdash$ est une relation rationnelle alors $\vdash$ satisfait les propriétés suivantes :

$$DR \quad \frac{\alpha \vee \beta \vdash \gamma \quad \alpha \not\vdash \gamma}{\beta \vdash \gamma} \quad \text{NR} \quad \frac{\alpha \vdash \beta \quad \alpha \wedge \gamma \not\vdash \beta}{\alpha \wedge \neg \gamma \vdash \beta}$$

Les règles ci-dessus sont la règle de Shoham (S), la coupure (CUT), la disjonction rationnelle (DR) et la négation rationnelle (NR).

Les propriétés du système P semblent être les propriétés minimales exigibles d'une relation pour la considérer comme une relation d'inférence. Ajouter la règle de monotonie rationnelle permet également d'obtenir des propriétés intéressantes pour une relation d'inférence tout en restant non monotone. L'intérêt de ces deux familles de relations est qu'elles disposent toutes deux d'une sémantique claire.

Définition 17 *Une structure $\mathcal{M}$ est un triplet $\langle S, i, \prec \rangle$ où S est un ensemble d'objets quelconques (appelés états), $\prec$ est un ordre strict (i.e. une relation transitive et irréflexive) sur S et i est une fonction (la fonction d'interprétation) qui associe un monde à chaque état i.e. $i : S \rightarrow \mathcal{W}$.*

Soit une formule α , on définit $Mod_{\mathcal{M}}(\alpha) = \{s \in S : i(s) \models \alpha\}$ et $\min_{\mathcal{M}}(\alpha) = \min(Mod_{\mathcal{M}}(\alpha), \prec)$.

Définition 18 *Soit une structure $\mathcal{M} = \langle S, i, \prec \rangle$ et soit $T \subseteq S$. On dit que T est smooth s'il satisfait la propriété suivante :*

$$\forall s \in T \setminus \min(T, \prec) \exists s' \in \min(T, \prec) \text{ t.q. } s' \prec s$$

On dit que $\mathcal{M}$ est un modèle préférentiel si $Mod_{\mathcal{M}}(\alpha)$ est smooth pour toute formule α .

On peut associer une relation d'inférence à tout modèle préférentiel de la manière suivante :

Définition 19 *Soit un modèle préférentiel $\mathcal{M} = \langle S, i, \prec \rangle$. La relation d'inférence $\vdash_{\mathcal{M}}$ est définie par :*

$$\alpha \vdash_{\mathcal{M}} \beta \text{ ssi } \min_{\mathcal{M}}(\alpha) \subseteq Mod_{\mathcal{M}}(\beta)$$

Et Kraus, Lehmann et Magidor [KLM90] ont prouvé le théorème de représentation suivant :

Théorème 19 $\vdash$ est une relation préférentielle si et seulement si il existe un modèle préférentiel $\mathcal{M} = \langle S, i, \prec \rangle$ tel que $\vdash_{\mathcal{M}} = \vdash$. Si le langage est fini, on peut choisir un S fini.

L'intuition derrière ce théorème de représentation est qu'un agent (i.e. un ensemble d'assertions conditionnelles $\alpha \vdash \beta$) dispose d'un pré-ordre partiel (une préférence) sur les états possibles du monde. Un état s est plus petit qu'un état s' si s est préféré à s' pour cet agent. Un agent inférera β de α si tous les états possibles préférés qui satisfont α satisfont également β .

On dit que $\mathcal{M} = \langle S, i, \prec \rangle$ est un *modèle rangé* si c'est un modèle préférentiel et si $\prec$ est modulaire¹.

En d'autres termes, un modèle est rangé si l'on peut ranger ses états par niveaux.

Lehmann et Magidor [LM92] ont montré le théorème de représentation suivant pour les relations rationnelles :

Théorème 20 Une relation d'inférence $\vdash$ est rationnelle si et seulement si il existe un modèle rangé $\mathcal{M} = \langle S, i, \prec \rangle$ tel que $\vdash_{\mathcal{M}} = \vdash$.

C'est-à-dire que l'on peut représenter une relation rationnelle par son modèle rangé :

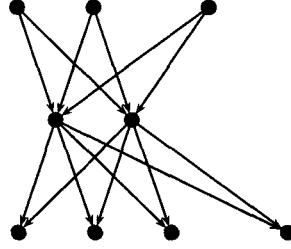


FIG. 1.6 – Représentation d'un modèle rangé

Sur la figure 1.6 est représenté un modèle rangé. Les points représentent les éléments de S et les flèches $\bullet_s \leftarrow \bullet_{s'}$ dénotent la relation $s \prec s'$.

Dans ce cas, la relation i est injective [Fre93, BMP97, PU99], c'est-à-dire que deux états différents correspondent à deux interprétations différentes. On peut alors remarquer que cette représentation est très proche des pré-ordres sur les interprétations des assignements fidèles (cf figure 1.2.4). La seule différence est que les pré-ordres des assignements fidèles sont totaux et donc deux interprétations à un même niveau sont équivalentes, alors que pour les modèles rangés, on utilise des ordres stricts modulaires et deux états à un même niveau sont incomparables.

Lorsque l'on regarde de plus près ces deux représentations, on peut donc entrevoir la correspondance entre relations rationnelles et opérateurs AGM que nous allons établir dans la section suivante.

1. si $x \not\prec y$, $y \not\prec x$ et $z \prec x$, alors $z \prec y$.

Correspondance relations d'inférence - opérateurs de révision

On peut définir une famille de relations rationnelles à partir d'un opérateur de révision de la façon suivante [MG89] :

Théorème 21 *Soient un opérateur de révision $*$ et une base de connaissance K . Si on définit la relation $\vdash_K$ comme suit :*

$$\alpha \vdash_K^* \beta \text{ ssi } \beta \in K * \alpha$$

Alors $\vdash_K^$ est une relation rationnelle qui satisfait la règle suivante, nommée préservation de la consistance :*

$$\text{Si } \alpha \vdash_K^* \perp, \text{ alors } \alpha \vdash \perp$$

Gärdenfors et Makinson ont également donné un résultat sur la construction inverse [MG89] :

Théorème 22 *Soit $\vdash$ une relation rationnelle qui préserve la consistance, il existe une base de connaissance K et un opérateur $*$ tels que $\vdash = \vdash_K^*$.*

A partir de ces résultats on peut donner un théorème de représentation qui exprime la correspondance entre un opérateur de révision et une famille de relations rationnelles :

On dit qu'une base de connaissance K correspond à une relation rationnelle $\vdash$ si $K = \{\alpha : \top \vdash \alpha\}$.

Théorème 23 *Un opérateur $*$ est un opérateur de révision si et seulement si à chaque base de connaissance K correspond une relation rationnelle $\vdash_K^*$ qui préserve la consistance telle que*

$$\alpha \vdash_K^* \beta \text{ ssi } \beta \in K * \alpha$$

1.2.6 Logique possibiliste

La logique possibiliste [DLP94, Lan91, Zad78] permet de modéliser de manière assez naturelle les informations imprécises et/ou incertaines. Cette distinction entre imprécision et incertitude, importante lorsque l'on tente de modéliser la connaissance d'un agent, est impossible à faire dans un cadre probabiliste.

Très grossièrement, on peut décrire la logique possibiliste comme un ensemble de formules auxquelles est attachée une information quantitative représentant la crédibilité de chaque formule.

Soit une relation $\geq_c$ sur les formules, $A \geq_c B$ signifie " A est au moins aussi certain que B ". $\geq_c$ est une *relation de nécessité qualitative* si elle vérifie les propriétés suivantes [Dub86, DP91] :

- | | |
|---|------------------------------|
| (D1) $A \geq_c A$ | (réflexivité) |
| (D2) $A \geq_c B$ ou $B \geq_c A$ | (totalité) |
| (D3) Si $A \geq_c B$ et $B \geq_c C$, alors $A \geq_c C$ | (transitivité) |
| (D4) $\top >_c \perp$ | (non trivialité) |
| (D5) $\top \geq_c A$ | (certitude de la tautologie) |
| (D6) Si $A \geq_c B$, alors $A \wedge C \geq_c B \wedge C$ | (stabilité conjonctive) |

Les conditions (D1)-(D3) expriment le fait que $\geq_c$ est un pré-ordre total. (D4) dit que les tautologies sont plus crédibles que les contradictions. Cette condition qui est la seule à contenir un ordre strict assure donc qu'il y a au moins 2 niveaux de certitudes, empêchant ainsi la trivialité $A \simeq_c B \forall A, B$. (D5) assure que les tautologies sont parmi les formules les plus crédibles. (D6) exprime le $\geq_c$ est stable pour la conjonction. C'est-à-dire que si A est plus crédible que B , avoir une information supplémentaire C ne doit pas nous faire changer d'avis.

Une *mesure de nécessité* est une fonction $N : \mathcal{L} \rightarrow [0,1]$ telle que :

- $N(\perp) = 0$,
- $N(\top) = 1$,
- $N(A \wedge B) = \min(N(A), N(B))$.

Une telle fonction associe donc à une proposition une mesure de nécessité et non pas un degré de véracité puisque $N(A) = 0$ n'indique pas que A est forcément fausse mais qu'elle ne bénéficie d'aucun support.

Il existe une notion duale à cette mesure de nécessité qu'est la mesure de possibilité :

Une *mesure de possibilité* est une fonction $\Pi : \mathcal{L} \rightarrow [0,1]$ telle que :

- $\Pi(\perp) = 0$,
- $\Pi(\top) = 1$,
- $\Pi(A \vee B) = \max(\Pi(A), \Pi(B))$.

Ces deux mesures sont interdéfinissables : $\Pi(A) = 1 - N(\neg A)$.

Une fonction f de $\mathcal{L}$ vers $[0,1]$ est dite compatible avec une relation $\geq_X$ sur $\mathcal{L}$ si et seulement si $A \geq_X B \Leftrightarrow f(A) \geq f(B)$.

Le résultat suivant donne alors la correspondance entre ces opérations quantitatives et la mesure qualitative [Dub86] :

Théorème 24 *Une relation de nécessité qualitative est compatible avec f si et seulement si f est une mesure de nécessité.*

Ces relations de nécessité qualitative sont très proches des enracinements épistémiques, comme l'interprétation intuitive de ces deux relations pouvait le laisser le supposer [DP91] :

Théorème 25 *L'ensemble d'axiomes (D2), (D3), (D5) et (D6) est équivalent à (EE1)-(EE4).*

La seule différence entre états épistémiques et relations de nécessité qualitative est que pour les premiers, on demande $\top > A$ (EE5), alors que les seconds n'imposent que $\top \geq_c A$ (D4).

Donc, pour les enracinements épistémiques, il est demandé que les seules formules maximales soient les tautologies, alors que pour les relations de nécessité qualitative, il est simplement demandé que les tautologies soient parmi les formules maximales.

Cette différence est nécessaire en logique possibiliste pour pouvoir coder des contraintes d'intégrité, c'est-à-dire des connaissances dont on est absolument sûr et que l'on ne veut pas voir remises en cause.

1.3 Critiques de la révision AGM

Nous avons vu à la section précédente que les postulats AGM correspondaient à un ensemble de méthodes de révision tout à fait intuitives. Nous avons vu également les rapports étroits entre révisions et d'autres domaines proches mais développés indépendamment tels que les relations d'inférence non monotones et la logique possibiliste. Ces postulats semblent donc disposer de justifications suffisantes en ce qui concerne leur fondement. D'un autre côté, aucun de ces postulats ne semble à l'abri d'une critique concernant telle ou telle situation. Pris individuellement, tous ces postulats sont critiquables et ils ont tous (ou presque) été critiqués. Le fait de critiquer ces postulats ne remet pas en cause le cadre édifié par Alchourrón, Gärdenfors et Makinson mais suggère des modifications ou des généralisations pour faire face à des cas de figure particuliers. Nous allons voir dans cette section les principales critiques qui ont été adressées à l'encontre des postulats AGM et les réponses apportées.

Le postulat AGM qui a fait couler le plus d'encre est sans aucun doute le postulat $(K \div 5)$, plus connu sous le nom de *recovery* (restauration). Ce postulat, bien que facilement critiquable intuitivement, semble nécessaire car, si l'on travaille avec des bases de connaissance closes déductivement, alors on ne peut pas définir d'opérateur acceptable ne satisfaisant pas *restauration*.

Cette constatation est l'une des motivations pour travailler avec des bases de connaissance non closes déductivement. Cette approche est détaillée section 1.3.2.

Une autre faiblesse du cadre AGM classique est qu'il ne prend pas suffisamment en compte l'itération du processus de révision. C'est-à-dire que les postulats AGM ne contraignent pas suffisamment deux révisions successives, et cela pose des problèmes de rationalité pour l'itération.

Dernièrement, le postulat de succès a été critiqué. En effet, nous n'avons pas toujours une confiance absolue en la nouvelle information et il peut être intéressant de disposer d'opérateurs n'incorporant pas forcément la nouvelle information telle quelle dans la base de connaissance. Hansson [Han97] a appelé ces opérateurs, opérateurs de semi-révision.

Enfin, une dernière remarque est que les postulats AGM ne caractérisent pas tous les types d'opérateurs de changement de la connaissance. En particulier, une distinction importante a été faite entre opérateurs de révision et opérateurs de mise à jour [KM91a, KW85]. Les premiers opérant une évolution des connaissances (incomplètes) à propos d'un monde statique, les seconds répercutant une évolution du monde sur les connaissances (obsolètes).

1.3.1 Restauration

Le plus critiqué des postulats AGM a sans doute été le postulat $(K \div 5)$, plus connu sous le nom de *restauration* (*recovery*) :

$(K \div 5)$ Si $A \in K$, alors $K \subseteq (K \div A) + A$ **(restauration)**

Ce postulat demande que lorsque l'on effectue la contraction d'une base K par une formule A suivie de l'expansion par cette même formule A , on doit retrouver alors la base de

connaissance de départ K (l'inclusion inverse de $(K \div 5)$ étant une conséquence de $(K \div 2)$).

Cette propriété semble naturelle puisqu'elle exprime la notion de changement minimal. Elle demande que la contraction par A n'enlève que le nécessaire pour ne plus impliquer A .

Mais des contre-exemples et des arguments contre ce postulat ont été donnés dans de nombreux travaux (voir par exemple [Han91b, Fuh91, Nie91, LR91]). Pour illustrer le problème soulevé par ce postulat, considérons l'exemple suivant donné par Hansson dans [Han91b] :

Exemple 1 *J'apprends dans un livre que Cléopâtre avait un fils et une fille. J'ajoute donc à l'ensemble de mes connaissances p et q dénotant respectivement le fait que Cléopâtre avait un fils et une fille. Un ami m'apprend que le livre que je lisais n'était pas un livre d'histoire mais un roman. Je dois donc revoir mes connaissances sur la maternité de Cléopâtre et supprime donc de mes connaissances le fait que Cléopâtre avait un enfant $p \vee q$. Peu après, dans un livre d'histoire j'apprends que Cléopâtre avait un enfant, j'ajoute donc à mes connaissances $p \vee q$. Dois-je ajouter à mes connaissances que Cléopâtre avait un fils et une fille (ce qu'impose restauration) ?*

Il semble donc que restauration induise un comportement plus que discutable pour les opérateurs de contraction AGM. Même Makinson reconnaît dans [Mak87] que :

"[recovery is] the only one among the six $[(K \div 1)-(K \div 6)]$ that is open to query from the point of view of acceptability under its intended reading."

Ce problème est interne à la contraction et cette critique de *restauration* n'affecte en rien la révision AGM puisque, comme noté à la remarque 1, ce postulat n'intervient pas lorsque l'on définit un opérateur de révision à partir d'un opérateur de contraction et de l'identité de Levi. Lorsque l'on définit un opérateur de contraction comme base d'un opérateur de révision on peut donc laisser de côté ce postulat. Mais le problème se pose lorsque l'on étudie les opérateurs de contraction en tant que tels.

Makinson nomme opérateurs d'*effacement* (withdrawal) les opérateurs satisfaisant les postulats $(K \div 1)-(K \div 4)$ et $(K \div 6)$, c'est-à-dire tous les postulats de base sauf *restauration*. Le problème est que *restauration* était le principal postulat de base en ce qui concerne le principe de conservation. Donc les opérateurs de withdrawal pèchent de ce point de vue, comme noté par Hansson [Han91a], un opérateur de withdrawal peut obéir à la propriété suivante :

$$K \div A = Cn(\emptyset) \quad \forall A \in K$$

Ce qui enlève de la base de connaissance des informations qui n'ont rien à voir avec la formule que l'on efface.

Dans [Han91b], Hansson explore les différentes manières d'affaiblir *restauration* afin de retrouver ce principe de conservation. Il propose alors le postulat de conservation du noyau (core-retainment) :

Si $A \in K$ et $A \notin K \div B$, alors $\exists K'$ tel que $K' \subset K$ et $B \notin Cn(K')$ et $B \in Cn(K' \cup \{A\})$
(conservation du noyau)

Cette propriété dit que si une formule n'est plus dans la base de connaissance après la contraction, c'est qu'elle contribuait à produire l'information que l'on voulait effacer. Ce postulat semble tout à fait raisonnable et moins fort que *restauration* pourtant [Han91b] :

Théorème 26 *Soit $\div$ un opérateur opérant sur une base de connaissance close déductivement K . $\div$ est un opérateur de contraction par intersection partielle si et seulement si il satisfait $(K \div 1)$, $(K \div 2)$, $(K \div 4)$, $(K \div 6)$ et conservation du noyau.*

Donc, lorsque l'on travaille avec des bases closes déductivement, *conservation du noyau* est équivalent à *restauration* ! Le problème est que les postulats de withdrawal et la propriété de conservation du noyau semblent intuitivement difficilement critiquables. Comme l'écrit Hansson [Han91a] :

*“Since it does not seem sensible for [a withdrawal operator] to violate core-retainment or any of $[(K \div 1)-(K \div 4)$ and $(K \div 6)]$, a reasonable withdrawal (contraction) operator without the recovery postulate does not seem possible in the AGM framework. The **pertinacity of the recovery property** is a prominent feature of the AGM framework.”*

Cette persistance de *restauration* dans le cadre de bases de connaissance closes déductivement est une des raisons qui ont poussé vers l'étude d'opérateurs de révision opérant sur des bases non closes déductivement. Nous appelons de tels opérateurs des opérateurs de révision syntaxique. Ceux-ci sont présentés plus en détail section suivante.

Eduardo Fermé a étudié dans le détail les implications de *restauration* pour les différentes méthodes de construction d'opérateurs de contraction (postulats, contractions par intersection partielle, enracinements épistémiques, contraction sûre et système de sphères) [Fer99a, Fer99b].

1.3.2 Révision syntaxique

Le cadre AGM requiert de travailler avec des bases de connaissance closes déductivement (bien que ce ne soit pas nécessaire pour tous les résultats). Bien que travailler avec des bases de connaissance closes déductivement est une idéalisation compréhensible pour développer une théorie de la révision, ce choix est discutable d'un point de vue algorithmique aussi bien que philosophique.

Du point de vue algorithmique, si l'on veut implanter un système de révision, travailler avec des bases closes déductivement semble simplement impossible d'un point de vue complexité.

Et, comme noté dans [Neb94] par Nebel, même si l'on suppose que l'on représente les bases de connaissance par des ensembles finis de formules équivalents logiquement aux bases de connaissance, les méthodes de révision usuelles (enracinement épistémique, révision par intersection partielle,...) utilisent des relations sur l'ensemble de toutes les formules d'une théorie close déductivement, ce qui pose des problèmes de représentation.

D'un point de vue philosophique il semble justifiable, lorsque l'on considère une base de connaissance A , de faire une distinction entre les informations qui ont été explicitement ajoutées dans la base de connaissance, ce qui forme le noyau de la base, et les informations

dérivées de ce noyau $Cn(A) \setminus A$, qui ne sont dans la base que comme des effets des informations du noyau. Il s'agit, en quelque sorte, de considérer les informations du noyau comme plus re-tranchées (au sens de l'enracinement épistémique) que les informations dérivées. Considérons l'exemple suivant donné par Hansson dans [Han91a] (voir aussi [Han89]) :

Exemple 2 *Nous sommes un jour férié et je me promène dans une ville qui compte deux fast-food. Je vois quelqu'un passer en mangeant un hamburger, j'en déduis donc qu'un des deux fast-food est ouvert $a \vee b$. En me dirigeant vers l'un des fast-food je vois de loin que l'éclairage de celui-ci fonctionne, j'en déduis donc que ce restaurant est ouvert a . Mes connaissances à ce moment peuvent être représentées par la base de connaissance $\{a, a \vee b\}$. Néanmoins, en arrivant à ce fast-food je lis une affiche indiquant que le restaurant est fermé aujourd'hui. L'éclairage ne fonctionnait que pour une personne faisant le nettoyage. La révision de ma base de connaissance doit donc contenir $\neg a$ mais également $a \vee b$ car j'ai toujours de bonnes raisons de penser que l'un des restaurants est ouvert.*

Imaginons à présent un scénario similaire, je me promène dans la ville mais ne croise pas de mangeur de hamburger. Lorsque je vois les lumières du fast-food ouvertes, ma base de connaissance est alors $\{a\}$. Lorsque je lis l'affiche m'annonçant que le restaurant est fermé, je n'ai aucune raison de penser $a \vee b$, c'est-à-dire que l'un des deux restaurants est ouvert.

Cet exemple illustre le fait qu'il peut être nécessaire de faire une distinction entre deux bases de connaissance (ici $\{a\}$ et $\{a, a \vee b\}$), bien que leur clôture déductive soit la même.

Le fait de donner une plus grande importance aux informations explicites, c'est-à-dire aux informations ayant un support direct, semble rapprocher cette idée de l'approche fondamentaliste de la révision. On oppose généralement ce que l'on appelle l'*approche cohérente*, qui considère les connaissances comme ne formant qu'un seul bloc, tel que le fait le cadre AGM, à l'*approche fondamentaliste*, qui considère les connaissances et leurs justifications, tel que les *truth-maintenance systems* [Doy79, dK86]. Les deux approches sont justifiables et elles sont en fait très liées [Doy92] (voir aussi [Neb89, del97]). Ces révisions syntaxiques semblent être un compromis entre ces deux approches. On s'autorise une distinction entre les informations explicites et implicites mais il n'y a pas de liste de justifications pour chaque information.

La solution est alors de développer des méthodes de révision opérant sur des bases de connaissance non closes déductivement. Ce problème a principalement été traité par Fuhrmann [Fuh91], Hansson [Han91a, Han91b, Han93] et Nebel [Neb89, Neb91, Neb94]. Hansson et Nebel appellent *belief set* une base de connaissance close déductivement et *belief base* une base non nécessairement close. Le problème posé alors est que les opérateurs ainsi définis n'obéissent plus au principe d'indépendance de syntaxe (K*6). C'est pour cette raison que nous appelons cela le cadre syntaxique.

Voir [Han91a] pour un inventaire détaillé des avantages de ces révisions syntaxiques. En particulier travailler avec des bases non closes déductivement fait gagner en expressivité, puisque deux bases $\{a\}$ et $\{a, a \vee b\}$ logiquement équivalentes (i.e. équivalentes statiquement), n'auront pas le même comportement après révision, elles ne sont donc pas forcément équivalentes dynamiquement. Il est intéressant de remarquer que l'on a une distinction similaire entre équivalence statique et dynamique dans le cadre de la révision itérée (cf section 2.1.1).

Un autre avantage des bases de connaissance non closes déductivement est que, dans ce cadre, il y a plusieurs bases inconsistantes. En effet, le problème lorsque l'on travaille avec des théories est qu'il n'y a qu'une seule base inconsistante, celle qui contient toutes les formules. Et lorsque l'on introduit une inconsistance "locale" (i.e. a et $\neg a$), cela "contamine" toute la base, interdisant toute réparation ultérieure. En revanche, lorsque l'on travaille avec des bases non closes déductivement, il y a plusieurs bases inconsistantes, par exemple la base $\{a, b, c, d, e, \neg a\}$ est différente de la base $\{a, \neg b, \neg c, \neg d, \neg e, \neg a\}$, alors que $Cn(\{a, b, c, d, e, \neg a\}) = Cn(\{a, \neg b, \neg c, \neg d, \neg e, \neg a\})$. Cette nuance permet de définir des opérateurs utilisant des bases de connaissance temporairement inconsistantes en évitant des trivialités [Han93].

Hansson nomme ses travaux sur les opérateurs de changement de bases de connaissance non closes déductivement *Belief Base Dynamics* (BBD). Hansson (re)définit les notions de contraction et de révision dans le cadre des BBD. Il inclut également dans sa caractérisation la notion de supersélecteurs qui sont des fonctions qui associent à chaque base de connaissance une fonction de sélection. Cette définition, proche des idées exposées au chapitre 3 permet donc également de supporter la révision itérée.

Une autre généralisation que supporte le cadre BBD est que la nouvelle information peut être un ensemble de formules, et non pas une seule formule comme dans le cadre AGM usuel. Mais, pour simplifier, on peut supposer que la nouvelle information est une formule unique. Voir [Han91a, Han89] pour plus de détails sur la révision/contraction par des ensembles de formules. Dans ce cas on définit naturellement $K \perp A$ comme étant l'ensemble des sous-ensembles maximaux de K qui n'impliquent aucun des éléments de A . De même, il faut définir ce qu'est la négation d'un ensemble de formules :

Définition 20 Soit A un ensemble fini de formules. On définit $\neg A$ par la formule :

- $\neg \emptyset = \perp$
- Si A est un singleton, $A = \{a\}$, alors $\neg A = \neg a$
- Si $A = \{a_1, \dots, a_n\}$, alors $\neg A = \neg a_1 \vee \dots \vee \neg a_n$

Nous allons présenter la caractérisation de Hansson [Han91a, Han93] mais il nous faut d'abord définir les supersélecteurs (appelés également *two-place selection functions* [Han93])

Définition 21 Un supersélecteur est une fonction f qui associe à chaque base de connaissance K une fonction de sélection $f(K) = \gamma_K$.

Un supersélecteur f est unifié si et seulement si pour toutes bases de connaissance K_1 et K_2 :

$$\text{Si } K_1 \perp A_1 = K_2 \perp A_2 \neq \emptyset, \text{ alors } (f(K_1))(K_1 \perp A_1) = (f(K_2))(K_2 \perp A_2)$$

Cette notion de supersélecteur unifié exige donc une certaine rationalité dans la définition de cette fonction. Un supersélecteur est relationnel (respectivement relationnel transitif) si et seulement si l'ensemble des fonctions de sélection qu'il associe aux bases de connaissance sont relationnelles (resp. relationnelles transitives).

Un opérateur de contraction par intersection partielle syntaxique est défini exactement comme dans le cadre AGM classique.

Un opérateur de contraction par intersection partielle syntaxique satisfait les conditions suivantes :

- (H÷1) $K \div A \subseteq K$ (inclusion)
- (H÷2) Si $x \in K \setminus (K \div A)$, alors il existe un K' avec $K \div A \subseteq K' \subseteq K$, tel que $A \cap Cn(K') = \emptyset$
et $A \cap Cn(K' \cup \{x\}) \neq \emptyset$ (pertinence)
- (H÷3) Si $A \cup Cn(\emptyset) = \emptyset$, alors $A \cup Cn(K \div A) = \emptyset$ (succès)
- (H÷4) Si pour chaque sous-ensemble K' de K on a $A \cup Cn(K') = \emptyset$ si et seulement si $B \cup Cn(K') = \emptyset$, alors $K \div A = K \div B$ (uniformité)
- (H÷5) Si $A \cup Cn(\emptyset) = \emptyset$, et si chaque élément de B implique un élément de A , alors $K \div A = (K \cap B) \div A$ (redondance)

Ces postulats sont une adaptation des postulats AGM dans le cadre de bases de connaissance non closes déductivement. On peut remarquer que la propriété de *pertinence* a déjà été citée section 1.3.1 comme une alternative au controversé postulat de restauration.

Théorème 27 *Un opérateur $\div$ est un opérateur de contraction par intersection partielle syntaxique si et seulement si il satisfait (H÷1)-(H÷4).*

Un opérateur $\div$ est un opérateur de contraction par intersection partielle syntaxique unifié² si et seulement si il satisfait (H÷1)-(H÷5).

Pour définir un opérateur de révision à partir d'un opérateur de contraction dans le cadre AGM, nous ne disposons que de l'identité de Levi :

$$K * A = (K \div \neg A) + A \quad \text{(Identité de Levi)}$$

Lorsque l'on veut définir dans le cadre des BBD un opérateur de révision à partir d'un opérateur de contraction on peut bien sûr utiliser la même identité, mais on peut également renverser cette identité [Han93], ce qui est impossible dans le cadre AGM classique car cette définition nécessite de passer par une base $K + A$ potentiellement inconsistante :

$$K * A = (K + A) \div \neg A \quad \text{(Identité de Hansson)}$$

Hansson nomme *révision interne* un opérateur de révision obtenu grâce à l'identité de Levi et *révision externe* un opérateur obtenu grâce à l'identité de Hansson. Plus formellement :

Définition 22 *Un opérateur de révision par intersection partielle interne est défini par*

$$K * A = (\cap f(K)(K \perp \neg A)) \cup A$$

Un opérateur de révision par intersection partielle externe est défini par

$$K \pm A = \cap f(K \cup A)((K \cup A) \perp \neg A)$$

Nous allons à présent énumérer un ensemble de propriétés pour les opérateurs de révision opérants sur des bases non closes déductivement :

- (H*0) $K * A$ est consistant si A est consistant (consistance)

2. Un opérateur de contraction est unifié si la fonction de sélection le définissant est unifiée.

- (H*1) $K * A \subseteq K \cup A$ (inclusion)
- (H*2) Si $x \in K \setminus (K * A)$, alors $\exists K'$ tel que $K * A \subseteq K' \subseteq K \cup A$, K' est consistant et $K' \cup \{x\}$ est consistant (pertinence)
- (H*3) $A \subseteq K * A$ (succès)
- (H*4) Si, pour tout $K' \subseteq K$, $K' \cup A$ est inconsistant ssi $K' \cup B$ est inconsistant, alors $K \cap (K * A) = K \cap (K * B)$ (uniformité)
- (H*4w) Si A et B sont inclus dans K et, pour tout $K' \subseteq K$, $K' \cup A$ est inconsistant ssi $K' \cup B$ est inconsistant, alors $K \cap (K * A) = K \cap (K * B)$ (uniformité faible)
- (H*5) Si A est consistant, et $A \cup \{b\}$ est inconsistant pour chaque $b \in B$, alors $K * A = (K \cup B) * A$ (redondance)
- (H*6) $K * A = (K \cup A) * A$ (pré-expansion)

Le théorème suivant résume les théorèmes de représentation donnés dans [Han93] :

Théorème 28 *Un opérateur $*$ est un opérateur de révision par intersection partielle interne si et seulement si il satisfait (H*0)-(H*4).*

Un opérateur $$ est un opérateur de révision par intersection partielle interne unifié si et seulement si il satisfait (H*0)-(H*5).*

Un opérateur $$ est un opérateur de révision par intersection partielle externe si et seulement si il satisfait (H*0)-(H*3), (H*4w) et (H*6).*

Un opérateur $$ est un opérateur de révision par intersection partielle externe unifié si et seulement si il satisfait (H*0)-(H*3), (H*4w), (H*5) et (H*6).*

Dans [Han98a], Hansson étudie de façon plus systématique les différentes propriétés logiques des opérateurs AGM et BBD et les rapports entre ces approches.

1.3.3 Itération

Une critique envers le cadre AGM est qu'il ne permet pas d'itérer le processus de révision de façon satisfaisante. Si l'on considère la caractérisation logique, les deux seuls postulats qui parlent d'itération sont (K*7) et (K*8). Et si l'on considère les deux méthodes usuelles de construction d'opérateurs de révision, à savoir les fonctions de révision par intersection partielle et les enracinements épistémiques, on se rend compte que lorsque l'on révisé la base de connaissance, il n'y a aucune indication sur la nouvelle fonction de sélection ou le nouvel enracinement épistémique à utiliser pour itérer le processus.

Les postulats AGM ne permettent pas d'assurer un bon comportement du processus d'itération de la révision parce qu'ils ne permettent pas d'assurer le maintien des informations conditionnelles.

Ces informations conditionnelles sont assez proches des *conditionnels* (counterfactuals) qui ont déjà été intensément étudiés (voir e.g. [Lew73, Sta68]). Un conditionnel peut être exprimé par une phrase du type : “Si α était vrai, alors β le serait aussi”, et il est généralement noté $\alpha > \beta$ ou à la manière d'une probabilité conditionnelle $\beta|\alpha$.

Les conditionnels sont très proches de la révision de la connaissance, ce lien a été souligné par de nombreux auteurs [Bou92a, Bou92b, Lev88], on peut d'ailleurs interpréter les relations d'inférence rationnelles $\alpha \vdash \beta$ comme des conditionnels, et la relation entre relations rationnelles et révision est bien connue (voir [MG89] et section 1.2.5).

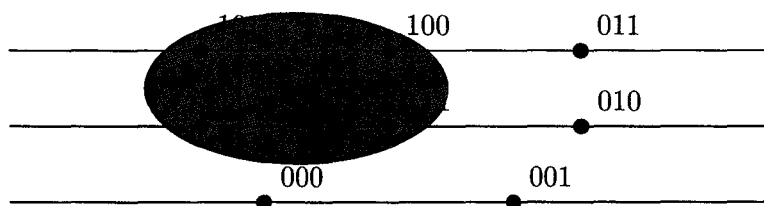


FIG. 1.7 – Information conditionnelle

Ce lien est illustré par le *test de Ramsey* [Ram31], énoncé par Stalnaker [Sta68] de la façon suivante :

“First add the antecedent (hypothetically) to your stock of beliefs; second make whatever adjustments are required to maintain consistency (without modifying the hypothetical belief in the antecedent); finally, consider whether or not the consequent is true.”

C'est-à-dire que pour tester si le conditionnel $\beta|\alpha$ appartient à la base de connaissance φ , il suffit de réviser la base de connaissance par α et voir si β est dans le résultat :

$$\beta|\alpha \text{ ssi } \varphi \circ \alpha \vdash \beta \quad (\text{Test de Ramsey})$$

Pour illustrer ce que l'on appelle informations conditionnelles dans le cadre de la révision, il suffit de considérer le pré-ordre de la figure 1.7. Si l'on nomme respectivement a, b, c les 3 variables propositionnelles, la base de connaissance associée à ce pré-ordre est $\varphi = \neg a \wedge \neg b$.

Si on apprend que a est vrai, les modèles minimaux de la nouvelle information sont $\{1,1,0\}$ et $\{1,1,1\}$ c'est-à-dire que $\varphi \circ a \vdash b$. Donc, si l'on apprend que a est vrai, on en conclura que b est vrai, c'est-à-dire que la base de connaissance φ contient l'information conditionnelle $b|a$.

Il est donc possible de représenter un pré-ordre par l'ensemble des informations conditionnelles correspondantes et la révision de la base de connaissance est alors l'ensemble des “conséquences” de ces conditionnels.

Les conditionnels sont donc une autre manière de coder les croyances (préférences) d'un agent et il serait souhaitable que la révision par une nouvelle information change les connaissances de l'agent, mais, autant que possible, ne modifie pas les croyances de l'agent. Le problème avec le modèle AGM est qu'il ne permet pas d'assurer le maintien de ces croyances. Par exemple avec la révision précédente, tout ce que l'on sait du nouvel état épistémique (le nouveau pré-ordre) de l'agent est que les modèles préférés (les connaissances) seront $\{1,1,0\}$ et $\{1,1,1\}$ mais on n'a aucune condition sur l'interclassement des autres interprétations. Il n'y a donc pas de conservation des informations conditionnelles.

Le fait est que dans le cadre AGM classique on ne s'intéresse pas à une stratégie de révision mais à une seule étape du processus. La question est simplement : étant données une connaissance et une nouvelle évidence, quelle est la nouvelle connaissance incorporant la nouvelle évidence?

Le problème lorsque l'on veut étudier un système d'information est que ce processus doit pouvoir s'itérer. Et cette itération pose de nouveaux problèmes de rationalité.

Pour montrer que le cadre AGM ne s'intéresse pas à la stratégie de révision, nous allons définir, figure 1.8, un opérateur de révision AGM utilisant le hasard lors du processus de révision.

-
- Initialiser la table de mémoire à vide.
 - Initialiser la base de connaissance à $\top$ et le pré-ordre correspondant est le pré-ordre plat où toutes les interprétations sont équivalentes.
 - Enregistrer cette base de connaissance et le pré-ordre correspondant dans la table de mémoire.
 - **tant que**(non fin_du_monde)
 - faire** - Acquérir la nouvelle information.
 - Calculer les modèles minimaux de la nouvelle information pour le pré-ordre courant.
 - La nouvelle base de connaissance est la formule (à équivalence logique près) ayant comme modèles ces modèles minimaux.
 - **si** la nouvelle base de connaissance apparaît déjà dans la table de mémoire
 - alors** - le pré-ordre courant est le pré-ordre correspondant à la base.
 - sinon** - Tirer au sort le nouveau pré-ordre courant ^a.
 - Enregistrer la base de connaissance et le pré-ordre correspondant dans la table de mémoire.
-

^ai.e. définir un pré-ordre où les modèles de la base de connaissance sont minimaux, et où les autres interprétations sont placées aléatoirement.

FIG. 1.8 – Opérateur de révision AGM aléatoire

Le fait de calculer aléatoirement les croyances de l'agent d'une révision à l'autre peut faire douter de la rationalité de cet opérateur en ce qui concerne une quelconque stratégie de révision. Pourtant cet opérateur est bien un opérateur AGM, puisque cette procédure construit incrémentalement un assignement fidèle.

La lacune des postulats AGM en ce qui concerne l'itération est que les opérateurs AGM semblent endogènes, car ils associent à une ancienne base de connaissance, une nouvelle base de connaissance de même nature. Mais l'endogénéité n'est qu'apparente, puisque la révision a nécessité une information supplémentaire sur les croyances de l'agent (enracinement épistémique, fonction de sélection, assignement fidèle...) mais ne produit aucune information de ce type. De manière fonctionnelle, si on considère la nouvelle information α comme étiquetant la transition, la révision d'une base de connaissance φ s'interprète comme :

$$(\varphi, \leq_{\varphi}) \xrightarrow{\alpha} (\varphi \circ \alpha, ?)$$

où $\leq_{\varphi}$ représente les croyances (préférences) de l'agent.

D'où la nécessité de considérer l'état épistémique d'un agent non pas comme une unique base de connaissance, mais comme un couple (base de connaissance, croyances), où "croyances" code l'information conditionnelle utilisée pour la révision suivante. En codant ces

croyances, on peut donc exprimer des conditions sur les stratégies de révision et leur rationalité.

On peut dire, en quelque sorte, que les opérateurs AGM classiques n'ont pas de mémoire et que l'on peut coder cette mémoire dans les croyances.

Nous verrons aux chapitres 2 et 3 les méthodes proposées pour résoudre les problèmes posés par l'itération.

1.3.4 Semi-révision

Récemment c'est le postulat de succès (K^*2) qui a été critiqué [Han97, Mak98, Han98b]. Ce principe, souvent appelé *primauté de la nouvelle information* (*primacy of update*) [Dal88a, Dal88b], dit que la nouvelle information est plus fiable que la base de connaissance actuelle et l'on donne donc une priorité supérieure à cette nouvelle information.

Ce n'est bien sûr pas toujours le cas, et il existe des applications où l'information la plus récente n'est pas forcément la plus fiable. Si on considère que l'agent est un robot et que les nouvelles informations à intégrer sont des observations données par des capteurs. Si ces capteurs sont parfaitement fiables, il n'y a pas de problème à appliquer (K^*2). En revanche, si les capteurs ne sont pas fiables, on dispose d'observations plus ou moins douteuses mais qui apportent tout de même des informations sur l'état du monde. Dans ce cas une méthode drastique pourrait être de ne pas tenir compte de cette nouvelle information peu fiable, mais, bien que l'on n'ait pas une confiance absolue en la nouvelle information, elle recelle tout de même un contenu informatif que l'on ne peut généralement pas totalement ignorer.

Dans [Mak98], Makinson propose une opération de révision filtrée (*screened revision*) qui n'accepte la nouvelle information que si celle-ci est compatible avec un noyau de connaissance qui ne peut être remis en cause par révision.

Plus formellement, il définit une *contraction protégeant* ξ :

$$K \div_{\xi} A = \cap \gamma(K \perp_{\xi} A)$$

où $K \perp_{\xi} A$ est l'ensemble des sous-ensembles maximaux de K qui contiennent $\xi \cap K$ et qui n'impliquent pas A .

L'opérateur de *révision filtrée* $\#_{\xi}$ est alors défini par :

$$K \#_{\xi} A = \begin{cases} Cn((K \div_{\xi} \neg A) \cup \{A\}) & \text{si } A \text{ est consistant avec } \xi \cap K \\ K & \text{sinon} \end{cases}$$

L'information ξ représente donc une connaissance de base qui ne peut être remise en doute. Une révision n'est autorisée que si elle ne remet pas en cause les éléments de $\xi \cap K$. Le problème étant que si $\xi \cap K$ est inconsistent, on ne pourra pas sortir de cette inconsistance par révision.

En généralisant un peu cette idée, Makinson propose alors un opérateur de révision filtrée relationnelle, où la connaissance de base n'est plus une base de connaissance, mais une relation entre formules, $A < B$ signifiant "*a priori*, A est moins crédible que B ". L'opérateur de

révision filtrée relationnelle $\#_{<}$ est alors défini par :

$$K\#_{<}A = \begin{cases} K\#_{\{B:A<B\}}A & \text{si } A \text{ est consistant avec } \{B : A < B\} \cap K \\ K & \text{sinon} \end{cases}$$

Dans [Han97], Hansson propose une série de postulats s'appliquant aux opérateurs de révision n'obéissant pas au postulat de succès. Il appelle ces opérateurs les opérateurs de *semi-révision*. Il s'intéresse plus particulièrement aux opérateurs travaillant avec des bases de connaissance non closes déductivement.

Par similitude aux opérateurs de révision externe (cf section 1.3.2) où l'on peut décomposer la révision de K par A en deux étapes :

1. Ajouter A à K ,
2. Contracter le résultat par $\neg A$.

On peut donner une définition similaire pour les opérateurs de semi-révision :

1. Ajouter A à K ,
2. Rendre le résultat consistant en supprimant des informations de A ou de K .

La différence se trouve donc dans la seconde étape. Lors d'une révision externe on ne supprime que des informations de K pour rétablir la consistance, alors que pour une semi-révision on peut également toucher à A . Cette définition n'est pas généralisable pour des bases closes déductivement, puisque le fait de n'avoir qu'une unique base inconsistante ne permet pas la construction d'opérations de ce genre.

Techniquement, les semi-révisions définies dans [Han97] utilisent pour cette seconde étape ce que Hansson appelle des *consolidations*, c'est-à-dire des contractions par $\perp$. Ces consolidations permettent d'obtenir une base de connaissance consistante à partir d'une base inconsistante. Nous ne détaillerons pas plus ici ces opérateurs. Mais il faut remarquer que la définition de ceux-ci :

1. Union de deux informations
2. Rétablissement de la cohérence

ne semble plus être de la révision mais de la fusion de deux bases de connaissance. En effet, de nombreuses méthodes de fusion suivent ce principe [BKM91, BKMS92, BDL⁺98]. Il est clair que révision et fusion de la connaissance sont des opérations très proches, et on voit ici à quel point la frontière est floue.

Pour insister sur ce fait, on peut également citer les travaux de Schlechta [Sch98] sur des opérateurs de semi-révision basés sur des distances. Schlechta a également étudié des opérateurs de révision basés sur des distances [SLM96, Sch98]. Intuitivement, le résultat d'une révision de K par A est l'ensemble des modèles de A les plus proches des modèles de K au sens de la distance choisie. Similairement le résultat d'une semi-révision est l'ensemble des interprétations appartenant aux couples composés d'un modèle de K et d'un modèle de A tels que la distance entre ces deux modèles est égale à la distance minimum entre un modèle de K et un modèle de A . Il est alors facile de rapprocher cette définition de semi-révision

basée sur une distance des opérateurs de fusion de Liberatore et Schaerf (voir [LS95, LS98] et section 5.2.2).

Fermé et Hansson ont proposé un autre type d'opérateurs de semi-révision, les opérateurs de *révision sélective* [FH99]. Ces opérateurs définissent des révisions où seulement une partie de la nouvelle information est acceptée, c'est-à-dire qu'un opérateur de révision sélective est défini par $K *_f A = K * f(A)$ où f est une fonction, typiquement telle que $f(A) \subseteq A$. Dans [FH99], les propriétés que doit satisfaire cette fonction f pour assurer à ces opérateurs de bonnes propriétés sont étudiées.

Finalement on peut également citer d'autres travaux où la nouvelle information n'est pas toujours acceptée, avec des approches plus quantitatives pour modéliser la confiance en la nouvelle information [Spo87, Wil94, BFH98].

1.3.5 Mise à jour

Katsuno et Mendelzon [KM91a], et avant eux Keller et Winslett [KW85], ont souligné que les postulats AGM ne s'appliquaient pas à tous les types de changements de la connaissance. En effet, les postulats de révision AGM ne s'appliquent qu'à la révision proprement dite et pas à la mise à jour. La différence entre ces deux opérations est qu'une révision permet d'incorporer une connaissance à propos du monde, c'est-à-dire d'améliorer sa connaissance du monde. Alors que la mise à jour permet de reporter un changement apporté au monde, c'est-à-dire d'actualiser sa connaissance du monde.

La révision permet donc de travailler sur un monde statique : l'état du monde ne change pas, c'est notre connaissance à propos de ce monde qui évolue. La mise à jour permet de travailler sur un monde dynamique : l'état du monde change et ces changements sont répercutés sur notre connaissance. Considérons l'exemple suivant pour illustrer ce propos [KM91a] :

Exemple 3 *Supposons que ma base de connaissance décrive deux objets, un livre et un magazine, dans une pièce. Il y a une table dans cette pièce et les objets peuvent être ou non sur la table. La formule l signifie "le livre est sur la table" et m "le magazine est sur la table". Je me rappelle que lorsque j'ai quitté cette pièce pour la dernière fois, il n'y avait qu'un seul objet sur la table, mais je n'ai pas pu voir lequel. Ma base de connaissance est $(l \wedge \neg m) \vee (\neg l \wedge m)$. J'envoie un robot dans cette pièce avec comme instruction de mettre le magazine sur la table. C'est-à-dire que je vais incorporer m à ma base de connaissance. Quelle est alors ma nouvelle base de connaissance ?*

Si on utilise un opérateur de révision pour traiter ce type de changement la nouvelle base de connaissance sera $\neg l \wedge m$, ceci étant une conséquence du postulat de succès ($K * 2$). Mais ce résultat est insatisfaisant, la base de connaissance initiale marquait une incertitude sur l'état du monde, c'est-à-dire que l'on était soit dans l'état $I = l \wedge \neg m$, soit dans l'état $J = \neg l \wedge m$. Le fait d'incorporer m à la base de connaissance avec un opérateur de révision lève cette incertitude et "choisit" l'état J . En quelque sorte, envoyer le robot avec comme instruction de mettre le magazine sur la table nous renseigne sur l'état du monde avant que le robot n'entre dans la pièce !

Alors que la seule chose que nous sachions lorsque le robot entre dans la pièce est qu'à présent le magazine est sur la table, cela ne nous informe pas sur l'état initial de la pièce. La nouvelle base de connaissance doit alors refléter le fait que le magazine est sur la table et que le livre peut ou non être sur la table : $(l \wedge m) \vee (\neg l \wedge m)$.

Katsuno et Mendelzon ont proposé une caractérisation logique des opérateurs de mise à jour ainsi qu'un théorème de représentation en terme de famille de pré-ordres sur les interprétations. Ces opérateurs sont une généralisation de la *Possible Models Approach* (PMA) de Winslett [Win88, Win90].

Soit φ et μ deux formules d'un langage propositionnel fini. φ est la base de connaissance, μ la nouvelle information et $\varphi \diamond \mu$ est une formule qui dénote la nouvelle base de connaissance après mise à jour de φ par μ . $\diamond$ est un opérateur de *mise à jour* s'il satisfait les propriétés suivantes :

- (U1) $\varphi \diamond \mu \vdash \mu$
- (U2) Si $\varphi \vdash \mu$, alors $\varphi \diamond \mu \leftrightarrow \varphi$
- (U3) Si φ et μ sont consistants alors $\varphi \diamond \mu$ est consistant
- (U4) Si $\varphi_1 \leftrightarrow \varphi_2$ et $\mu_1 \leftrightarrow \mu_2$ alors $\varphi_1 \diamond \mu_1 \leftrightarrow \varphi_2 \diamond \mu_2$
- (U5) $(\varphi \diamond \mu) \wedge \phi \vdash \varphi \diamond (\mu \wedge \phi)$
- (U6) Si $\varphi \diamond \mu_1 \vdash \mu_2$ et $\varphi \diamond \mu_2 \vdash \mu_1$, alors $\varphi \diamond \mu_1 \leftrightarrow \varphi \diamond \mu_2$
- (U7) Si φ est une formule complète, alors $(\varphi \diamond \mu_1) \wedge (\varphi \diamond \mu_2) \vdash \varphi \diamond (\mu_1 \vee \mu_2)$
- (U8) $(\varphi_1 \vee \varphi_2) \diamond \mu \leftrightarrow (\varphi_1 \diamond \mu) \vee (\varphi_2 \diamond \mu)$

Les postulats (U1)-(U5) sont très proches des postulats (R1)-(R5) de la révision. (U2) et (U3) sont une version affaiblie respectivement de (R2) et (R3). Le postulat (R6) a été remplacé par les 3 postulats (U6)-(U8). (U6) dit que, si la mise à jour d'une base par une nouvelle information μ_1 implique μ_2 , et si la mise à jour de cette même base par la nouvelle information μ_2 implique μ_1 , alors les deux mises à jour donnent le même résultat. (U7) dit que lorsque l'on a une connaissance complète du monde, les modèles communs aux deux mises à jour $\varphi \diamond \mu_1$ et $\varphi \diamond \mu_2$ sont également modèles de $\varphi \diamond (\mu_1 \vee \mu_2)$. Le postulat le plus important de la liste est sans doute (U8), c'est lui qui assure que l'on examine séparément chaque monde possible (modèle) de la base de connaissance.

Katsuno et Mendelzon [KM91a] ont également défini des opérateurs d'*effacement* (erasure) qui sont aux opérateurs de mise à jour ce que les opérateurs de contraction sont aux opérateurs de révision.

Du point de vue sémantique on dispose du théorème de représentation suivant :

Théorème 29 *Un opérateur de mise à jour $\diamond$ satisfait les postulats (U1)-(U8) si et seulement si il existe un assignement fidèle qui associe à chaque base de connaissance φ un pré-ordre partiel $\leq_\varphi$ tel que*

$$Mod(\varphi \diamond \mu) = \bigcup_{I \models \varphi} \min(Mod(\mu), \leq_{\varphi|_I})$$

Intuitivement, les opérateurs de révision apportent un changement minimal à l'ensemble des mondes possibles (les modèles de la base de connaissance) pour trouver les mondes les plus crédibles au vu de la nouvelle information, alors que les opérateurs de mise à jour apportent un changement minimal à chaque monde possible de façon à tenir compte de la nouvelle information quel que soit l'état dans lequel le monde se trouvait. Donc, en examinant ce théorème et le théorème correspondant pour les opérateurs de révision (théorème 16), on en déduit facilement le résultat suivant :

Théorème 30 *Soit un opérateur de révision $\circ$ satisfaisant les postulats (R1)-(R6), alors l'opérateur $\diamond$ défini par :*

$$\varphi \diamond \mu = \bigvee_{I \models \varphi} \varphi_{\{I\}} \circ \mu$$

est un opérateur de mise à jour satisfaisant (U1)-(U8).

Il est donc possible de définir un opérateur de mise à jour à partir d'un opérateur de révision donné.

Le problème général soulevé par le travail de Katsuno et Mendelzon est celui de l'ontologie des opérateurs de changement de la connaissance. En effet, lorsque l'on définit un type d'opérateurs de changement, il faut soigneusement décrire le type d'applications et les hypothèses faites par ces opérateurs. Il est par exemple utile de spécifier la nature de la nouvelle information : observations, connaissances, croyances... Friedman et Halpern [FH96] ont critiqué ce manque dans le cadre AGM et ont souligné l'importance d'attacher une ontologie à chaque définition d'opérateurs.

La plupart des travaux dans le domaine de la mise à jour étant la définition d'opérateurs pour des applications particulières, le travail de Katsuno et Mendelzon est une tentative de caractérisation logique de ces opérateurs au même titre que celle proposée par AGM dans le cadre de la révision.

Et comme la révision, la caractérisation logique des opérateurs de mise à jour a également subi des critiques (voir par exemple [HR98, DdSCP95, dSC96]). Le problème est que dans ce cadre la caractérisation logique n'a pas été appuyée par un ensemble de théorèmes de représentation. Cette caractérisation est donc, dans ce sens, plus fragile que la caractérisation AGM.

Chapitre 2

Révision itérée

On ne traverse jamais deux fois la même rivière; c'est toujours
une nouvelle eau qui coule.

Héraclite

Nous avons vu section 1.3.3 que le cadre AGM classique ne permet pas de traiter le problème de l'itération du processus de révision d'une manière satisfaisante et que cela est dû au fait que ces opérateurs ne permettaient pas d'assurer le maintien des informations conditionnelles.

Une autre façon d'envisager le problème est de parler de stratégie de révision, puisque dans le cas des opérateurs AGM classiques, rien ne lie deux opérations de révision successives. Pourtant, pour avoir des propriétés de rationalité en ce qui concerne l'itération, il semble naturel de conditionner une révision à l'itération précédente. Pour coder cette information sur les conditionnels, une simple base de connaissance n'est pas suffisante et il a été proposé par Darwiche et Pearl de travailler avec un cadre plus "riche". C'est cette adaptation du cadre AGM classique que nous présentons dans une première section. Bien que la proposition de Darwiche et Pearl semble être la plus aboutie en terme de révision itérée, nous verrons qu'elle n'est pas exempte de défauts.

Ensuite nous passons en revue les autres propositions de la littérature en ce qui concerne la révision itérée :

- Nous présentons l'opérateur de révision naturelle de Boutilier. Cet opérateur obéit de manière ostentatoire au principe de changement minimal. C'est en fait un cas limite des opérateurs de Darwiche et Pearl, et ceux-ci ont montré que cet opérateur pouvait avoir un mauvais comportement.
- Les opérateurs de révision rangée de Lehmann sont une alternative intéressante à la famille définie par Darwiche et Pearl mais semblent avoir une fâcheuse tendance à l'oubli.
- Les fonctions ordinales conditionnelles (OCF) définies par Spohn et utilisées également par Williams semblent avoir de bonnes propriétés. Mais ce cadre nécessite une information numérique marquant le degré de confiance en la nouvelle information. Ces opérateurs peuvent donc suggérer des adaptations dans un cadre purement qualitatif mais celles-ci ne sont pas nombreuses.

- Enfin Nayak *et al.* ont défini des opérateurs de révision itérée travaillant sur les enrachements épistémiques.

Dans une dernière section, nous illustrons les différences et points communs entre les différentes familles d'opérateurs. Ceux-ci se cristallisent en particulier sur le comportement des opérateurs après un certain nombre d'itérations. En effet, nous identifions dans cette section deux cas limites pour les opérateurs de révision. Et tous les opérateurs de révision itérée semblent mener à l'un de ces deux cas limites.

2.1 Proposition de Darwiche et Pearl pour l'itération

On peut décomposer la proposition de Darwiche et Pearl en deux parties. Dans un premier temps ils étendent la caractérisation AGM à des objets plus riches que des bases de connaissance. On appellera ces objets états épistémiques. Ceux-ci permettent de coder la stratégie de révision de l'agent. Dans un second temps Darwiche et Pearl proposent des postulats supplémentaires pour capturer les propriétés propres à l'itération.

2.1.1 AGM pour états épistémiques

La caractérisation AGM n'est pas suffisamment contraignante pour appréhender la révision itérée. Darwiche et Pearl [DP94, DP97] ont reformulé les postulats de Katsuno et Mendelzon [KM91b] en termes d'états épistémiques. Pour Darwiche et Pearl un état épistémique est un objet abstrait, représentant les croyances d'un agent, dont on peut extraire la connaissance actuelle.

Plus précisément, à chaque état épistémique Φ est associé une base de connaissance $Bel(\Phi)$ qui est une formule propositionnelle et qui représente la partie objective (*i.e.* les connaissances) de Φ . Les modèles d'un état épistémique Φ sont les modèles de sa base de connaissance associée $Mod(\Phi) = Mod(Bel(\Phi))$.

Pour alléger les notations on abrégera $Bel(\Phi) \vdash \mu$, $Bel(\Phi) \wedge \mu$ et $I \models Bel(\Phi)$ respectivement par $\Phi \vdash \mu$, $\Phi \wedge \mu$ et $I \models \Phi$.

A partir d'ici nous appellerons donc état épistémique un tel objet. Darwiche et Pearl ne précisent pas plus la nature de cet état épistémique, mais, à la lumière de ce qui a été dit à la section précédente, il est possible de décrire un état épistémique comme un couple (base de connaissance, ensemble conditionnel) où la base de connaissance est obtenue par l'opérateur de "projection" Bel et l'ensemble conditionnel représente l'information conditionnelle (*i.e.* l'enracinement épistémique, le pré-ordre total sur les interprétations, l'ensemble de conditionnels...) qui code la stratégie de révision.

On dispose de deux sortes d'équivalences entre états épistémiques. Une équivalence "faible" qui indique une identité statique entre les états épistémiques, c'est-à-dire que leurs bases de connaissance sont logiquement équivalentes mais, après une révision par une même information, il n'y a *a priori* aucune relation entre les états épistémiques résultats.

Une équivalence "forte" qui indique une identité dynamique entre états épistémiques, dans le sens où leurs bases de connaissance sont logiquement équivalentes et, après une révision par une même information, les bases de connaissance associées aux deux nouveaux états épistémiques seront également logiquement équivalentes.

En d'autres termes l'équivalence faible dénote l'équivalence des bases de connaissance alors que l'équivalence forte dénote l'équivalence des stratégies de révision.

On appellera simplement équivalence l'équivalence faible et égalité l'équivalence forte, ceci est traduit plus formellement dans la définition suivante.

Définition 23 Deux états épistémiques sont équivalents, noté $\Phi \leftrightarrow \Phi'$, si et seulement si leurs bases de connaissance sont des formules équivalentes $Bel(\Phi) \leftrightarrow Bel(\Phi')$. Deux états épistémiques Φ et Φ' sont égaux, noté $\Phi = \Phi'$, si et seulement si ils sont identiques.

Les postulats AGM pour états épistémiques sont les suivants :

Soient un état épistémique Φ et une formule μ dénotant la nouvelle information. $\Phi \circ \mu$ représente l'état épistémique résultat de la révision de Φ par μ . L'opérateur $\circ$ est un opérateur de révision s'il satisfait les propriétés suivantes :

- (R*1) $\Phi \circ \mu \vdash \mu$
- (R*2) Si $\Phi \wedge \mu$ est consistant, alors $\Phi \circ \mu \leftrightarrow \Phi \wedge \mu$
- (R*3) Si μ est consistant, alors $\Phi \circ \mu$ est consistant
- (R*4) Si $\Phi_1 = \Phi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\Phi_1 \circ \mu_1 \leftrightarrow \Phi_2 \circ \mu_2$
- (R*5) $(\Phi \circ \mu) \wedge \varphi \vdash \Phi \circ (\mu \wedge \varphi)$
- (R*6) Si $(\Phi \circ \mu) \wedge \varphi$ est consistant, alors $\Phi \circ (\mu \wedge \varphi) \vdash (\Phi \circ \mu) \wedge \varphi$

Cette formulation ressemble beaucoup à celle de Katsuno et Mendelzon. La principale différence est que l'on travaille avec des états épistémiques au lieu de bases de connaissance. Au niveau des postulats, seul (R*4) est plus faible que sa contrepartie usuelle. C'est en fait le seul postulat où l'on utilise la partie "conditionnelle" des états épistémiques. Voir [DP97, FH96] pour plus de détails sur la motivation de cette définition.

Il existe également un théorème de représentation montrant comment ces opérateurs peuvent être caractérisés en termes de familles de pré-ordres sur les interprétations. Avant de donner cette caractérisation sémantique, il est nécessaire de définir ce qu'est un assignement fidèle dans le cadre d'états épistémiques.

Définition 24 Une fonction qui associe à chaque état épistémique Φ un pré-ordre $\leq_\Phi$ sur les interprétations est un assignement fidèle si et seulement si :

1. Si $I \models \Phi$ et $J \models \Phi$, alors $I \simeq_\Phi J$
2. Si $I \models \Phi$ et $J \not\models \Phi$, alors $I <_\Phi J$
3. Si $\Phi_1 = \Phi_2$, alors $\leq_{\Phi_1} = \leq_{\Phi_2}$

A présent il est possible de reformuler le théorème de représentation de Katsuno et Mendelzon [KM91b] en termes d'états épistémiques (cf [DP97]) :

Théorème 31 Un opérateur de révision $\circ$ satisfait les postulats (R*1)-(R*6) si et seulement si il existe un assignement fidèle qui associe à chaque état épistémique Φ un pré-ordre total $\leq_\Phi$ tel que :

$$Mod(\Phi \circ \mu) = \min(Mod(\mu), \leq_\Phi)$$

Il faut noter ici que ce théorème ne contraint que la partie objective du nouvel état épistémique et ne donne aucune information sur la partie conditionnelle.

2.1.2 Les postulats supplémentaires de Darwiche et Pearl

Nous avons vu dans l'introduction qu'une forte limitation des postulats de la révision AGM est qu'ils imposent de trop faibles contraintes sur l'itération du processus de révision. Darwiche et Pearl [DP94, DP97] ont proposé des postulats pour la révision itérée. Le but de ces postulats est de garder le plus possible de connaissances conditionnelles de l'ancienne base de connaissance. En plus des postulats (R*1)-(R*6), un opérateur de révision doit satisfaire les postulats suivants :

- (C1) Si $\alpha \vdash \mu$, alors $(\Phi \circ \mu) \circ \alpha \leftrightarrow \Phi \circ \alpha$
- (C2) Si $\alpha \vdash \neg\mu$, alors $(\Phi \circ \mu) \circ \alpha \leftrightarrow \Phi \circ \alpha$
- (C3) Si $\Phi \circ \alpha \vdash \mu$, alors $(\Phi \circ \mu) \circ \alpha \vdash \mu$
- (C4) Si $\Phi \circ \alpha \not\vdash \neg\mu$, alors $(\Phi \circ \mu) \circ \alpha \not\vdash \neg\mu$

L'explication des postulats est la suivante : (C1) dit que si deux informations sont incorporées successivement et si la deuxième implique la première, alors incorporer seulement la seconde aurait donné la même base de connaissance. (C2) dit que lorsque deux informations contradictoires arrivent, la seconde seule donnerait le même résultat. (C3) dit qu'une information doit être gardée si l'on effectue une révision par une information qui, étant donnée la base de connaissance, implique la première. (C4) dit qu'aucune information ne peut contribuer à son propre rejet.

Darwiche et Pearl ont donné un théorème de représentation pour leurs postulats.

Théorème 32 *Soit un opérateur de révision qui vérifie (R*1)-(R*6). L'opérateur vérifie (C1)-(C4) si et seulement si l'opérateur et l'assignement fidèle correspondant vérifient :*

- (CR1) Si $I \models \mu$ et $J \models \mu$, alors $I \leq_{\Phi} J$ ssi $I \leq_{\Phi \circ \mu} J$
- (CR2) Si $I \models \neg\mu$ et $J \models \neg\mu$, alors $I \leq_{\Phi} J$ ssi $I \leq_{\Phi \circ \mu} J$
- (CR3) Si $I \models \mu$ et $J \models \neg\mu$, alors $I <_{\Phi} J$ seulement si $I <_{\Phi \circ \mu} J$
- (CR4) Si $I \models \mu$ et $J \models \neg\mu$, alors $I \leq_{\Phi} J$ seulement si $I \leq_{\Phi \circ \mu} J$

Les conditions (CR1)-(CR4) posent des relations sur les pré-ordres totaux avant et après révision. (CR1) impose que l'interclassement des modèles de la nouvelle information est préservé. (CR2) impose la même condition sur les contre-modèles de la nouvelle information. C'est-à-dire que le classement relatif des modèles et contre-modèles est préservé. (CR3) et (CR4) disent qu'après révision le classement entre modèles et contre-modèles peut être modifié mais simplement en faveur des modèles. C'est-à-dire que la révision augmente la crédibilité des modèles de la nouvelle information.

Dans [DP94] les postulats (C1)-(C4) ont d'abord été donnés comme complément aux postulats usuels (R1)-(R6).

Freund et Lehmann [FL94] ont montré que (C2) est inconsistant avec les postulats AGM. De plus Lehmann [Leh95] a montré que les postulats (C1) et (R1)-(R6) impliquent (C3) et (C4).

Dans [DP97] Darwiche et Pearl ont reformulé leurs postulats (et les postulats AGM) en termes d'états épistémiques (cf section 2.1.1) et, de ce fait, ont enlevé cette contradiction et ces redondances.

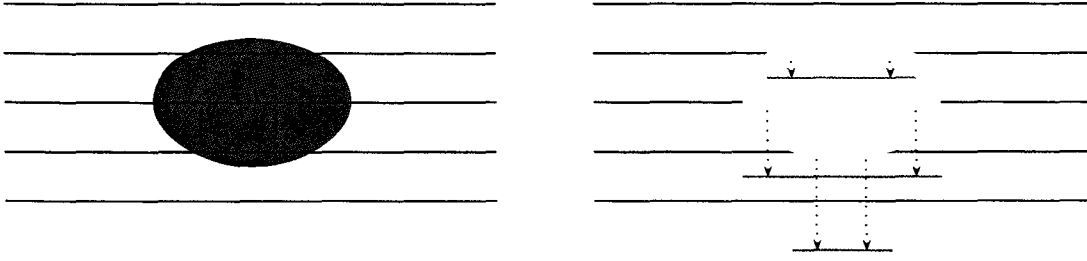


FIG. 2.1 – Révision Darwiche et Pearl

Nous illustrons le comportement des opérateurs de Darwiche et Pearl sur la figure 2.1. Ces figures représentent des pré-ordres totaux. Les interprétations ne sont pas représentées mais les lignes dénotent les “niveaux” du pré-ordre et il faut donc imaginer les interprétations distribuées sur chacun de ces niveaux.

La nouvelle information est représentée figure de gauche par la zone grisée et les contraintes imposées par les postulats de Darwiche et Pearl illustrées dans la figure de droite. C’est-à-dire que l’interclassement des modèles de la nouvelle information et celui de ses contre-modèles sont préservés, et on effectue une “descente” des modèles par rapport aux contre-modèles.

Bien que la proposition de Darwiche et Pearl semble être la plus aboutie en matière de révision itérée, elle souffre de quelques défauts. Elle est en effet à la fois trop et pas assez contraignante.

Le premier contre-exemple montre que les contraintes imposées par le postulat (C2) sont trop fortes [Kon98] :

Exemple 4 *Considérons un circuit composé d’un additionneur et d’un multiplieur. Nous n’avons initialement aucune information à propos de ce circuit i.e. $\Phi = \top$. Nous apprenons que l’additionneur et le multiplieur fonctionnent, en considérant les deux variables propositionnelles *adder_ok* et *multiplier_ok* dénotant respectivement que l’additionneur (resp. le multiplieur) fonctionne, la nouvelle information est donc $\mu = \text{adder_ok} \wedge \text{multiplier_ok}$. Quelqu’un teste alors le circuit et il s’avère que l’additionneur ne fonctionne pas correctement $\alpha = \neg \text{adder_ok}$. Doit-on alors “oublier” que le multiplieur fonctionne comme l’impose (C2) : $\alpha \vdash \neg \mu$ donc $(\Phi \circ \mu) \circ \alpha \leftrightarrow \Phi \circ \alpha \leftrightarrow \alpha$?*

Le problème avec le postulat (C2) est qu’il impose, lorsqu’une observation précédente est partiellement remise en doute, de ne pas tenir compte du tout de cette information. Ce comportement semble un peu excessif. Dans l’exemple, l’information μ est remise en question par α , il faut donc également remettre en question la partie de μ non remise en cause par α . Imaginons une application où la nouvelle information est composée d’une conjonction de dizaines de formules, si l’une d’entre elles est ultérieurement remise en cause, cela doit-il nous faire “perdre” l’information contenue dans les autres formules ?

Lehmann [Leh95], dont les opérateurs (définis section 2.3) ont le même comportement, illustre ceci par la remarque suivante dont il attribue l'idée à P. Y. Schobbens.

“Suppose an agent learns, first, a long conjunction $a \wedge b \wedge \dots \wedge z$ and then its negation $\neg a \vee \neg b \vee \dots \vee \neg z$. Postulate (I7)¹ implies it will forget about the first information, since it has been contradicted by the second one. But one could argue that this is not the right thing to do. Granted, the first information is incorrect, but it could be almost correct. If one makes this assumption, upon receiving the second information, one will conclude that few components, perhaps only one, of the conjunction are false, most of them still believed to hold true. This analysis distinguishes, I think two kinds of [revisions]. In the first one, one retracts a proposition because the source from which it has been obtained is now known to be unreliable. In this case, there is no reason to suppose the proposition is approximately correct. In the second kind of [revisions], one retracts a proposition because some new information came to contradict it. In this case, one may have reason to believe the proposition is still approximately correct. This distinction should lead to two different sets of postulates for [revisions].”

L'hypothèse sous-jacente de (C2) est que chaque information est produite par une source et que remettre en question la véracité d'une information remet en question la confiance en cette source. C'est une condition qui peut être utile pour certaines applications mais, outre le fait que Darwiche et Pearl ne mentionnent pas ce comportement dans leur article, cette hypothèse semble trop forte dans un cadre général.

Le second contre-exemple montre que, d'un autre côté, les postulats de Darwiche et Pearl ne contraignent pas suffisamment la révision itérée [NFPS96] :

Exemple 5 *Je crois que Tweety est un oiseau qui chante. Mais, comme il n'y a pas de forte corrélation entre le fait de chanter et celui d'être un oiseau, je suis tout de même prêt à penser que Tweety chante même s'il s'avère que Tweety n'est pas un oiseau. De même, si l'on m'apprend que Tweety ne chante pas, je penserai tout de même que Tweety est un oiseau.*

Supposons à présent que j'apprends que Tweety n'est pas un oiseau et, ensuite, qu'il ne chante pas. Avec ce scénario, il est raisonnable de s'attendre à ce que ma nouvelle connaissance soit que Tweety est un animal qui n'est pas un oiseau et qui ne chante pas.

Mais cela n'est pas garanti par les postulats de Darwiche and Pearl.

2.2 Révision naturelle

Boutilier propose [Bou93, Bou96] un opérateur de *révision naturelle* dont le but est d'avoir de bonnes propriétés en ce qui concerne l'itération. Cet opérateur peut être considéré comme accomplissant un changement minimal dans le pré-ordre associé aux états épistémiques.

Lorsque l'agent incorpore une nouvelle information, il considère les modèles minimaux de la nouvelle information pour le pré-ordre correspondant à l'ancien état épistémique. Et le pré-ordre associé au nouvel état épistémique est exactement le même que l'ancien, la seule différence est que les modèles minimaux de la nouvelle information sont les nouveaux modèles minimaux pour le pré-ordre. Ceci est illustré figure 2.2.

1. cf section 2.3.

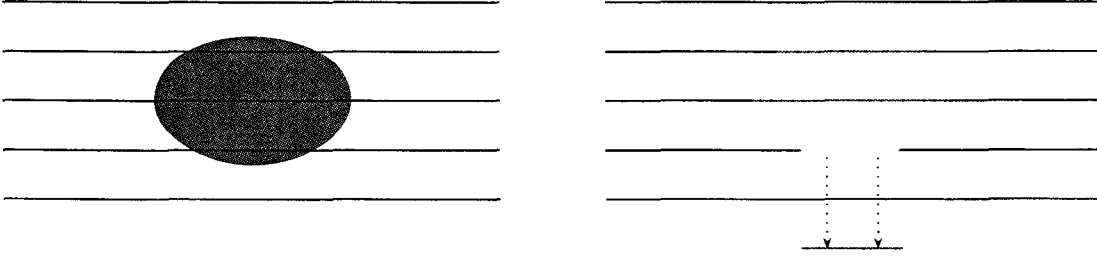


FIG. 2.2 – Révision naturelle

Ce principe est appelé *minimisation absolue* par Darwiche et Pearl et peut être caractérisé par [Bou93] :

(CB) Si $\Phi \circ \psi \vdash \neg \mu$, alors $(\Phi \circ \psi) \circ \mu \leftrightarrow \Phi \circ \mu$

Définition 25 $\circ$ est un opérateur de révision naturelle s'il vérifie (R*1)-(R*6) et (CB).

Au niveau sémantique, la condition de minimisation absolue correspond à la condition suivante :

(CBR) Si $I, J \models \neg(\Phi \circ \mu)$, alors $I \leq_{\Phi} J$ ssi $I \leq_{\Phi \circ \mu} J$

Et le théorème de représentation est le suivant [DP97] :

Théorème 33 Soit $\circ$ un opérateur de révision satisfaisant (R*1)-(R*6). $\circ$ vérifie (CB) si et seulement si l'opérateur et son assignement fidèle vérifient (CBR).

L'opérateur de révision naturelle est un cas particulier des opérateurs de Darwiche et Pearl :

Théorème 34 Si un opérateur $\circ$ vérifie (R*1)-(R*6) et (CB), alors il vérifie (C1)-(C4).

En examinant l'interprétation sémantique, on s'aperçoit que cet opérateur est un cas limite des opérateurs de Darwiche et Pearl. C'est en fait l'opérateur qui opère un changement minimal du pré-ordre associé à l'état épistémique. La justification d'un tel opérateur vient donc du principe du changement minimal : cet opérateur garde le maximum de conditionnels de l'ancienne base de connaissance.

Mais Darwiche et Pearl ont montré que suivre ce principe pouvait mener à des résultats discutables. Le problème est que cela induit un mauvais comportement pour l'itération du processus de révision car cette conservation des conditionnels peut remettre en cause des connaissances [DP97] :

Exemple 6 Je rencontre un étrange animal qui semble être un oiseau, je crois donc que cet animal est un oiseau. Cet animal se rapproche et je vois que cet animal est rouge, je pense donc que cet animal est rouge. Un expert en animaux étranges passe par là et m'apprend qu'en fait cet animal n'est pas un oiseau mais un mammifère. Je révisé donc mes croyances et pense à présent que cet animal est un mammifère. Que dois-je croire à propos de la couleur de l'animal ?

D'après l'opérateur de révision naturelle, on ne peut plus croire que l'animal est rouge (il suffit de prendre $\Phi = \top \circ \text{oiseau}$, $\mu = \neg \text{oiseau}$ et $\psi = \text{rouge}$). Bien que la couleur et l'espèce de l'animal ne soient *a priori* absolument pas liées, l'ordre dans lequel on apprend ces informations les "conditionnent" en quelque sorte pour l'opérateur de révision naturelle. Dans l'exemple ci-dessus on illustre bien le fait que, pour l'opérateur de révision naturelle, la croyance "l'animal est rouge" dépend de la croyance "l'animal est un oiseau", puisque remettre en cause cette dernière invalide la première. Alors que si l'on avait appris la couleur de l'animal avant d'avoir les informations sur son espèce, on aurait maintenu cette information sur sa couleur.

On peut illustrer ce problème directement sur la condition (CB) puisque lorsque $\Phi \circ \psi \vdash \neg \mu$, cette condition semble supposer que l'incompatibilité entre μ et $\Phi \circ \psi$ provient de ψ alors que μ peut très bien être parfaitement compatible avec ψ mais incompatible avec Φ .

Mais, malgré ce défaut, cet opérateur reste tout de même un opérateur avec de bonnes propriétés logiques et une définition simple qui semble en accord avec le principe de changement minimal.

2.3 Révision rangée

Lehmann [Leh95] propose des postulats pour les opérateurs de révision qui sont censés assurer de bonnes propriétés en ce qui concerne l'itération. Sa proposition est basée sur la notion de séquences de révisions (qui jouent le rôle d'états épistémiques). Soient Φ et Φ' , deux séquences de révisions, i.e. $\Phi = \varphi_1 \circ \dots \circ \varphi_n$, et soient φ et μ deux formules consistantes :

- (I1) $Bel(\Phi)$ est une théorie consistante
- (I2) $\Phi \circ \varphi \vdash \varphi$
- (I3) Si $\Phi \circ \varphi \vdash \mu$, alors $\Phi \vdash \varphi \rightarrow \mu$
- (I4) Si $\Phi \vdash \mu$, alors $\Phi \circ \Phi' \leftrightarrow \Phi \circ \mu \circ \Phi'$
- (I5) Si $\varphi \vdash \mu$, alors $\Phi \circ \mu \circ \varphi \circ \Phi' \leftrightarrow \Phi \circ \varphi \circ \Phi'$
- (I6) Si $\Phi \circ \varphi \not\vdash \neg \mu$, alors $\Phi \circ \varphi \circ \mu \circ \Phi' \leftrightarrow \Phi \circ \varphi \circ (\varphi \wedge \mu) \circ \Phi'$
- (I7) $\Phi \circ \neg \varphi \circ \varphi \subseteq Cn(Bel(\Phi) \cup \{\varphi\})$

Contrairement aux postulats Darwiche et Pearl, ces postulats ne sont pas donnés comme une extension des postulats AGM. Lehmann montre dans [Leh95] qu'ils capturent les postulats AGM.

Le postulat (I1) correspond aux postulats (K*1) et (K*5). (I1) semble plus fort que (K*5) car, dans ce cadre, Lehmann n'autorise des révisions que par des formules consistantes. Les postulats (I2) et (I3) correspondent exactement aux postulats (K*2) et (K*3). Le postulat (I4) dit qu'apprendre quelque chose que l'on connaît déjà ne modifie pas l'état épistémique. Le postulat (I5) exprime le fait qu'une révision peut être ignorée si elle est suivie de l'incorporation d'une information plus forte. (I6) dit que si l'on effectue deux révisions successives et si la seconde n'est pas une révision sévère, on peut alors remplacer cette séquence par la révision par la conjonction des deux formules. Ce postulat capture les postulats (K*7) et (K*8). (I7) est, d'après Lehmann, une expression du principe de changement minimal. Ce postulat peut

rappeler celui de restauration pour la contraction (cf section 1.1.3). Il implique qu'aucune modification apportée par l'incorporation de $\neg K$ ne doit être maintenue si l'on révisé ensuite par K .

La plupart des postulats donnés par Lehmann ne sont donc qu'une généralisation des postulats AGM au cas où l'on s'intéresse aux séquences de révision, ce qui permet d'appréhender la révision itérée. Les deux postulats qui concernent uniquement l'itération du processus sont (I5) et (I7).

Lehmann donne également une caractérisation sémantique de ces opérateurs sous forme de modèles rangés (widening ranking models²).

Définition 26 Soit $\mathcal{N}$ un segment initial de la classe des ordinaux. Un modèle rangé est une fonction $M : \mathcal{N} \mapsto 2^{\mathcal{W}} \setminus \emptyset$ telle que :

1. pour tout $n, m \in \mathcal{N}$, si $n \leq m$, alors $M(n) \subseteq M(m)$
2. pour tout $I \in \mathcal{W}$, $\exists n \in \mathcal{N}$ tel que $I \in M(n)$

C'est-à-dire qu'un modèle rangé est une succession de strates d'interprétations telle que chaque strate contient la strate qui lui est immédiatement inférieure et telle que chaque interprétation apparaît dans au moins une strate.

L'ordinal associé à chaque strate dénote le degré d'implausibilité de la strate.

A chaque séquence de révision Φ correspond un rang $r(\Phi)$ et un ensemble d'interprétations $p(\Phi) \subseteq M(r(\Phi))$.

La définition de la procédure de révision est itérative :

Définition 27 Soit un modèle rangé M . On procède par induction sur la longueur de Φ :

- Si $\Phi = []^3$, alors $r(\Phi) = 0$ et $p(\Phi) = M(0)$.
- S'il y a dans $p(\Phi)$ des modèles de φ , alors $r(\Phi \circ \varphi) = r(\Phi)$ et $p(\Phi \circ \varphi)$ est l'ensemble des éléments de $p(\Phi)$ modèles de φ . Si aucun élément de $p(\Phi)$ n'est modèle de φ , alors $r(\Phi \circ \varphi)$ est le plus petit $n > r(\Phi)$ tel qu'au moins un élément de $M(n)$ est modèle de φ . $p(\Phi \circ \varphi)$ est l'ensemble des éléments de $M(n)$ modèles de φ .

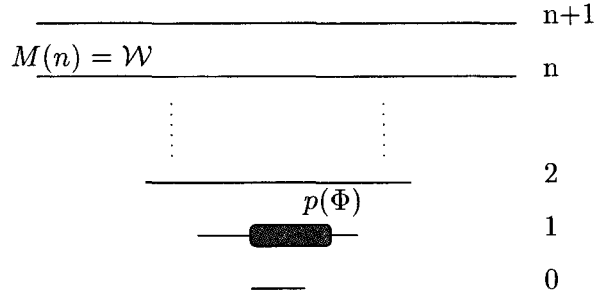
Pour toute séquence Φ , $Bel(\Phi)$ est la formule (à équivalence logique près) qui est satisfaite par tout modèle de $p(\Phi)$.

Les deux théorèmes suivants tiennent lieu de théorème de représentation [Leh95] :

Théorème 35 L'opérateur de révision défini à partir d'un modèle rangé (suivant la définition 27) satisfait les postulats (I1)-(I7).

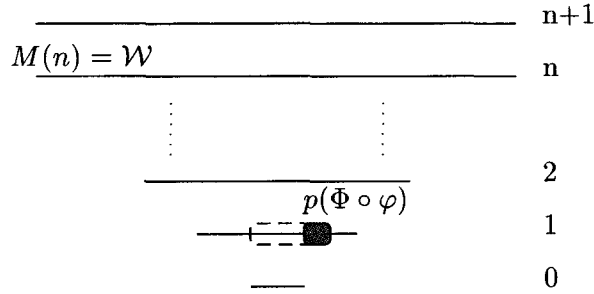
Le résultat inverse est donné par le théorème suivant :

Théorème 36 Un opérateur de révision qui vérifie les postulats (I1)-(I7) est défini par un modèle rangé.

FIG. 2.3 – *Modèle rangé*

On peut représenter les strates données par un modèle rangé. Rappelons que chacune de ces strates est incluse dans la strate qui lui est immédiatement supérieure.

La figure 2.3 illustre l'état épistémique Φ d'un agent, qui se "trouve" actuellement au niveau 1 de son modèle rangé. Lorsqu'une nouvelle information φ doit être incorporée, si elle est compatible avec Φ , i.e. s'il y a des modèles communs à $p(\Phi)$ et φ , alors on restreint la nouvelle connaissance à ces modèles communs (cf figure 2.4).

FIG. 2.4 – *modèle rangé : nouvelle information compatible*

Si ce n'est pas le cas, on change de strate jusqu'à trouver une strate qui contient des modèles de la nouvelle information. Supposons ici qu'il y a de tels modèles à la strate 2 (cf figure 2.5).

La représentation précédente était une illustration du fonctionnement des opérateurs de révision rangée, nous donnons à présent figure 2.6 une représentation graphique similaire aux autres approches, pour pouvoir établir une comparaison.

Bien que la proposition de Lehmann ne semble pas avoir soulevé de critique dans la littérature, nous pensons qu'elle souffre de quelques faiblesses :

1. Tout d'abord, il semble critiquable qu'un modèle rangé puisse avoir une strate initiale quelconque. C'est-à-dire que, d'après la définition, la séquence vide est associée à une connaissance non triviale. Ce qui signifie qu'en l'absence de toute information l'agent possède une information *a priori*.

2. littéralement *modèles rangés emboîtés*.

3. i.e. la séquence est la séquence vide.

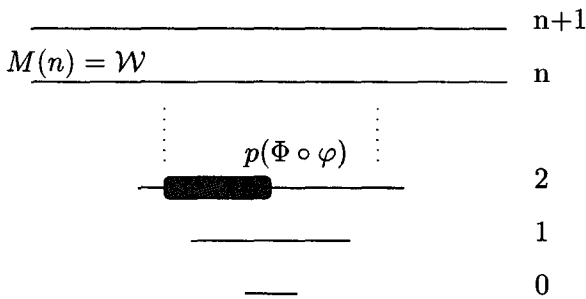


FIG. 2.5 – modèle rangé : nouvelle information non compatible

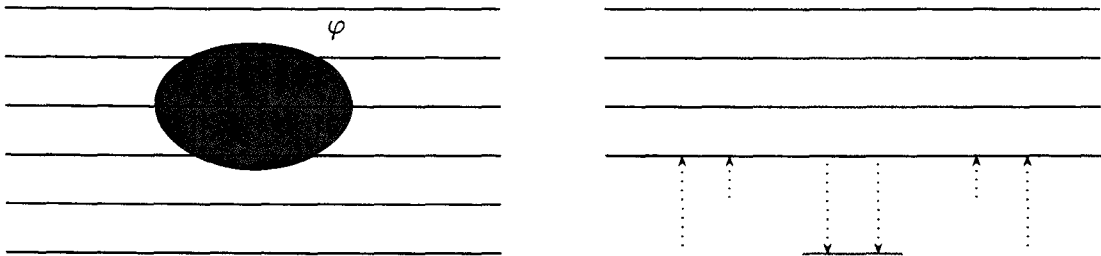


FIG. 2.6 – révision rangée

2. Une autre caractéristique, liée à la précédente, est qu’il existe donc une préférence *a priori* de l’agent donnée par son modèle rangé qui ne varie pas avec le temps (*i.e.* après révision). C’est-à-dire que les croyances de l’agent sont fixes.
3. Une autre conséquence liée à ce modèle rangé est que, lorsque la nouvelle information n’est pas compatible avec l’information courante, on change de strate, mais l’ancienne information n’intervient absolument pas dans la détermination de la nouvelle connaissance. Ce qui semble aller à l’encontre du principe de conservation.
4. La dernière et la plus grave critique que l’on peut adresser à ce modèle est que, comme illustré sur les figures 2.3 et 2.5, à partir d’un certain ordinal n^4 , toutes les strates sont égales et contiennent toutes les interprétations. C’est-à-dire que, lorsque l’on atteint cette strate n , toute révision sévère (*i.e.* non compatible avec la connaissance actuelle) par une information φ a comme résultat cette information φ . On peut voir cela également en utilisant la représentation usuelle (figure 2.6) car le nombre de “niveaux” diminue à chaque fois que la nouvelle information contredit les croyances actuelles. Donc, à partir d’un certain moment, on n’a plus qu’un seul niveau, *i.e.* toutes les interprétations sont aussi crédibles. C’est-à-dire que l’on atteint comme cas limite l’opérateur révision par intersection totale (*cf* définition 5), qui est considéré comme insatisfaisant pour une opération de révision. Voir la section 2.6 pour plus de précisions quant aux cas limites des opérations de révision itérée.

4. lorsque le langage est fini.

2.4 Fonctions ordinales conditionnelles et transmutations

Spohn [Spo87] propose de modéliser l'état épistémique d'un agent par une *fonction ordinale conditionnelle* (Ordinal Conditional Function ou OCF), qui peut être considérée comme un classement des interprétations.

Définition 28 Une *fonction ordinale conditionnelle* κ est une fonction de l'ensemble des mondes possibles $\mathcal{W}$ vers la classe des ordinaux telle que au moins un monde est associé à 0.

On nomme $\mathcal{C}$ l'ensemble des OCF. L'ordinal $\kappa(I)$ associé à un monde possible I peut être vu comme le *degré d'incrédulité* (degree of disbelief) de ce monde dans l'état épistémique représenté par l'OCF.

Cette fonction sur les mondes possibles peut être directement étendue aux propositions de la façon suivante :

$$\kappa(\varphi) = \min_{I \models \varphi} \kappa(I)$$

Une formule φ est crue pour un état épistémique représenté par κ si $\kappa(\varphi) = 0$. Sinon le degré d'incrédulité de φ est $\kappa(\varphi)$. Cette notion permet donc de différencier les formules qui ne sont pas crues dans l'état épistémique courant. En effet, on peut dire que la formule ayant le plus petit degré d'incrédulité est la plus vraisemblable dans l'état actuel.

Mais cette notion permet également d'établir une distinction entre les formules crues. On définit le *degré de confiance* (degree of firmness) d'une formule de la façon suivante : φ est crue avec une confiance α relativement à l'OCF κ si et seulement si :

- soit $\kappa(\varphi) = 0$ et $\alpha = \kappa(\neg\varphi)$,
- ou $\kappa(\varphi) > 0$ et $\alpha = -\kappa(\varphi)$.

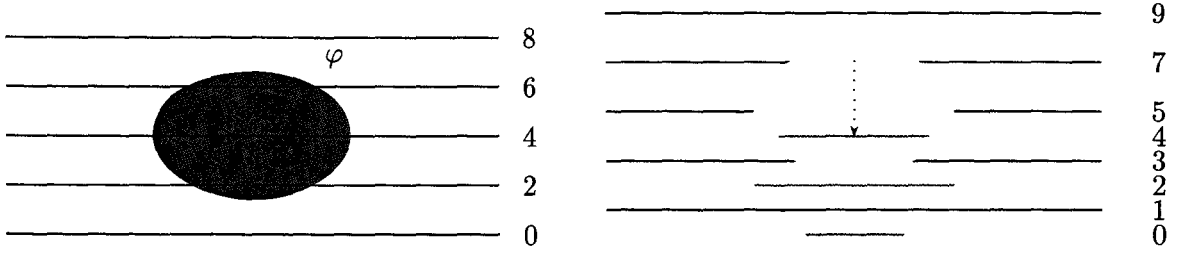
La révision d'un OCF définie par Spohn est appelée (φ, α) -conditionnalisation de κ , où κ est l'OCF courant, φ est la nouvelle information et α est le *degré de confiance* de la nouvelle information. Intuitivement, ce degré de confiance exprime la vraisemblance de la nouvelle information. Plus cette valeur est élevée, plus la nouvelle information est fiable. Le résultat de la conditionnalisation est un nouvel OCF où la nouvelle information est crue avec une confiance de α .

Définition 29 Soient un OCF κ , une formule φ , un ordinal α , la (φ, α) -conditionnalisation de κ est l'OCF $\kappa_{\varphi, \alpha}$ satisfaisant l'équation suivante :

$$\kappa_{\varphi, \alpha}(I) = \begin{cases} \kappa(I) - \kappa(\varphi) & \text{si } I \models \varphi \\ \kappa(I) - \kappa(\neg\varphi) + \alpha & \text{si } I \models \neg\varphi \end{cases}$$

Où $a - b$ représente l'unique ordinal c tel que $b + c = a$.

Ces opérateurs peuvent être vus comme un cas particulier des opérateurs de Darwiche et Pearl, puisque, comme illustré par la figure 2.7, la conditionnalisation par une formule φ correspond à une “descente” des modèles de φ par rapport aux contre-modèles. La contrainte supplémentaire est que cette descente est “normée”, car les opérateurs Darwiche et Pearl conservent simplement l'ordre relatif des modèles et des contre-modèles, alors qu'une conditionnalisation conserve les différences de confiance (au sens des degrés de confiance) relatives des modèles et des contre-modèles.

FIG. 2.7 – *Fonction Ordinale Conditionnelle : $(\varphi, 1)$ -conditionnalisation*

Il est clair que les opérateurs définis par Spohn satisfont les postulats de la révision AGM (si $\alpha \neq 0$) et ceux de la contraction AGM (si $\alpha = 0$), et que l'information ordinaire supplémentaire permet de définir des notions plus subtiles que celles dont on dispose dans le cadre AGM classique [Gär88]. Mais le principal inconvénient de cette proposition est qu'elle nécessite un degré de confiance pour la nouvelle information. Pour quelques applications cette mesure est fournie par le système d'information, mais, pour la plupart des applications, on ne dispose pas de cette information numérique sur la fiabilité de la nouvelle information. De plus, il est difficile d'apprécier la distinction entre de tels degrés de confiance. Est-il sensé de faire une distinction entre une nouvelle information d'un degré de confiance de 1000 et une autre d'un degré de confiance de 1001?

Dans les cas où l'on ne peut attacher une telle information numérique à la nouvelle information, il semble qu'il n'existe que deux solutions dans l'esprit des OCF. Spohn présente ces deux solutions comme des cas limites de son approche et critique ces deux approches. Puis il présente ses OCF comme le choix intermédiaire.

Il s'avère qu'un des cas limites correspond à la révision naturelle de Boutilier qui a déjà été largement étudiée [Bou93, Bou96]. Mais l'autre cas limite n'a jamais été étudié plus sérieusement. En fait, l'autre cas limite correspond à l'opérateur de révision avec mémoire basique qui sera présenté au chapitre 3.

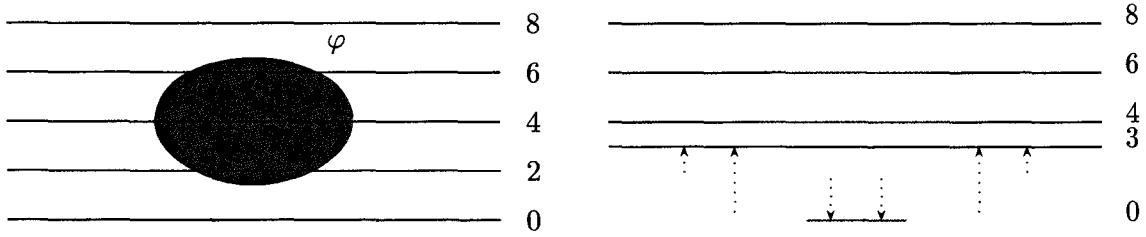
Williams [Wil94] reprend ces OCF comme moyen de représenter les états épistémiques et donne d'autres opérations de modifications de ces OCF. Williams nomme *transmutation* toute transition d'un état épistémique à un autre. En plus de la conditionnalisation définie par Spohn, elle propose la définition de transmutation et d'ajustement des OCF.

Définition 30 *Un schéma de transmutation d'un OCF est un opérateur*

$*$: $\mathcal{C} \times \{2^{\mathcal{W}} \setminus \{\emptyset, \mathcal{W}\}\} \times On^5 \rightarrow \mathcal{C}$ *qui envoie (κ, φ, i) dans $\kappa^*(\varphi, i)$ et qui vérifie :*

1. $\kappa^*(\varphi, i)(\neg\varphi) = i$
2. $Mod(Bel(\kappa^*(\varphi, i))) = \begin{cases} \{I \models \varphi : \kappa(I) = \kappa(\varphi)\} & \text{si } i > 0 \\ \{I \in \mathcal{W} : \text{soit } \kappa(I) = 0 \text{ ou } I \models \neg\varphi \text{ avec } \kappa(I) = \kappa(\neg\varphi)\} & \text{sinon} \end{cases}$

5. On est la classe des ordinaux.

FIG. 2.8 – Fonction Ordinale Conditionnelle : $(\varphi, 3)$ -ajustement

On appelle $\kappa^*(\varphi, i)$ la (φ, i) -transmutation de κ . Clairement, une conditionnalisation est un cas particulier de transmutation.

Une autre transmutation particulière est l'ajustement :

Définition 31 Soit φ une formule ($\varphi \neq \perp$, $\varphi \neq \top$) et i un ordinal. Le (φ, i) -ajustement d'un OCF κ est l'opérateur suivant :

$$\kappa^*(\varphi, i) = \begin{cases} \kappa^-(\varphi) & \text{si } i = 0 \\ (\kappa^-(\varphi))^\times(\varphi, i) & \text{si } 0 < i < \kappa(\neg\varphi) \\ \kappa^\times(\varphi, i) & \text{sinon} \end{cases}$$

où

$$\kappa^-(\varphi)(I) = \begin{cases} 0 & \text{si } I \models \neg\varphi \text{ et } \kappa(I) = \kappa(\neg\varphi) \\ \kappa(I) & \text{sinon} \end{cases}$$

$$\kappa^\times(\varphi, i)(I) = \begin{cases} 0 & \text{si } I \models \varphi \text{ et } \kappa(I) = \kappa(\varphi) \\ \kappa(I) & \text{si soit } I \models \varphi \text{ et } \kappa(I) \neq \kappa(\varphi) \text{ ou } I \models \neg\varphi \text{ et } \kappa(I) > i \\ i & \text{sinon} \end{cases}$$

D'après Williams, un ajustement est une transmutation qui induit un changement minimal dans l'OCF, c'est-à-dire qu'il obéit au principe de changement minimal. L'idée défendue ici est donc similaire à celle justifiant l'opérateur de révision naturelle de Boutilier. La différence est que le cadre est plus quantitatif (l'ajustement nécessite un degré de confiance en la nouvelle information) et que l'opération d'ajustement permet d'effectuer une révision, une contraction ou une restructuration, suivant la valeur du degré de confiance.

En fait, une (φ, i) -conditionnalisation tente de préserver au maximum l'ancien OCF lorsque l'on se restreint à φ et à $\neg\varphi$, alors qu'un ajustement tente de préserver au maximum l'ancien OCF dans son ensemble.

On appelle *restructuration* l'opération modifiant un état épistémique sans modifier sa base de connaissance associée. Cela permet de modifier la confiance relative entre les connaissances. On peut donc rendre une formule φ plus ou moins crédible, au vu d'une nouvelle information, sans remettre en cause les connaissances de l'agent. Plus formellement, Williams définit une

restructuration dans le cadre des OCF de la façon suivante :

Définition 32 Soit un OCF κ et φ une formule différente de $\perp$ et $\top$ tels que $\kappa(\neg\varphi) > 0$. Si i est un ordinal non nul, alors la (φ, i) -transmutation de κ est une restructuration.

Ajustement et conditionnalisation sont donc deux sortes de transmutations particulières. Williams a montré [Wil94] que toute transmutation pouvait, lorsque le langage est fini, s'exprimer comme une séquence d'ajustements ou comme une séquence de conditionnalisations. Ces deux notions semblent être des notions basiques puisqu'elles permettent d'exprimer toute transmutation.

Dans [Wil94], Williams montre que ces opérations sur les OCF correspondent à des opérations sur des enracinements épistémiques ordinaux (Ordinal Epistemic Entrenchment (OEE) également appelés Entrenchment Rankings dans [WFPS95]), qui sont une modification des enracinements épistémiques qui permet de tenir compte de l'information numérique additionnelle apportée par ce degré de confiance.

Définition 33 Une fonction d'enracinement épistémique ordinal est une fonction E de l'ensemble des formules $\mathcal{L}$ vers la classe des ordinaux qui satisfait les conditions suivantes : soient A et B deux formules de $\mathcal{L}$

- (OEE1) Si $A \vdash B$, alors $E(A) \leq E(B)$
- (OEE2) $E(A) \leq E(A \wedge B)$ ou $E(B) \leq E(A \wedge B)$
- (OEE3) $\vdash A$ si et seulement si $E(A) = \mathcal{O}^6$
- (OEE4) Si A est inconsistent, alors $E(A) = 0$

Williams définit donc également transmutation, conditionnalisation et ajustement pour ces OEE et montre l'équivalence entre OCF et OEE.

2.5 Entrenchment Kinematics

Un opérateur de révision étant équivalent à un enracinement épistémique, Nayak, Foo, Pagnucco et Sattar ont proposé des opérateurs de révision opérant directement sur ces enracinements épistémiques [NFPS94, NFPS96]. C'est-à-dire qu'ils modélisent un état épistémique par un enracinement épistémique et un opérateur de révision définit donc des transitions entre enracinement épistémiques.

Nayak *et al.* proposent une définition modifiée des enracinements épistémiques pour permettre leur révision. Nous appellerons ces relations des *enracinements épistémiques faibles* (WEE) :

- | | |
|---|----------------|
| (WEE1) Si $A \leq B$ et $B \leq C$, alors $A \leq C$ | (transitivité) |
| (WEE2) Si $A \vdash B$, alors $A \leq B$ | (domination) |
| (WEE3) $A \leq A \wedge B$ ou $B \leq A \wedge B$ | (conjonction) |
| (WEE4) Si $\exists C$ t.q. $\perp < C$, alors si $B \leq A \forall B$, alors $\vdash A$ | (maximalité) |

6. où $\mathcal{O}$ est un ordinal suffisamment grand, i.e. $\nexists C \in \mathcal{L}$ t.q. $E(C) > \mathcal{O}$.

Cette définition est un affaiblissement des enracinements épistémiques habituels puisque la condition de *minimalité* a disparu pour donner la possibilité d'atteindre une connaissance inconsistante. Pour la même raison, la condition de *maximalité* a été sensiblement modifiée.

On associe alors une base de connaissance à un WEE de la façon suivante :

Définition 34 Soit $\leq$ un enracinement épistémique faible.

$$Bel(\leq) = \begin{cases} \{A : \perp < A\} & \text{si } \perp \text{ n'est pas maximal pour } \leq \\ K_{\perp} & \text{sinon} \end{cases}$$

La révision s'effectue donc au niveau de l'enracinement épistémique :

Définition 35 Soit $\leq$ un WEE, on note $\leq_C^{\odot}$ le nouvel WEE résultat de la révision de $\leq$ par C défini de la façon suivante :

$A \leq_C^{\odot} B$ si et seulement si une des trois conditions suivantes est satisfaite :

1. $C \vdash \perp$
2. $C \vdash A$, $C \vdash B$ et soit $\vdash B$, soit ($\not\vdash A$ et $A \leq B$)
3. $C \not\vdash A$ et $(C \rightarrow A) \leq (C \rightarrow B)$

On peut alors lier ces opérateurs définis sur les enracinements épistémiques à des opérateurs travaillant sur les bases de connaissance de la façon suivante :

Définition 36 Soit une WEE $\leq$. On définit $*_{\leq}$ par :

$$B \in K *_{\leq} A \text{ ssi } \begin{cases} (A \rightarrow \neg B) < (A \rightarrow B) & \text{ou} \\ \vdash \neg A & \text{ou} \\ A \vdash B \end{cases}$$

Inversement, si $*$ est un opérateur de révision et K une base de connaissance, on définit $\leq_{*,K}$ par :

$$A \leq_{*,K} B \text{ ssi } \begin{cases} A \notin K * (\neg A \vee \neg B) & \text{ou} \\ \vdash A \wedge B & \text{ou} \\ K = K_{\perp} \end{cases}$$

Les opérateurs de révision ainsi définis satisfont (K*1)-(K*6) et les trois propriétés suivantes :

(K*0) $K_{\perp} * A = Cn(\{A\})$ (**absurdité**)

(K*7new) Si $A \wedge B \not\vdash \perp$, alors $(K * A) * B = K * (A \wedge B)$ (**conjonction**)

(K*8new) Si $A \wedge B \vdash \perp$, alors $(K * A) * B = K * B$ (**(DP2')**)

(K*0) indique comment sortir de la base inconsistante, grâce à une révision. Cette condition n'est pas une conséquence des postulats AGM.

La propriété (K*7new) (Conjonction) est une généralisation de la propriété de *récalcitrance* également proposée dans [NFPS96] :

Si $A \wedge B \not\vdash \perp$, alors $A \in (K * A) * B$ (récalcitrance)

Théorème 37 *En présence de (K^*1) - (K^*6) , Conjonction implique Récalcitrance, (K^*7) , (K^*8) , $(C1)$, $(C3)$ et $(C4)$ ⁷.*

Le problème avec la caractérisation logique de Nayak *et al.* est qu'ils ne considèrent pas d'états épistémiques, ils ne travaillent qu'avec des bases de connaissance. Sous peine d'incohérence de la caractérisation logique, il faut alors supposer que deux opérateurs de révision apparaissant successivement dans un même postulat sont des opérateurs de révision différents. Ce qui, en quelque sorte, code l'état épistémique au niveau de l'opérateur de révision. C'est-à-dire que si dans un postulat apparaît $(K * A) * B$, il faut comprendre $(K * A) *^{|A} B$, où $*^{|A}$ est le nouvel opérateur de révision.

Les opérateurs de Nayak *et al.* sont donc les opérateurs vérifiant (K^*0) , (K^*1-K^*6) , (K^*7_{new}) et (K^*8_{new}) . Ce qui est dit dans les théorèmes suivants :

Théorème 38 *Soit un $WEE \leq$. Soient la base de connaissance $K = Bel(\leq)$ et $\leq_A^\odot$ le résultat de la révision de $\leq$ par A . Alors les opérateurs de révision $* = *_{\leq}$ et $*^{|A} = *_{\leq_A^\odot}$ satisfont les postulats (K^*0) , (K^*1-K^*6) , (K^*7_{new}) et (K^*8_{new}) pour la base de connaissance $K = Bel(\leq)$.*

Théorème 39 *Soient une formule A et deux opérateurs de révision $*$ et $*^{|A}$ qui vérifient les postulats (K^*0) , (K^*1-K^*6) , (K^*7_{new}) et (K^*8_{new}) . Soit une base de connaissance K , alors $\leq_{*,K}^\odot = \leq_{*^{|A}, K * A}$.*

2.6 Cas limites

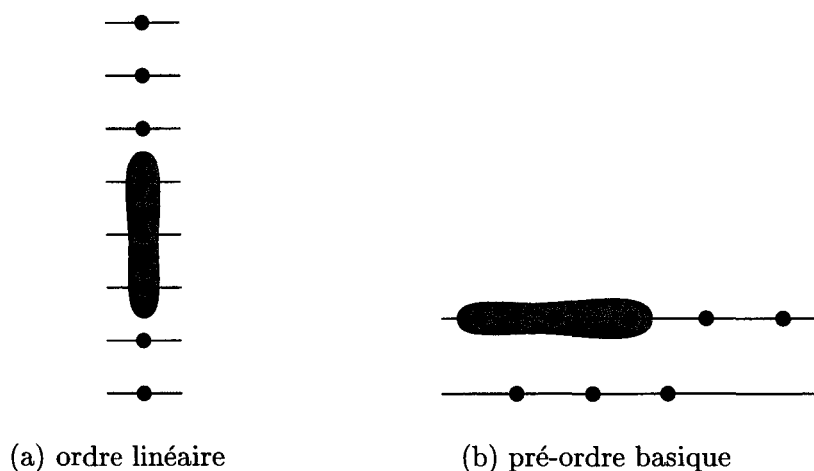
Lorsque l'on considère la représentation des différents opérateurs de révision itérée, on peut décrire ceux ci comme opérant une transformation sur un pré-ordre à chaque révision par une nouvelle information.

On peut citer deux types de pré-ordres particuliers que sont les pré-ordres linéaires et les pré-ordres basiques.

Les pré-ordres linéaires sont des pré-ordres totaux tels qu'il n'existe pas d'interprétations équivalentes, i.e. $\forall I, J$, si $I \neq J$, on a soit $I < J$ soit $J < I$. Ces pré-ordres sont donc associés à une base de connaissance complète (celle dont le seul modèle est l'interprétation minimum pour le pré-ordre) et toute révision d'un tel pré-ordre donnera également une base de connaissance complète comme résultat.

Les pré-ordres basiques sont les pré-ordres à deux niveaux, les modèles de la base associée au niveau minimum et les contre-modèles dans l'autre niveau. Toute révision sévère d'un tel pré-ordre, c'est-à-dire toute révision par une nouvelle information inconsistante avec la base de connaissance, se résumera à cette nouvelle information, i.e. si $\varphi \vdash \neg \mu$, alors $\varphi \circ \mu = \mu$.

⁷. formulés en terme de bases de connaissance et pas d'états épistémiques.

FIG. 2.9 – révision par φ

Il semble que tous les opérateurs de révision itérée vérifiant un certain nombre de propriétés logiques ont comme cas limite un de ces types de pré-ordres. C'est-à-dire qu'après un certain nombre d'itérations (par des formules aléatoires), soit tous les pré-ordres associés aux bases de connaissance sont des pré-ordres linéaires, soit ce sont tous des pré-ordres basiques.

Par exemple, il est facile de montrer que les opérateurs de Darwiche et Pearl, de Boutilier, les conditionnalisations de Spohn et les opérateurs de révision à mémoire (cf chapitre 3) mènent à des pré-ordres linéaires. L'argument est simplement que chaque révision soit conserve le nombre de niveaux du pré-ordre précédent, soit augmente le nombre de niveaux. Donc, après un certain nombre d'itérations (par des formules aléatoires), on atteint le plus grand nombre de niveaux possible, c'est-à-dire l'ordre linéaire.

De même, on peut montrer facilement que les révisions rangées de Lehmann mènent à des pré-ordres basiques, puisque le nombre de niveaux des pré-ordres diminue. Cela se voit facilement puisque les niveaux de crédibilité sont imbriqués, donc il existe forcément⁸ un niveau contenant toutes les interprétations, et tous les niveaux supérieurs sont égaux. Lorsque l'on a atteint ces niveaux de crédibilité, les pré-ordres associés aux bases de connaissance sont des pré-ordres basiques.

Tous les opérateurs existants mènent donc à un de ces deux cas limites et il semble que l'on ne puisse pas échapper à l'un de ces deux comportements. La question est alors de savoir si l'un de ces deux cas est préférable à l'autre.

Seuls les ajustements de Williams semblent échapper à ces deux cas limites mais rappelons que le cadre des OCF est plus quantitatif puisqu'il suppose l'existence d'une information numérique dénotant la crédibilité de la nouvelle information. Et il ne semble pas possible de donner une contrepartie de ces ajustements dans un cadre purement qualitatif. Plus exac-

8. dans le cas fini.

tement l'opérateur de révision naturelle semble être la seule alternative, mais cet opérateur mène à des pré-ordres linéaires.

On peut interpréter le résultat induit par le cas donnant des ordres linéaires comme un processus d'apprentissage. En effet, au bout d'un certain temps, l'expérience de l'agent (son historique de révisions) est suffisamment importante pour lui donner une image très précise du monde. Et, lorsqu'une nouvelle information arrive, il prend en compte son expérience passée pour établir ces nouvelles connaissances.

En revanche, mener à des pré-ordres basiques après un certain nombre d'itérations semble dénoter un processus d'oubli, puisque l'on perd alors toute préférence sur les mondes possibles. C'est-à-dire que ces opérateurs mènent à des révisions par intersection totale. Ce qui est donc critiquable.

En conclusion, s'il s'avère qu'effectivement aucun opérateur de révision, dans un cadre purement qualitatif, ne peut échapper à l'un des deux cas limites, ordres linéaires ou pré-ordres basiques, alors il semble naturel de préférer les opérateurs menant à des ordres linéaires. En effet, dans le cadre d'un système à bases de connaissance, il n'est pas souhaitable que l'opération de révision devienne triviale après un certain nombre d'itérations.

Chapitre 3

Principe de primauté forte

Nous allons présenter dans ce chapitre une famille d'opérateurs de révision qui obéit à ce que nous appelons le *principe de primauté forte* de la nouvelle information.

Le *principe de primauté* de la nouvelle information est habituellement reconnu comme une hypothèse de base pour la révision. Ce principe est principalement capturé par le postulat (K*2) dans le cadre AGM, ce postulat exprime le fait que la nouvelle information est vraie après révision par cette information.

Le *principe de primauté forte* que nous définissons ici renforce cette contrainte en demandant qu'après la révision par une nouvelle information tous les modèles de cette nouvelle information deviennent plus crédibles que les contre-modèles.

Nous verrons comment exprimer logiquement cette contrainte et nous montrerons que les opérateurs vérifiant cette propriété ont un comportement satisfaisant en ce qui concerne la révision itérée.

Nous pensons que ce renforcement du principe de primauté usuel n'est pas trop exigeant et dans beaucoup de cas assez naturel. Nous allons tenter d'expliquer le bien fondé de ce *principe de primauté forte* sur un petit exemple :

Exemple 7 *Supposons que les nouvelles informations que je reçois sont des observations à propos du monde et que j'ai une confiance entière en ces informations. J'apprends que quatre faits atomiques sont vrais $\mu = a \wedge b \wedge c \wedge d$. J'effectue donc une révision de mes connaissances par cette nouvelle information μ et, comme l'exige le principe de primauté, μ est vrai dans ma nouvelle base de connaissance.*

Mais considérons également mes croyances, c'est-à-dire la crédibilité que j'accorde aux autres mondes possibles. Comme j'ai toute confiance en la nouvelle information, les mondes que je considère les plus crédibles (en dehors de mes connaissances actuelles) sont les mondes où la nouvelle information est vraie. C'est-à-dire que je préférerais toujours un monde où la nouvelle information est vraie à un monde où cette information est fausse.

Les autres mondes ne satisfont donc pas μ mais ils ne sont pas tous équivalents. Puisque j'ai toute confiance en la nouvelle information, je trouverai plus crédible un monde qui satisfait trois de quatre faits atomiques qu'un monde qui n'en satisfait qu'un, voire aucun.

C'est cette attitude que nous nommons *principe de primauté forte*. Nous pensons que, psychologiquement, cette hypothèse est assez réaliste et, par conséquent, que ce n'est donc pas une condition trop contraignante.

L'hypothèse sous-jacente ici est que les faits atomiques sont indépendants, dans le sens où il n'existe aucun lien logique entre eux. C'est une hypothèse assez faible, tout à fait justifiable si la base de connaissance représente notre connaissance sur l'état d'un monde.

L'objection que l'on pourrait faire est qu'il existe tout de même généralement des liens logiques entre faits atomiques. Par exemple, les faits atomiques "*il pleut*" et "*la route est humide*" ne sont indéniablement pas indépendants. Mais cette dépendance peut-être vue comme une contrainte d'intégrité du système. En effet, il n'y a aucune raison que deux faits atomiques a et b soient logiquement liés *a priori*. La seule explication pour laquelle on considère "*il pleut*" et "*la route est humide*" comme liés est que nous avons une connaissance de base (la contrainte d'intégrité de notre système) qui contient la règle $\text{il pleut} \rightarrow \text{la route est humide}$.

Le cas où les faits atomiques sont liés logiquement peut donc être considéré comme un cas particulier de révision où l'on dispose d'une base de contraintes d'intégrité sur le système. Il est également vrai que, dans certains cas, il existe des liens logiques entre faits atomiques que nous ne soupçonnons pas. Mais cela revient simplement à dire que nous avons une connaissance imparfaite de la base de contraintes d'intégrité. L'hypothèse que les faits atomiques sont indépendants est donc une idéalisation tout à fait satisfaisante pour une étude de la révision itérée.

Nous étudions également dans ce chapitre la notion d'historique de révisions. Nous soutenons que maintenir un historique des révisions est un moyen de garantir un comportement rationnel en ce qui concerne l'itération du processus de révision. Le point important est que l'on ne peut pas modéliser les croyances d'un agent simplement avec ses connaissances actuelles mais on doit se soucier de comment cet agent a acquis ces connaissances. Un état épistémique ne peut donc pas être représenté par une base de connaissance seule. Une information supplémentaire est nécessaire comme, par exemple, l'historique des bases de connaissance acceptées par l'agent. Considérons l'exemple suivant pour illustrer cette idée.

Exemple 8 Considérons les deux séries de révisions suivantes : $\Phi_1 = a * (a \rightarrow b)$ et $\Phi_2 = (a \rightarrow b) * a$, nous obtenons alors la même base de connaissance $a \wedge b$. Mais si plus tard nous apprenons que b n'est pas vrai, doit-on obtenir $\Phi_1 * \neg b = \Phi_2 * \neg b$?

Nous pensons que non. Lorsque l'on apprend $\neg b$ nous ne pouvons garder à la fois a et $a \rightarrow b$. Selon le principe de primauté de la nouvelle information, dans Φ_1 l'information a est moins fiable que $a \rightarrow b$. Nous sommes donc enclins à accepter la fausseté de la plus ancienne information et nous obtenons donc $\Phi_1 * \neg b = \neg a \wedge \neg b$. Inversement dans Φ_2 , c'est $a \rightarrow b$ qui est moins fiable, donc lorsque l'on apprend $\neg b$ cela donne $\Phi_2 * \neg b = a \wedge \neg b$.

Lorsque l'on a effectué la première révision, nous avons obéi au principe de primauté de la nouvelle information, et donc nous avons considéré que la nouvelle information était plus fiable que l'ancienne. Pourquoi à l'étape suivante (i.e. au moment d'incorporer $\neg b$) devrions-nous oublier cela comme le demande le cadre AGM classique?

Comme souligné section 1.3.3, il y a une différence entre bases de connaissance et états épistémiques. Un état épistémique contient en général une information beaucoup plus com-

plexe qu'une simple base de connaissance. Cette information plus complexe peut, par exemple, être un ensemble de conditionnels. Ainsi, dans l'exemple précédent, $\Phi_1 * \neg b = \neg a \wedge \neg b$ signifie que Φ_1 contenait le conditionnel $\neg a | \neg b$ ("Si $\neg b$ était vrai alors $\neg a$ le serait aussi"), alors que dans Φ_2 le conditionnel était $a | \neg b$ ("Si $\neg b$ était vrai alors a le serait aussi").

3.1 De la notion d'état épistémique

Il existe une ambivalence de la notion d'état épistémique qu'il faut éclaircir ici sous peine de confusion.

Philosophiquement, l'état épistémique d'un agent représente son état cognitif à un moment donné. C'est-à-dire les connaissances et les croyances de l'agent. Peter Gärdenfors les définit ainsi [Gär88]:

"The first and most fundamental factor in an epistemological theory is a class of epistemic states or states of belief. The intended interpretation is that such a model is a representation of a person's knowledge and belief at a certain point in time [...]. However, the epistemic states here are not seen as psychological entities; they are presented as rational idealizations of psychological states. This means that a state in a computer program may also be seen as a model of an epistemic state."

Un état épistémique est donc une idéalisation de l'état cognitif d'un agent. Mais la façon de représenter cette idéalisation peut donc être quelconque. Dans les cadres probabilistes c'est une mesure de probabilité. Dans le cadre AGM classique c'est typiquement une base de connaissance ou un ensemble de mondes possibles.

Les propositions pour la révision itérée ont souligné la nécessité d'utiliser une information plus complexe qu'un simple ensemble de propositions pour modéliser l'état cognitif d'un agent. Cette information plus complexe a été (malheureusement) nommée état épistémique, mais peut se réduire de manière équivalente à un pré-ordre sur les interprétations, à un ensemble de conditionnels, à un enracinement épistémique, ou à un historique de révision.

On peut néanmoins imaginer des états épistémiques plus complexes encore que ces derniers. Nous en donnons un exemple section 3.5. Pour éviter les confusions entre les deux acceptions, on nommera Etat Epistémique (avec des majuscules) la notion de modélisation d'état cognitif et état épistémique (avec des minuscules) les Etats Epistémiques particuliers utilisés dans les divers travaux en révision itérée.

Pour pouvoir étudier les propriétés logiques d'un état épistémique, il est nécessaire de pouvoir isoler les connaissances associées à un état épistémique. On suppose donc l'existence d'une fonction de projection qui associe, à un état épistémique, une base de connaissance. Donc un état épistémique peut être vu comme une base de connaissance représentant la partie logique (base de connaissance) de l'état épistémique, et une information additionnelle, que nous appelons *ensemble conditionnel*, qui représente les croyances, la stratégie de révision de

l'agent. C'est-à-dire que cette information code la confiance relative de l'agent en les connaissances qu'il ne croit pas actuellement (i.e. qui ne font pas partie de sa base de connaissance) et définit donc les nouvelles connaissances de l'agent si ses connaissances actuelles sont remises en doute. Cette information additionnelle peut prendre différentes formes tel qu'un pré-ordre sur les interprétations, un ensemble d'assertions conditionnelles, un enracinement épistémique, un historique de révision, etc.

Lorsque l'on considère des bases de connaissance comme Etats Epistémiques, on dispose d'une syntaxe et d'une sémantique claire. En revanche, lorsque l'on utilise des états épistémiques, on ne dispose pas de formalisation précise. Pour être précis nous proposons donc une telle formalisation ici.

Soit $\mathcal{E}$ l'ensemble des Etats Epistémiques. Un élément particulier de $\mathcal{E}$ est Ξ , qui dénote l'Etat Epistémique correspondant à un état cognitif vierge, c'est-à-dire l'Etat Epistémique qui ne contient aucune connaissance et aucune croyance (i.e. aucune stratégie de révision).

On définit une fonction de projection π . Cette fonction permet d'obtenir la base de connaissance associée à un Etat Epistémique.

Définition 37 Une fonction de projection π est une fonction qui associe à chaque Etat Epistémique une base de connaissance $\pi : \mathcal{E} \rightarrow \mathcal{L}$ et telle que $\pi(\Xi) = K_{\top}$.

Ces Etats Epistémiques, en tant qu'idéalisations d'états cognitifs, ont donc très peu de structure. Pour pouvoir manipuler ces états nous avons besoin de contraindre la nature de ces objets. C'est ce que nous faisons à présent.

Soit $\bullet$ une fonction, nommée constructeur, qui, à un Etat Epistémique et à une formule, associe un Etat Epistémique : $\bullet : \mathcal{E} \times \mathcal{L} \rightarrow \mathcal{E}$

On définit alors l'ensemble des Etats Epistémiques librement engendrés par $\bullet$ et Ξ , noté $\mathcal{EE}$, que l'on nomme l'ensemble des états épistémiques. Cet ensemble est obtenu par induction à partir d'un constructeur et de l'Etat Epistémique Ξ .

Définition 38 Tout état épistémique Φ de $\mathcal{EE}$ est construit par application des deux règles suivantes un nombre fini de fois :

1. $\Xi \in \mathcal{EE}$
2. Si $\Phi \in \mathcal{EE}$ et $\varphi \in \mathcal{L}$, alors $(\Phi \bullet \varphi) \in \mathcal{EE}$

C'est-à-dire que, syntaxiquement, un état épistémique est une séquence (finie) de bases de connaissance. L'état épistémique Ξ dénotant la séquence vide.

Lorsque $\Phi = (\dots (\Xi \bullet \varphi_1) \bullet \dots \bullet \varphi_n)$, on omettra les parenthèses marquant l'associativité à gauche de $\bullet$ et on écrira simplement $\Phi = \Xi \bullet \varphi_1 \bullet \dots \bullet \varphi_n$. L'état épistémique initial d'un agent est toujours Ξ . Si l'on désire considérer un agent dans un état épistémique initial quelconque, il suffit de partir de l'état Ξ et de reconstruire l'état initial. Comme tout état épistémique est initié par Ξ , on omettra cet état initial, i.e. si $\Phi = \Xi \bullet \varphi_1 \bullet \dots \bullet \varphi_n$, on notera simplement $\Phi = \varphi_1 \bullet \dots \bullet \varphi_n$.

On peut alors définir l'équivalence et l'égalité entre états épistémiques.

Définition 39 Deux états épistémiques sont équivalents, noté $\Phi \leftrightarrow \Phi'$, si et seulement si leurs bases de connaissance sont équivalentes $\pi(\Phi) \leftrightarrow \pi(\Phi')$.

Avec cette définition syntaxique des états épistémiques, on peut à présent donner une définition plus précise de l'égalité entre états épistémiques que celle donnée section 2.1.1 comme identité des deux états. La notion d'égalité d'états épistémiques n'est pas une identité syntaxique mais une identité d'états cognitifs. En effet, il y a plusieurs façons d'atteindre un même état épistémique. Donc si $\Phi = \varphi_1 \bullet \dots \bullet \varphi_n$ et $\Phi' = \varphi_1 \bullet \dots \bullet \varphi_n$, alors ces deux états épistémiques sont égaux. Mais l'implication inverse n'est pas vraie.

On peut donner deux définitions de l'égalité d'états épistémiques. La première caractérise l'égalité comme une équivalence extensionnelle, la seconde comme une équivalence dynamique.

Définition 40 Soient deux états épistémiques Φ et Φ' , ces deux états épistémiques sont égaux, noté $\Phi = \Phi'$, si et seulement si pour toute formule φ , $\Phi \bullet \varphi \leftrightarrow \Phi' \bullet \varphi$.

Cette définition tente de capturer le fait que les états cognitifs représentés par les deux états épistémiques sont les mêmes. La deuxième définition suit la même idée mais en insistant sur l'idée d'équivalence dynamique.

Définition 41 Soient deux états épistémiques Φ et Φ' , ces deux états épistémiques sont égaux, noté $\Phi = \Phi'$, si et seulement si pour toute séquence de révision $\varphi_1 \bullet \dots \bullet \varphi_n$, $\Phi \bullet \varphi_1 \bullet \dots \bullet \varphi_n \leftrightarrow \Phi' \bullet \varphi_1 \bullet \dots \bullet \varphi_n$.

On peut remarquer que la définition d'égalité de la définition 41 est strictement plus forte que celle de la définition 40. En fait, avec les sémantiques considérées, ces deux définitions coïncident. Mais on prendra pour la suite comme définition de l'égalité la définition 40 puisqu'elle est plus aisée à vérifier pour les preuves.

Pour donner une sémantique à un état épistémique il suffit de définir une interprétation pour la structure $\langle \mathcal{EE}, \mathcal{L}, \bullet, \pi \rangle$. Cette structure sera également nommée *espace épistémique*.

On définit une fonction d'interprétation pour chaque élément de la structure. Soient deux ensembles E et L , on définit :

- $\mathcal{I}_{\mathcal{EE}} : \mathcal{EE} \rightarrow E$
- $\mathcal{I}_{\mathcal{L}} : \mathcal{L} \rightarrow L$
- $\mathcal{I}_{\bullet} = \circ$ où $\circ : E \times L \rightarrow E$
- $\mathcal{I}_{\pi} = Bel$ où $Bel : E \rightarrow L$ est une fonction telle que $\pi(\Phi) = \varphi$ ssi $Bel(\mathcal{I}_{\mathcal{EE}}(\Phi)) = \mathcal{I}_{\mathcal{L}}(\varphi)$

Dans la suite on supposera, pour simplifier, que dans toutes les interprétations le langage reste le même, i.e. $L = \mathcal{L}$ et $\mathcal{I}_{\mathcal{L}} = id$.

Une interprétation $\langle \mathcal{I}_{\mathcal{EE}}, \mathcal{I}_{\mathcal{L}}, \mathcal{I}_{\bullet}, \mathcal{I}_{\pi} \rangle$ est modèle de la structure $\langle \mathcal{EE}, \mathcal{L}, \bullet, \pi \rangle$ si

1. $\Phi = \Phi'$ ssi $\mathcal{I}_{\mathcal{EE}}(\Phi) = \mathcal{I}_{\mathcal{EE}}(\Phi')$
2. $\pi(\Phi) \vdash \mu$ ssi $Bel(\mathcal{I}_{\mathcal{EE}}(\Phi)) \vdash \mathcal{I}_{\mathcal{L}}(\mu)$
3. $(\pi(\Phi) \wedge \mu) \leftrightarrow \varphi$ ssi $(Bel(\mathcal{I}_{\mathcal{EE}}(\Phi)) \wedge \mathcal{I}_{\mathcal{L}}(\mu)) \leftrightarrow \mathcal{I}_{\mathcal{L}}(\varphi)$

On peut alors voir que les différentes notions proposées pour les états épistémiques correspondent à la notion de modèle sémantique de notre syntaxe.

Par exemple, si on fait correspondre E avec l'ensemble des pré-ordres totaux, L avec l'ensemble des ensembles d'interprétations, $\circ$ une fonction qui, à un pré-ordre total et un ensemble de modèle, associe un pré-ordre total, et Bel avec la fonction min, alors on peut retrouver les opérateurs de Darwiche et Pearl, celui de Boutilier et ceux de Lehmann.

On dit qu'une interprétation est modèle de l'ensemble d'axiomes A si et seulement si cette interprétation est modèle de la structure et les axiomes A sont vérifiés par la structure.

On pose les notations suivantes : $\Phi \vdash \mu$, $\Phi \wedge \mu$ et $I \models \Phi$ dénotent respectivement $\pi(\Phi) \vdash \mu$, $\pi(\Phi) \wedge \mu$ et $I \models \pi(\Phi)$.

Lorsque $\Phi = \varphi_1 \bullet \dots \bullet \varphi_n$ on notera $\Phi \bullet \varphi$ l'état épistémique correspondant à l'addition de φ à la séquence de révisions, i.e. $\Phi \bullet \varphi = \varphi_1 \bullet \dots \bullet \varphi_n \bullet \varphi$. De même lorsque $\Phi' = \varphi'_1 \bullet \dots \bullet \varphi'_n$, on notera par abus $\Phi \bullet \Phi'$ la séquence $\varphi_1 \bullet \dots \bullet \varphi_n \bullet \varphi'_1 \bullet \dots \bullet \varphi'_n$.

Les postulats AGM pour états épistémiques sont donc alors les suivants :

- (R*1) $\Phi \bullet \mu \vdash \mu$
- (R*2) Si $\Phi \wedge \mu$ est consistant, alors $\Phi \bullet \mu \leftrightarrow \Phi \wedge \mu$
- (R*3) Si μ est consistant, alors $\Phi \bullet \mu$ est consistant
- (R*4) Si $\Phi_1 = \Phi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\Phi_1 \bullet \mu_1 \leftrightarrow \Phi_2 \bullet \mu_2$
- (R*5) $(\Phi \bullet \mu) \wedge \varphi \vdash \Phi \bullet (\mu \wedge \varphi)$
- (R*6) Si $(\Phi \bullet \mu) \wedge \varphi$ est consistant, alors $\Phi \bullet (\mu \wedge \varphi) \vdash (\Phi \bullet \mu) \wedge \varphi$

La formulation est donc la même que celle proposée par Darwiche et Pearl, mais on travaille ici à un niveau purement syntaxique. Ces postulats peuvent donc être considérés comme les axiomes d'un espace épistémique particulier. Un opérateur de révision est alors un modèle de l'espace épistémique. Le théorème de représentation pour les opérateurs satisfaisant les postulats (R*1)-(R*6) peut alors s'énoncer ainsi :

Théorème 40 *Soit un opérateur $\circ$ défini par une interprétation $\langle \mathcal{I}_{\mathcal{E}}, \mathcal{I}_{\mathcal{L}}, \circ, Bel \rangle$. $\circ$ est un modèle de (R*1)-(R*6) si et seulement si il existe un assignement fidèle tel que*

$$Mod(Bel(\mathcal{I}_{\mathcal{E}}(\Phi) \circ \mu)) = \min(Mod(\mu), \leq_{\Phi})$$

Un opérateur de révision peut donc être considéré comme une interprétation particulière de l'espace épistémique.

Finalement, on peut remarquer que si on identifie $\mathcal{E}\mathcal{E}$ à $\mathcal{L}$ et π à l'identité, on recouvre le cadre AGM classique en associant $\bullet$ à un opérateur de révision.

3.2 Révision à mémoire

Nous donnons à présent la caractérisation logique des opérateurs de révision à mémoire :

Définition 42 *Un espace épistémique est dit à mémoire si les axiomes suivants sont satisfaits :*

- (H1) $\Phi \bullet \varphi \vdash \varphi$

- (H2) Si $\Phi \wedge \varphi$ est consistant, alors $\Phi \bullet \varphi \leftrightarrow \Phi \wedge \varphi$
(H3) Si φ est consistant, alors $\Phi \bullet \varphi$ est consistant
(H4) Si $\Phi_1 = \Phi_2$ et $\varphi_1 \leftrightarrow \varphi_2$, alors $\Phi_1 \bullet \varphi_1 = \Phi_2 \bullet \varphi_2$
(H5) $(\Phi \bullet \varphi) \wedge \mu \vdash \Phi \bullet (\varphi \wedge \mu)$
(H6) Si $(\Phi \bullet \varphi) \wedge \mu$ est consistant, alors $\Phi \bullet (\varphi \wedge \mu) \vdash (\Phi \bullet \varphi) \wedge \mu$
(H7) $\Phi \bullet \Phi' \leftrightarrow \Phi \bullet \pi(\Phi')$

(H1-H6) sont exactement les postulats (R*1-R*6), excepté (H4) qui est plus fort que (R*4). Le postulat (H7) exprime la confiance forte en la nouvelle information en exigeant une associativité à droite des opérateurs à mémoire (au niveau des bases de connaissance).

Nous avons laissé les postulats (H5) et (H6) pour montrer la similitude de ces opérateurs avec les postulats (R*1)-(R*6) mais, en fait, il est facile de voir que les postulats (H5) et (H6) sont une conséquence de (H7) :

Théorème 41 *Les postulats (H1), (H2), (H4) et (H7) impliquent les postulats (H5) et (H6)*

Preuve : Lorsque $(\Phi \bullet \varphi) \wedge \mu$ n'est pas consistant alors (H5) et (H6) sont vérifiés trivialement. Lorsque $(\Phi \bullet \varphi) \wedge \mu$ est consistant on a alors que $\varphi \wedge \mu$ est consistant (car si ce n'est pas le cas i.e. $\varphi \wedge \mu \vdash \perp$, d'après (H1) $\Phi \bullet \varphi \vdash \varphi$ donc $(\Phi \bullet \varphi) \wedge \mu \vdash \perp$. Contradiction), donc comme $\varphi \wedge \mu$ est consistant on a par (H2) $\varphi \bullet \mu \leftrightarrow \varphi \wedge \mu$. Par hypothèse et par (H2) on a $(\Phi \bullet \varphi) \wedge \mu \leftrightarrow (\Phi \bullet \varphi) \bullet \mu$. De (H7) on a $\Phi \bullet \varphi \bullet \mu \leftrightarrow \Phi \bullet \pi(\varphi \bullet \mu)$, c'est-à-dire par (H4) et les deux faits précédents $(\Phi \bullet \varphi) \wedge \mu \leftrightarrow \Phi \bullet (\varphi \wedge \mu)$. Donc (H5) et (H6) sont vérifiés. $\square$

Définition 43 *Les opérateurs de révision à mémoire sont les interprétations de l'espace épistémique qui sont modèles de (H1), (H2), (H3), (H4) et (H7).*

Nous allons donner à présent un théorème de représentation pour ces opérateurs à mémoire en termes de familles de pré-ordres sur les interprétations. Cela donnera une définition plus constructive de ces opérateurs.

Nous définissons d'abord ce qu'est un assignement conservatif.

Définition 44 *Une fonction qui associe à chaque état épistémique Φ un pré-ordre $\leq_\Phi$ sur les interprétations est un assignement conservatif si et seulement si :*

1. Si $I \models \Phi$ et $J \models \Phi$, alors $I \simeq_\Phi J$
2. Si $I \models \Phi$ et $J \not\models \Phi$, alors $I <_\Phi J$
3. $\Phi_1 = \Phi_2$ ssi $\leq_{\Phi_1} = \leq_{\Phi_2}$
4. Si $I <_\varphi J$, alors $I <_{\Phi \bullet \varphi} J$
5. Si $I \simeq_\varphi J$, alors $I \leq_{\Phi \bullet \varphi} J$ ssi $I \leq_\Phi J$

Il est possible de renforcer les propriétés de l'assignement conservatif en ajoutant la propriété suivante :

6. Si $I \not\models \varphi$ et $J \not\models \varphi$, alors $I \leq_{\Phi \bullet \varphi} J$ ssi $I \leq_\Phi J$

Cette condition est la condition (CR2) de Darwiche and Pearl. Nous appellerons de tels assignements des *assignements conservatifs forts*.

Nous avons le théorème de représentation suivant :

Théorème 42 *Soit un opérateur de révision $\circ$. $\circ$ est modèle de (H1)-(H7) si et seulement si il existe un assignement conservatif qui associe à chaque état épistémique Φ un pré-ordre total $\leq_\Phi$ tel que :*

$$Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \leq_\Phi)$$

Preuve :

(Seulement si) Soit $\circ$ un opérateur qui satisfait (H1)-(H7). Soit l'assignement qui à chaque état épistémique Φ associe un pré-ordre $\leq_\Phi$ tel que $\forall I, J \in \mathcal{W}$, $I \leq_\Phi J$ si et seulement si $I \models \Phi \circ \varphi_{\{I, J\}}$. Nous montrons que $\leq_\Phi$ est un pré-ordre total, que l'assignement ainsi défini est un assignement conservatif et que $Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \leq_\Phi)$.

Nous montrons d'abord que $\leq_\Phi$ est un pré-ordre total :

Réflexivité : De (H1) et (H3) on a que $\forall I$ $\Phi \circ \varphi_{\{I\}} \leftrightarrow \varphi_{\{I\}}$, c'est-à-dire $I \leq_\Phi I$.

Transitivité : Les cas où au moins deux des trois interprétations coïncident sont triviaux. Supposons donc que les trois interprétations sont différentes et supposons que $I \leq_\Phi J$ et $J \leq_\Phi L$. On montre alors $I \leq_\Phi L$.

Cas 1 : $\Phi \circ \varphi_{\{I, J, L\}} \wedge \varphi_{\{I, L\}}$ n'est pas consistant. Alors $\Phi \circ \varphi_{\{I, J, L\}} \leftrightarrow \varphi_{\{J\}}$. Donc $\Phi \circ \varphi_{\{I, J, L\}} \wedge \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{J\}}$. Par (H5) et (H6) il s'ensuit que $\Phi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{J\}}$, donc par définition $I \not\leq_\Phi J$. Contradiction.

Cas 2 : $\Phi \circ \varphi_{\{I, J, L\}} \wedge \varphi_{\{I, L\}}$ est consistant. Par l'absurde. Supposons que $I \not\leq_\Phi L$, i.e. de (H1) et (H3), $\Phi \circ \varphi_{\{I, L\}} \leftrightarrow \varphi_{\{L\}}$. Par (H5) et (H6) on a $\Phi \circ \varphi_{\{I, J, L\}} \wedge \varphi_{\{I, L\}} \leftrightarrow \Phi \circ \varphi_{\{I, L\}} \leftrightarrow \varphi_{\{L\}}$. Donc $I \not\models \Phi \circ \varphi_{\{I, J, L\}}$. De (H1) et (H3) on a soit $\Phi \circ \varphi_{\{I, J, L\}} \leftrightarrow \varphi_{\{J, L\}}$ ou $\Phi \circ \varphi_{\{I, J, L\}} \leftrightarrow \varphi_{\{L\}}$. Dans le premier cas on obtient $\Phi \circ \varphi_{\{I, J, L\}} \wedge \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{J\}}$, par (H5) et (H6) on conclut $\Phi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{J\}}$, c'est-à-dire $I \not\leq_\Phi J$. Contradiction. Dans le second cas on obtient $\Phi \circ \varphi_{\{I, J, L\}} \wedge \varphi_{\{J, L\}} \leftrightarrow \varphi_{\{L\}}$, par (H5) et (H6) on conclut $\Phi \circ \varphi_{\{J, L\}} \leftrightarrow \varphi_{\{L\}}$, c'est-à-dire $J \not\leq_\Phi L$. Contradiction.

Totalité : $\forall I, J \in \mathcal{W}$, de (H1) on a que $\Phi \circ \varphi_{\{I, J\}} \vdash \varphi_{\{I, J\}}$ et de (H3) que $\Phi \circ \varphi_{\{I, J\}} \not\vdash \perp$, donc soit $I \models \Phi \circ \varphi_{\{I, J\}}$ ou $J \models \Phi \circ \varphi_{\{I, J\}}$, i.e. $I \leq_\Phi J$ ou $J \leq_\Phi I$.

On vérifie à présent les conditions de l'assignement conservatif :

1. Si $I \models \Phi$ et $J \models \Phi$, alors par (H2) on a $\Phi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I, J\}}$. Ce qui par définition donne $I \leq_\Phi J$ et $J \leq_\Phi I$, donc $I \simeq_\Phi J$.
2. Si $I \models \Phi$ et $J \not\models \Phi$, alors par (H2) on obtient $\Phi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I\}}$. C'est-à-dire par définition $I \leq_\Phi J$ et $J \not\leq_\Phi I$, donc $I <_\Phi J$.
3. Direct par (H4).
4. Si $I <_\varphi J$, alors par définition $\varphi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I\}}$ et de (H7) $\Phi \circ \varphi \circ \varphi_{\{I, J\}} \leftrightarrow \Phi \circ Bel(\varphi \circ \varphi_{\{I, J\}})$. De plus de (H1) et (H3) on a $\Phi \circ Bel(\varphi \circ \varphi_{\{I, J\}}) \leftrightarrow \varphi_{\{I\}}$, ce qui donne $\Phi \circ \varphi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I\}}$. C'est-à-dire par définition $I <_{\Phi \circ \varphi} J$.

5. Si $I \simeq_\varphi J$, alors par définition $\varphi \circ \varphi_{\{I,J\}} \leftrightarrow \varphi_{\{I,J\}}$. D'où $\Phi \circ \text{Bel}(\varphi \circ \varphi_{\{I,J\}}) \leftrightarrow \Phi \circ \varphi_{\{I,J\}}$. De (H7) on a $\Phi \circ \varphi \circ \varphi_{\{I,J\}} \leftrightarrow \Phi \circ \text{Bel}(\varphi \circ \varphi_{\{I,J\}})$, ce qui donne $\Phi \circ \varphi \circ \varphi_{\{I,J\}} \leftrightarrow \Phi \circ \varphi_{\{I,J\}}$. D'où directement par définition: $I \leq_{\Phi \circ \varphi} J$ si et seulement si $I \leq_\Phi J$.

Finalement on montre que $\text{Mod}(\Phi \circ \varphi) = \min(\text{Mod}(\varphi), \leq_\Phi)$. D'abord pour l'inclusion $\text{Mod}(\Phi \circ \varphi) \subseteq \min(\text{Mod}(\varphi), \leq_\Phi)$ on suppose que $I \models \Phi \circ \varphi$ et par l'absurde supposons que $I \notin \min(\text{Mod}(\varphi), \leq_\Phi)$. Donc on peut trouver un $J \models \varphi$ tel que $J <_\Phi I$. alors $I \not\models \Phi \circ \varphi_{\{I,J\}}$. Comme $\varphi_{\{I,J\}} \vdash \varphi$, de (H5) et (H6) on a $(\Phi \circ \varphi) \wedge \varphi_{\{I,J\}} \leftrightarrow \Phi \circ \varphi_{\{I,J\}}$. Mais $I \not\models \Phi \circ \varphi_{\{I,J\}}$, d'où $I \not\models \Phi \circ \varphi$. Contradiction.

Inversement on suppose $I \in \min(\text{Mod}(\varphi), \leq_\Phi)$. On souhaite montrer que $I \models \Phi \circ \varphi$. Comme $I \in \min(\text{Mod}(\varphi), \leq_\Phi)$, alors $\forall J \models \varphi \ I \leq_\Phi J$. C'est-à-dire $I \models \Phi \circ \varphi_{\{I,J\}}$. Comme $(\Phi \circ \varphi) \wedge \varphi_{\{I,J\}}$ est consistant, de (H5) et (H6) on déduit $(\Phi \circ \varphi) \wedge \varphi_{\{I,J\}} \leftrightarrow \Phi \circ \varphi_{\{I,J\}}$. Et comme $I \models \Phi \circ \varphi_{\{I,J\}}$, $I \models \Phi \circ \varphi$.

(Si) On considère l'assignement conservatif qui associe à chaque état épistémique Φ un pré-ordre total $\leq_\Phi$ et on suppose que l'opérateur $\circ$ vérifie $\text{Mod}(\Phi \circ \varphi) = \min(\text{Mod}(\varphi), \leq_\Phi)$. On montre alors que $\circ$ satisfait (H1-H7).

- (H1) Par définition $\text{Mod}(\Phi \circ \varphi) \subseteq \text{Mod}(\varphi)$, donc $\Phi \circ \varphi \vdash \varphi$.
- (H2) Supposons que $\Phi \wedge \varphi$ est consistant. On veut montrer que $\min(\text{Mod}(\varphi), \leq_\Phi) = \text{Mod}(\Phi \wedge \varphi)$. Notons d'abord que si $I \models \Phi$ alors d'après les conditions 1 et 2 $I \in \min(\mathcal{W}, \leq_\Phi)$. Donc si $I \models \Phi \wedge \varphi$ alors $I \in \min(\text{Mod}(\varphi), \leq_\Phi)$. D'où $\min(\text{Mod}(\varphi), \leq_\Phi) \supseteq \text{Mod}(\Phi \wedge \varphi)$. Pour l'autre inclusion on considère $I \in \min(\text{Mod}(\varphi), \leq_\Phi)$. Vers une contradiction supposons que $I \not\models \Phi \wedge \varphi$. Puisque $I \not\models \Phi$ par la condition 2 on a que $\forall J \models \Phi \ J <_\Phi I$. En particulier $\forall J \models \Phi \wedge \varphi \ J <_\Phi I$. Donc $I \notin \min(\text{Mod}(\varphi), \leq_\Phi)$. Contradiction.
- (H3) Si φ est consistant, alors $\text{Mod}(\varphi) \neq \emptyset$ et, comme il n'y a qu'un nombre fini d'interprétations, il n'y a pas de chaînes infinies d'inégalités strictes, donc $\min(\text{Mod}(\varphi), \leq_\Phi) \neq \emptyset$. Donc $\Phi \circ \varphi$ est consistant.
- (H4) Direct d'après la condition 3.
- (H7) On a directement par définition $\Phi' \leftrightarrow \perp$ si et seulement si $\Phi' = \Phi'' \circ \perp$. De là on a directement $\Phi \circ \Phi'$ est inconsistant si et seulement si $\Phi \circ \text{Bel}(\Phi')$ est inconsistant. Ensuite, si $\Phi \circ \Phi'$ est consistant. Par les conditions 1 et 2 on a $I \models \Phi \circ \Phi'$ si et seulement si $I \models \min(\mathcal{W}, \leq_{\Phi \circ \Phi'})$. Mais $I \in \min(\mathcal{W}, \leq_{\Phi \circ \Phi'})$ peut être réécrit en utilisant les conditions 4 et 5 en $I \in \min(\min(\mathcal{W}, \leq_{\Phi'}), \leq_\Phi)$. Avec les conditions 1 et 2 on obtient $\min(\min(\mathcal{W}, \leq_{\Phi'}), \leq_\Phi) = \min(\text{Bel}(\Phi'), \leq_\Phi)$, et par définition, $\min(\text{Bel}(\Phi'), \leq_\Phi) = \text{Bel}(\Phi \circ \text{Bel}(\Phi'))$. On obtient donc $I \models \Phi \circ \Phi'$ si et seulement si $I \models \Phi \circ \text{Bel}(\Phi')$.

□

Le théorème précédent dit simplement que les bases de connaissance obtenues après chaque révision par la méthode syntaxique et l'interprétation sémantique coïncident. Contrairement au théorème de représentation classique (cf section 1.2.4) où l'opérateur et la nouvelle base de connaissance sont parfaitement définis par l'assignement correspondant, on ne définit pas dans ce théorème, à proprement parler, l'opérateur syntaxique $\circ$; on se contente de le caractériser.

Mais on peut également énoncer ce théorème directement en terme de sémantique associé au constructeur $\circ$. Cette formulation montre que ce théorème peut tout de même être considéré comme un théorème de représentation.

Théorème 43 *Soit une interprétation de la structure $\langle \mathcal{EE}, \mathcal{L}, \circ, \pi \rangle$ qui envoie $\mathcal{EE}$ dans l'ensemble des pré-ordres totaux. Cette interprétation est un modèle des axiomes (H1)-(H7) si et seulement si l'interprétation est construite à l'aide d'un assignement conservatif (i.e. si $\Phi = \varphi_1 \bullet \dots \bullet \varphi_n$, alors le pré-ordre $\leq_\Phi = \mathcal{I}_{\mathcal{EE}}(\Phi)$ vérifie les conditions 1 – 5).*

Nous allons à présent nous intéresser à quelques propriétés logiques de cette famille d'opérateurs de révision.

Théorème 44 *Si un opérateur $\circ$ est un modèle de (H1)-(H6), alors il satisfait (H7) si et seulement si il satisfait les postulats suivants :*

(H'7) *Si $\varphi \bullet \mu \leftrightarrow \mu$, alors $\Phi \bullet \varphi \bullet \mu \leftrightarrow \Phi \bullet \mu$*

(H'8) *Si $\Phi' \vdash \mu$, alors $\Phi \bullet \Phi' \vdash \mu$*

Preuve :

(H7) implique directement (H'7), et (H7) et (H1) impliquent (H'8). Pour montrer que (H'7) et (H'8) impliquent (H7) nous montrons que l'ensemble de postulats (H1)-(H6), (H'7) et (H'8) correspondent aux conditions de l'assignement conservatif.

Soit $\circ$ un opérateur qui satisfait (H1)-(H6) et (H'7)-(H'8). On définit un assignement qui associe à chaque état épistémique Φ un pré-ordre $\leq_\Phi$ défini par $\forall I, J \in \mathcal{W}$, $I \leq_\Phi J$ si et seulement si $I \models \Phi \circ \varphi_{\{I, J\}}$. La preuve est la même que celle du théorème 42, il reste simplement à montrer que les conditions 4 et 5 de l'assignement sont satisfaites :

4. Si $I <_\varphi J$, alors par définition $\varphi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I\}}$ et d'après (H'8) si $\varphi \circ \varphi_{\{I, J\}} \vdash \varphi_{\{I\}}$, alors $\Phi \circ \varphi \circ \varphi_{\{I, J\}} \vdash \varphi_{\{I\}}$. De plus d'après (H3) on a que $\Phi \circ \varphi \circ \varphi_{\{I, J\}} \not\vdash \perp$, donc on a $\Phi \circ \varphi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I\}}$. Ce qui par définition est $I <_{\Phi \circ \varphi} J$.
5. si $I \simeq_\varphi J$, alors par définition $\varphi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I, J\}}$. D'après (H'7) comme $\varphi \circ \varphi_{\{I, J\}} \leftrightarrow \varphi_{\{I, J\}}$, alors $\Phi \circ \varphi \circ \varphi_{\{I, J\}} \leftrightarrow \Phi \circ \varphi_{\{I, J\}}$. C'est-à-dire par définition : $I \leq_{\Phi \circ \varphi} J$ si et seulement si $I \leq_\Phi J$.

□

Le postulat (H7') est une généralisation du postulat (C1) de Darwiche et Pearl. Il dit que si les mondes possibles satisfaisant μ sont tous aussi crédibles pour φ , alors réviser un état épistémique par φ et μ donne la même base de connaissance que si l'on révisé uniquement par μ .

Le postulat (H8') est une version affaiblie de (H7), puisqu'il assure que si μ est impliqué par un certain historique de révisions Φ' alors, μ sera impliqué par tout historique ayant Φ' comme suffixe. Ce postulat exprime l'importance forte de la dernière information.

Il est également intéressant de noter que :

Théorème 45 *les postulats (H1)-(H7) impliquent la propriété suivante :*

(C) *Si $\varphi \wedge \mu$ est consistant, alors $\Phi \bullet \varphi \bullet \mu \leftrightarrow \Phi \bullet (\varphi \wedge \mu)$*

Preuve : D'après l'hypothèse $\varphi \wedge \mu$ est consistant, donc de (H2) on déduit $\varphi \bullet \mu \leftrightarrow \varphi \wedge \mu$. On a donc par (H4) $\forall \Phi \Phi \bullet (\varphi \bullet \mu) \leftrightarrow \Phi \bullet (\varphi \wedge \mu)$, ce qui d'après (H7) donne $\Phi \bullet \varphi \bullet \mu \leftrightarrow \Phi \bullet (\varphi \wedge \mu)$.

□

(C) exprime le fait que réviser successivement par deux informations compatibles revient à réviser par leur conjonction. Cela est proche du postulat de *Conjonction* proposé par Nayak *et al.*, mais (C) est plus faible que *Conjonction*, puisqu'il n'implique que l'équivalence entre les états épistémiques alors que *Conjonction* demande l'égalité. De ce fait, les deux états épistémiques ont la même base de connaissance mais peuvent avoir des ensembles conditionnels différents.

Théorème 46 *Un opérateur de révision à mémoire satisfait les postulats (C1), (C3) et (C4).*

Preuve : (C1) est une conséquence directe de (H7') puisque si $\alpha \vdash \mu$ alors d'après (H2) $\alpha \circ \mu \leftrightarrow \mu$ et d'après (H7') $\Phi \circ \mu \circ \alpha \leftrightarrow \Phi \circ \alpha$.

Pour (C3) supposons que $\Phi \circ \alpha \vdash \mu$, c'est-à-dire en particulier que $\alpha \wedge \mu$ est consistant, donc d'après la condition (C) $\Phi \circ \alpha \circ \mu \leftrightarrow \Phi \circ (\alpha \wedge \mu) \leftrightarrow \Phi \circ \mu \circ \alpha$. Or d'après (H'8) et (H1) $\Phi \circ \alpha \circ \mu \vdash \mu$ donc $\Phi \circ \mu \circ \alpha \vdash \mu$.

La preuve de (C4) est similaire. Supposons que $\Phi \circ \alpha \not\vdash \neg \mu$, c'est-à-dire en particulier $\alpha \wedge \mu \not\vdash \perp$ donc d'après la condition (C) $\Phi \circ (\alpha \wedge \mu) \leftrightarrow \Phi \circ \mu \circ \alpha$. Or $\alpha \wedge \mu \vdash \mu$ donc $\Phi \circ \mu \circ \alpha \not\vdash \neg \mu$.

□

Un opérateur de révision à mémoire ne satisfait généralement pas le postulat (C2). Nous avons en fait le théorème de représentation suivant :

Théorème 47 *Un opérateur de révision $\circ$ satisfait les postulats (H1-H7) et (C2) si et seulement si il existe un assignement conservatif fort qui associe à chaque état épistémique Φ un pré-ordre total $\leq_\Phi$ tel que :*

$$Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \leq_\Phi)$$

Preuve : La majeure partie de ce théorème est donnée par le théorème 42. Il reste simplement à montrer que le postulat (C2) correspond à la condition 6 sur l'assignement conservatif fort.

(Seulement si)

6. Si $I \not\models \varphi$ et $J \not\models \varphi$, alors $\varphi_{\{I,J\}} \vdash \neg \varphi$ donc d'après (C2) $\Phi \circ \varphi \circ \varphi_{\{I,J\}} \leftrightarrow \Phi \circ \varphi_{\{I,J\}}$. Ce qui par définition donne $I \leq_{\Phi \bullet \varphi} J$ ssi $I \leq_\Phi J$

(Si)

- (C2) Supposons que $\varphi \vdash \neg \mu$. Soit $I, J \models \varphi$, alors $I, J \not\models \mu$. Donc d'après la condition 6 on a $I \leq_{\Phi \bullet \mu} J$ ssi $I \leq_\Phi J$ et d'après la condition 5 $I \leq_{\Phi \bullet \mu \bullet \varphi} J$ ssi $I \leq_{\Phi \bullet \varphi} J$. Ce qui par définition nous donne $\Phi \circ \mu \circ \varphi \leftrightarrow \Phi \circ \varphi$ □

Ce théorème annonce qu'une sous-classe d'opérateurs de révision à mémoire satisfait (C2). Nous verrons bientôt qu'il n'existe qu'un seul opérateur (l'opérateur basique) qui satisfait cette condition.

Le postulat (C2) a été montré inconsistent avec les postulats AGM par Freund et Lehmann [FL94]. Darwiche et Pearl ont donc modifié les postulats AGM, les énonçant en termes d'états épistémiques, et ont de ce fait supprimé cette contradiction. Mais nous soutenons que (C2) n'est pas toujours souhaitable pour autant et que satisfaire ce postulat peut mener à des résultats contre-intuitifs. Nous avons montré, section 2.1.2, que dans certains cas le postulat (C2) induit exactement la même sorte de mauvais comportement qu'il est censé prévenir.

3.3 Opérateurs de révision à mémoire et opérateurs AGM

Nous donnons dans cette section un exemple de construction des opérateurs de révision à mémoire à partir d'un assignement fidèle donné. Cela permet d'illustrer les motivations de la caractérisation logique. Il faut aussi noter que cette construction montre également (via les théorèmes de représentation) comment construire un opérateur de révision à mémoire à partir d'un opérateur de révision AGM usuel.

Nous associons à chaque base de connaissance φ un pré-ordre total $\leq_\varphi$ qui représente les modèles de la base de connaissance au niveau le plus bas et la crédibilité *a priori* des autres mondes possibles. Cela peut-être vu comme l'état épistémique qu'aurait l'agent s'il n'avait aucune connaissance et apprenait cette information. Cela doit être interprété comme la politique de révision de l'agent. Nous imposons très peu de contraintes à cet assignement, supposant que c'est un assignement fidèle normalisé, c'est-à-dire :

Définition 45 *Un assignement fidèle est dit normalisé s'il satisfait la propriété suivante :*

4. Si $\varphi = \perp$, alors $\forall I, J \in \mathcal{W} \ I \simeq_\varphi J$

Dans la formulation usuelle de Katsuno et Mendelzon, le pré-ordre associé à la base contradictoire est un pré-ordre quelconque. Mais, si notre but est de construire un système de connaissance qui peut supporter de manière satisfaisante une information contradictoire, il semble naturel de demander que celle-ci n'ait pas de conséquence sur les prochaines révisions. C'est ce qui motive la définition de ces assignements fidèles normalisés. Cette contrainte n'est pas nécessaire dans la définition des opérateurs de révision à mémoire mais nous la trouvons souhaitable pour la plupart des systèmes à base de connaissance. D'ailleurs, tous les exemples d'opérateurs que nous donnerons seront normalisés. Il faut remarquer que cette contrainte correspond exactement au postulat d'*absurdité* proposé par Nayak *et al.* .

Notons que nous utilisons ici deux définitions d'assignement fidèle: une sur les états épistémiques et l'autre sur les bases de connaissance. Cet assignement sur les bases de connaissance induit un assignement sur les états épistémiques de la façon suivante :

Définition 46 *Soit un assignement fidèle qui associe à chaque base de connaissance φ un pré-ordre $\leq_\varphi$. Soit un état épistémique $\Phi = \varphi_1 \bullet \dots \bullet \varphi_n$. On définit $\leq_\Phi$ comme :*

- Si $n = 1$, alors $\leq_\Phi = \leq_{\varphi_1}$

- Sinon $I \leq_{\Phi} J$ ssi $I <_{\varphi_n} J$ ou
 $I \simeq_{\varphi_n} J$ et $I \leq_{\varphi_1 \bullet \dots \bullet \varphi_{n-1}} J$

Le pré-ordre $\leq_{\Phi}$ est appelé ensemble conditionnel de Φ et la base de connaissance de Φ , notée $Bel(\Phi)$, est la formule (à équivalence logique près) telle que $Mod(Bel(\Phi)) = \min(\mathcal{W}, \leq_{\Phi})$ si φ_n est consistant, et $Bel(\Phi) = \perp$ sinon.

Donc le pré-ordre associé à un état épistémique est l'ordre lexicographique (inverse) sur la séquence des pré-ordres des bases de connaissance. Il est facile de voir que :

Théorème 48 Si la fonction qui associe, à chaque base de connaissance φ , un pré-ordre $\leq_{\varphi}$ est un assignement fidèle, alors la fonction qui associe à chaque état épistémique Φ un pré-ordre $\leq_{\Phi}$ selon la définition 46 est un assignement conservatif.

Corollaire 49 Lorsque $\leq_{\Phi}$ est défini grâce à la définition 46, l'opérateur de révision $\circ$ défini par $\mathcal{I}_{\mathcal{E}\mathcal{E}}(\Phi) = \leq_{\Phi}$ et $\leq_{\Phi} \circ \varphi = \leq_{\Phi \bullet \varphi}$ satisfait (H1)-(H7).

Les pré-ordres définis grâce à la définition 46 satisfont les propriétés de l'assignement conservatif. Inversement, lorsque l'on part d'un assignement conservatif, on retrouve directement l'ordre lexicographique de la définition 46 lorsque l'on part de l'assignement fidèle $\varphi \mapsto \leq_{\varphi}$ où $\leq_{\varphi} = \leq_{\exists \bullet \varphi}$.

Cela implique en particulier qu'un assignement conservatif est complètement déterminé par un unique assignement fidèle.

3.4 Quelques opérateurs de révision à mémoire

Nous donnons dans cette section quatre exemples d'opérateurs de révision à mémoire. Le premier est appelé opérateur basique à mémoire car il correspond à l'assignement le plus simple que l'on puisse définir. Nous montrons que, même avec un opérateur aussi simple, il est possible de construire des croyances complexes. Nous montrons également que cet opérateur est le seul à satisfaire (C2).

Nous donnons également un opérateur Dalal à mémoire qui permet d'illustrer le comportement des opérateurs de révision à mémoire lorsque les pré-ordres sont plus complexes.

Finalement, nous montrons comment généraliser les opérateurs de révision basés sur les OTP (Présentations Ordonnées de Théories) de Ryan [Rya94] de façon à obtenir des opérateurs de révision à mémoire.

3.4.1 Opérateur basique à mémoire

Soit l'assignement qui associe à chaque base de connaissance un pré-ordre de la façon suivante :

Définition 47 Soient une base de connaissance φ et deux interprétations I, J .

$$I \leq_{\varphi}^b J \text{ si et seulement si } \begin{cases} \varphi = \perp \\ \text{ou} \\ I \models \varphi \end{cases}$$

Nous obtenons ce que nous appelons un pré-ordre basique, c'est-à-dire un pré-ordre à (au plus) deux niveaux, avec les modèles de φ au niveau bas et les autres interprétations au niveau haut.

Il est facile de montrer que l'assignement qui associe à chaque base de connaissance φ un pré-ordre $\leq_\varphi^b$ est un assignement fidèle normalisé. Et nous définissons $\leq_\Phi^b$, de manière lexicographique inverse, avec la définition 46.

Définition 48 Soient un état épistémique Φ et une formule φ , l'opérateur de révision basique à mémoire est défini à partir de l'assignement fidèle normalisé $\leq_\varphi^b$ et de la définition 46. En particulier $Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \leq_\Phi^b)$.

Même avec cet ordre basique sur les bases de connaissance, il est possible de construire des états épistémiques très précis, grâce à l'historique de révisions. Nous allons illustrer le comportement de cet opérateur sur quelques exemples simples. Considérons un langage $\mathcal{L}$ avec seulement deux variables propositionnelles a et b . Voyons quelques exemples d'états épistémiques :

Nous noterons les interprétations seulement par les valuations, i.e. 10 dénotera l'interprétation qui envoie a à Vrai et b à Faux. Deux interprétations sont équivalentes, pour le pré-ordre, si elles apparaissent à un même niveau. Une interprétation I est "meilleure" qu'une interprétation J ($I <_\varphi^b J$) si elle apparaît à un niveau inférieur.

$\leq_{a \circ b} =$	00 10 01 11	$\leq_{b \circ a} =$	00 01 10 11	$\leq_{a \wedge b} =$	00 01 10 11
$\leq_{(a \wedge b) \circ a} =$	00 01 10 11	$\leq_{(a \wedge b) \circ a \circ \neg b} =$	01 11 00 10	$\leq_{(a \wedge b) \circ \neg b} =$	01 11 00 10

FIG. 3.1 – Exemples opérateur basique à mémoire

Sur les exemples donnés figure 3.4.1, on remarque bien que l'ordre d'arrivée des informations est très important pour les opérateurs à mémoire. Comme l'exprime la condition (C) les bases de connaissance associées aux trois historiques $a \circ b$, $b \circ a$ et $a \wedge b$ sont les mêmes. Mais on remarque également que les trois pré-ordres diffèrent.

Dans le même ordre d'idées, il est facile d'illustrer sur ces exemples que les opérateurs de révision à mémoire ne satisfont pas les postulats (I4), (I5) et (I6) de Lehmann. Prenons par exemple $(a \wedge b) \circ a \circ \neg b \not\rightarrow a \wedge b \circ \neg b$ comme contre-exemple à (I4) et (I5). Et prenons $a \circ b \circ \neg(a \wedge b) \not\rightarrow a \circ (a \wedge b) \circ \neg(a \wedge b)$ comme contre-exemple pour (I6). L'opérateur basique satisfait (I7) puisque, comme nous le montrons ci-dessous, il satisfait (C2) (voir section 3.4.2 pour un contre-exemple à (I7)).

Il est facile de montrer qu'après un nombre suffisant d'itérations (par des formules aléatoires), l'opérateur basique et tous les opérateurs que nous définissons dans ce chapitre mènent

à des pré-ordres linéaires sur les interprétations, i.e. des pré-ordres où il n'y a pas d'interprétations équivalentes. Avec de tels pré-ordres, les bases de connaissance associées aux états épistémiques sont des formules complètes. Comme souligné section 2.6, cela peut être considéré comme de l'apprentissage : lorsqu'un agent a un long historique de révisions, il garde des préférences sur les mondes possibles issues de son expérience passée.

Il faut noter qu'avec cette méthode de construction d'opérateurs de révision, on obtient des opérateurs ayant des propriétés intéressantes même si l'assignement fidèle de départ est très simple. En effet, l'assignement fidèle utilisé par exemple pour construire l'opérateur basique est l'assignement le plus simple pouvant être obtenu. Il se contente de faire la distinction entre modèles et contre-modèles. Notons que si l'on utilise cet assignement pour construire un opérateur de révision classique au moyen du théorème de représentation de Katsuno et Mendelzon (section 1.2.4), on obtient l'opérateur de révision par intersection totale qui n'est intuitivement pas très satisfaisant (section 1.2.1). Alors qu'avec la construction proposée ici on obtient un opérateur avec un comportement convenable.

Le pré-ordre associé à chaque état épistémique par le pré-ordre basique est un assignement conservatif fort. Il n'est pas dur de voir alors que le seul opérateur de révision à mémoire qui satisfait les conditions de l'assignement conservatif fort est l'opérateur basique :

Théorème 50 *Le seul opérateur de révision à mémoire à valeur dans des pré-ordres totaux (i.e. $\mathcal{IE}(\Phi)$) est un pré-ordre total, qui satisfait (H0)-(H7) et (C2) est l'opérateur de révision basique à mémoire.*

Preuve : Nous montrons que les conditions de l'assignement conservatif fort correspondent à un assignement défini à partir de la définition 46 et d'un assignement fidèle basique sur les bases de connaissance. Nous utilisons alors le théorème 47 pour conclure.

On montre d'abord qu'un assignement fidèle basique sur les bases de connaissance donne, grâce à la définition 46, un assignement fidèle conservatif fort. Les conditions 1 à 5 sont vérifiées d'après le théorème 48. Il reste à vérifier la condition 6. Si $I \not\models \varphi$ et $J \not\models \varphi$, alors comme l'assignement fidèle donne un pré-ordre basique $\leq_\varphi$ on a $I \simeq_\varphi J$ alors d'après la condition 5 $I \leq_{\Phi \circ \varphi} J$ ssi $I \leq_\Phi J$.

Pour l'autre implication considérons un assignement conservatif fort et montrons que le pré-ordre associé à chaque φ est un pré-ordre basique. C'est-à-dire que l'on voudrait montrer que si $I \not\models \varphi$ et $J \not\models \varphi$, alors $I \simeq_\varphi J$. Notons que pour deux interprétations I et J on a exclusivement un des trois faits suivants : soit $I <_\varphi J$, soit $I \simeq_\varphi J$, soit $I >_\varphi J$. Supposons que l'on se trouve dans le premier cas, i.e. $I <_\varphi J$, d'après la condition 4 on a $I <_{\Phi \circ \varphi} J$, alors que la condition 6 demande $I \leq_{\Phi \circ \varphi} J$ ssi $I \leq_\Phi J$ et $J \leq_{\Phi \circ \varphi} I$ ssi $J \leq_\Phi I$. C'est-à-dire que les conditions 4 et 6 impliquent $I <_\Phi J$. Or l'ordre $\leq_\Phi$ entre I et J est a priori quelconque, d'où contradiction. Le dernier cas est impossible pour la même raison. Le seul cas possible est donc $I \simeq_\varphi J$. C'est ce que nous voulions montrer.

□

On peut donc remarquer en étudiant la preuve ci dessus que la condition 6 de l'assignement conservatif fort est équivalente, en présence des conditions 1-5, à la condition 6' suivante :

6' Si $I \not\models \varphi$ et $J \not\models \varphi$, alors $I \simeq_\varphi J$

Dans [Spo87] Spohn critique cette méthode de révision en soulignant trois inconvénients de cette approche.

Tout d'abord cet opérateur n'est pas réversible, c'est-à-dire qu'étant donnés l'état épistémique actuel et la dernière information incorporée, il n'est pas possible de retrouver l'ancien état épistémique. Cela est vrai dans le formalisme de Spohn et dans celui présenté ici. Mais dans [BDP99, BKPP99] les états épistémiques sont codés sous forme de polynômes et l'opérateur basique est alors une "multiplication" sur les polynômes. Etant donnée cette représentation, l'opérateur basique est réversible.

Deuxièmement cet opérateur n'est pas commutatif, c'est-à-dire que si φ_1 et φ_2 sont deux propositions logiquement indépendantes¹, l'état épistémique résultant de l'incorporation de ces deux propositions n'est pas le même si φ_1 arrive avant φ_2 ou si φ_2 arrive avant φ_1 ou si les deux propositions arrivent ensembles. Selon Spohn (Nayak et al. donnent une propriété (conjonction) similaire), si deux propositions ne sont pas contradictoires entre elles, alors l'état épistémique résultat ne doit pas dépendre de l'ordre d'incorporation de ces informations. Il est clair que les opérateurs de révision à mémoire ne satisfont pas cette condition puisqu'ils donnent une grande priorité à la dernière information. Néanmoins une forme faible de cette condition est satisfaite puisque ces opérateurs satisfont la propriété (C) qui exprime cette idée sur la base de connaissance de l'état épistémique résultat.

Dernièrement Spohn souligne que l'hypothèse qu'être informé au sujet de φ rende tous les mondes possibles satisfaisant φ plus plausibles que ceux qui satisfont $\neg\varphi$ est une hypothèse forte. Il est vrai que l'hypothèse est forte, c'est pourquoi nous avons appelé cela le principe de *primauté forte de la nouvelle information*. Mais l'étude des propriétés de ces opérateurs manquait dans la littérature et considérer cet opérateur basique comme un cas particulier d'une classe d'opérateurs de révision aide à justifier cette approche. De plus, lorsque l'on ne dispose d'aucune information quantitative additionnelle attachée à la nouvelle information, ce qui est nécessaire pour le cadre de Spohn, les deux seules possibilités qu'évoque Spohn sont la révision naturelle de Boutilier et l'opérateur de révision à mémoire basique. Et seul l'opérateur de révision à mémoire basique est un cas particulier des opérateurs de conditionnalisation proposés par Spohn. Cela est visible sur la figure 3.2.

Finalement on peut également souligner le fait que Liberatore [Lib97] a montré que plusieurs problèmes sont plus simples (en complexité) pour l'opérateur basique à mémoire que pour les autres méthodes de révision itérée (y compris la révision naturelle de Boutilier, la révision rangée de Lehmann et les transmutations de Williams).

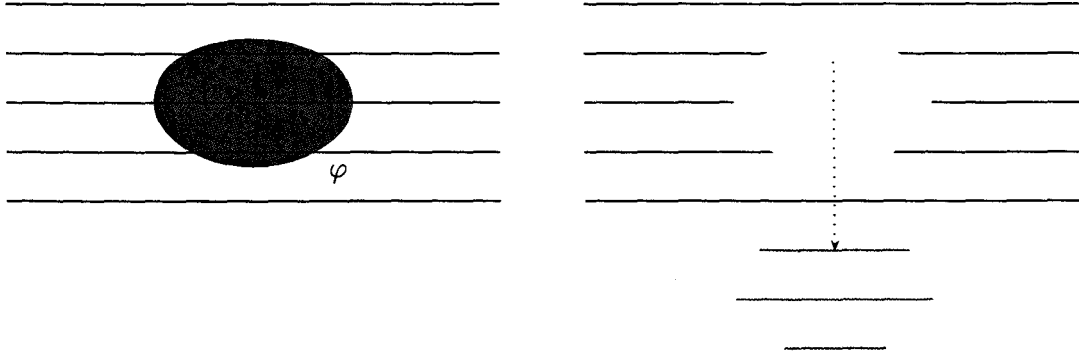
3.4.2 Opérateur Dalal à mémoire

Nous utiliserons dans cette section la distance de Dalal [Dal88a], qui est, en fait, la distance de Hamming entre interprétations :

Définition 49 Soient deux interprétations I et J , la distance de Dalal $\text{dist}(I, J)$ est définie comme le nombre de variables propositionnelles où ces deux interprétations diffèrent. Soit une base de connaissance φ , la distance de Dalal entre une interprétation I et cette base est :

$$d(I, \varphi) = \min_{J \models \varphi} (\text{dist}(I, J))$$

1. i.e. t.q. $\varphi_1 \not\models \varphi_2$, $\varphi_2 \not\models \varphi_1$, $\varphi_1 \not\models \neg\varphi_2$, $\varphi_2 \not\models \neg\varphi_1$

FIG. 3.2 – Opérateur basique à mémoire : révision par φ

On définit alors l'assignement qui associe à chaque base de connaissance un pré-ordre comme suit :

Définition 50 Soient une base de connaissance φ et deux interprétations I et J .

$$I \leq_{\varphi}^d J \text{ si et seulement si } d(I, \varphi) \leq d(J, \varphi)$$

Nous obtenons donc un pré-ordre avec les modèles de φ au niveau le plus bas et les autres interprétations distribuées sur les autres niveaux.

On utilise la définition 46 pour construire l'assignement fidèle $\leq_{\Phi}^d$ sur les états épistémiques à partir de l'assignement $\leq_{\varphi}^d$.

Définition 51 Soient un état épistémique Φ et une formule φ , l'opérateur de révision Dalal à mémoire est défini à partir de l'assignement fidèle normalisé $\leq_{\varphi}^d$ et de la définition 46. En particulier $Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \leq_{\Phi}^d)$.

D'après le théorème 50 nous savons que l'opérateur de Dalal à mémoire ne satisfait pas (C2). Cela s'illustre simplement sur l'exemple 4. Soient $\Phi = \top$, $\mu = \text{add_ok} \wedge \text{multiplier_ok}$ et $\alpha = \neg \text{add_ok}$. Les pré-ordres correspondants sont représentés figure 3.4.2.

$$\begin{array}{lll}
 \leq_{\Phi} = & \begin{array}{c} 00 \\ 01 \ 10 \\ 11 \end{array} & \leq_{\mu} = \begin{array}{c} 00 \\ 01 \ 10 \\ 11 \end{array} & \leq_{\alpha} = \begin{array}{c} 10 \ 11 \\ 00 \ 01 \end{array} \\
 & & & \\
 \leq_{\Phi \circ \alpha} = & \begin{array}{c} 10 \ 11 \\ 00 \ 01 \end{array} & \leq_{\Phi \circ \mu \circ \alpha} = & \begin{array}{c} 10 \\ 00 \\ 11 \\ 01 \end{array}
 \end{array}$$

FIG. 3.3 – Exemples opérateur Dalal à mémoire

On a donc $\alpha \models \neg \mu$ mais $\Phi \circ \mu \circ \alpha \leftrightarrow \neg \text{add_ok} \wedge \text{multiplier_ok}$ alors que $\Phi \circ \alpha \leftrightarrow \neg \text{add_ok}$.

Nous avons vu section précédente que les opérateurs à mémoire ne satisfaisaient pas beaucoup de postulats de Lehmann. Alors que l'opérateur basique satisfait (I7), l'opérateur de Dalal à mémoire ne satisfait pas ce postulat. Par exemple, $\top \wedge \neg(a \wedge b) \not\models \top \circ a \wedge b \circ \neg(a \wedge b)$.

3.4.3 Les opérateurs OTP

Ryan a proposé d'appliquer ses *Présentations Ordonnées de Théories* (Ordered Presentations of Theories ou OTP) à la révision de la connaissance [Rya94]. Très grossièrement, un OTP est un multi-ensemble de formules muni d'un pré-ordre partiel. Ce pré-ordre représente la confiance relative des sources d'information de chaque formule. Nous ne donnerons pas ici la définition exacte d'un OTP, le lecteur intéressé pourra se reporter à e.g [Rya91, Rya92] pour plus de détails. Nous n'introduirons que les notions nécessaires à la définition de l'opérateur de révision de Ryan.

Nous verrons que l'opérateur de Ryan ne satisfait pas les propriétés désirées pour un opérateur de révision. Nous donnerons alors deux modifications de cet opérateur qui satisfont les propriétés demandées.

Ces opérateurs sont intéressants car, il n'y a pas de pré-ordres *a priori*. Cette information est fournie par la formule elle-même d'une manière (syntaxique) très naturelle.

Nous devons d'abord définir les monotonies d'une formule.

Définition 52 Soient une interprétation I et une variable propositionnelle p . Alors $I^{[p]}$ (respectivement $I^{[\neg p]}$) dénote l'interprétation qui est identique à I sur chaque variable propositionnelle exceptée (peut-être) sur p qui est mis à Vrai (respectivement Faux).

Définition 53 Soient une formule φ différente de $\perp$ et p une variable propositionnelle.

1. φ est monotone en p si pour toute interprétation I , $I \models \varphi$ implique $I^{[p]} \models \varphi$.
2. φ est anti-monotone en p si pour toute interprétation I , $I \models \varphi$ implique $I^{[\neg p]} \models \varphi$.

L'ensemble des variables propositionnelles sur lesquelles φ est monotone (respectivement anti-monotone) est noté φ^+ (resp. φ^-). Ces deux ensembles sont appelés les monotonies de φ . Si $\varphi = \perp$, alors $\varphi^+ = \varphi^- = \emptyset$.

De cette définition, Ryan définit une relation d'inférence qu'il appelle *inférence naturelle*.

Définition 54 φ infère naturellement μ , noté $\varphi \vdash_N \mu$, si

1. $\varphi \vdash \mu$, et
2. $\varphi^+ \subseteq \mu^+$, et
3. $\varphi^- \subseteq \mu^-$.

Cette relation a quelques propriétés intéressantes. En particulier, et contrairement à l'inférence classique, elle ne permet pas d'ajouter des disjonctifs non pertinents dans les conclusions (par exemple $p \not\vdash_{NP} p \vee q$). Voir [Rya91, Rya92] pour plus de détails sur cette relation d'inférence.

Finalement la relation de préférence associée à une formule φ est fournie par l'ensemble des conséquences naturelles que les interprétations satisfont :

Définition 55 Soient une formule φ et deux interprétations I, J . La relation $\preceq_\varphi$ est définie comme : $I \preceq_\varphi J$, si pour chaque μ , $\varphi \vdash_N \mu$ implique $(J \models \mu \Rightarrow I \models \mu)$.

Donc une interprétation est meilleure qu'une autre si elle satisfait plus de conséquences naturelles. On peut noter que la relation $\preceq_\varphi$ est un pré-ordre partiel.

Enfin, l'opérateur de Ryan est défini en utilisant cette relation de préférence à la place de l'assignement fidèle sur les bases de connaissance dans la définition 46 de façon à définir la relation $\preceq_\Phi$ et :

Définition 56 Soient un état épistémique Φ et une formule φ , l'opérateur de révision de Ryan est défini à partir de l'assignement $\preceq_\Phi$ et de la définition 46. En particulier $Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \preceq_\Phi)$.

Puisque cet assignement est défini à partir de pré-ordres partiels, il ne satisfait donc pas tous les postulats :

Théorème 51 L'opérateur de révision de Ryan satisfait les postulats (H0), (H1), (H3), (H4), (H5) et (H7), mais il ne satisfait pas (H2) et (H6).

Un contre-exemple à (H2) et (H6), donné dans [Rya94], est le suivant :

Soient $\varphi_1 = p \vee q \vee r$, $\varphi_2 = \neg p \wedge \neg q \wedge \neg r$ et $\varphi_3 = (p \leftrightarrow q) \wedge \neg r$. Alors pour (H2), prenons $\Phi = \varphi_1 \circ \varphi_2$ et $\varphi = \varphi_3$. Alors $Mod(\Phi) = \{011, 101, 110\}$ et $Mod(\varphi) = \{000, 001, 010, 100, 110, 111\}$, donc $Mod(\Phi \wedge \varphi) = \{110\}$ alors que $Mod(\Phi \circ \varphi) = \{110, 001\}$. Le même contre-exemple peut être utilisé pour (H6) en prenant $\Phi = \varphi_1$, $\varphi = \varphi_2$ et $\mu = \varphi_3$.

Ces deux violations des postulats semblent très gênantes. En particulier (H2) semble difficilement critiquable. Nous montrons donc ci-après comment modifier cet opérateur pour satisfaire les propriétés.

La manière la plus facile pour modifier l'opérateur de Ryan afin de satisfaire les propriétés des opérateurs de révision à mémoire est de "compléter" les pré-ordres partiels $\preceq_\varphi$ en pré-ordres totaux. Cela peut être obtenu de deux façons.

Clôture du pré-ordre

En s'inspirant de la construction de la clôture rationnelle d'une base de conditionnels [LM92, Pea90], on peut imaginer une déformation fainéante du pré-ordre. C'est-à-dire une déformation qui transforme un pré-ordre partiel en un pré-ordre total avec un effort minimum.

Définition 57 Soit $\rho_\varphi(I)$ la "distance de I à φ " dans le sens suivant :

1. Si $I \in \min(\mathcal{W}, \preceq_\varphi)$ alors $\rho_\varphi(I) = 0$,
2. sinon $\rho_\varphi(I) = a$, où a est la longueur de la plus longue chaîne d'inégalités strictes $I_0 \prec_\varphi \dots \prec_\varphi I_n$ avec $I_0 \in \min(\mathcal{W}, \preceq_\varphi)$ et $I_n = I$.

Cette "distance" définit un pré-ordre total sur les interprétations :

Définition 58 $I \leq_\varphi^{\text{OTP}_1} J$ si et seulement si $\rho_\varphi(I) \leq \rho_\varphi(J)$.

Il est facile de voir que l'assignement ainsi défini est un assignement fidèle normalisé. Enfin, on définit l'opérateur de révision grâce à la définition 46.

Définition 59 Soient un état épistémique Φ et une formule φ , l'opérateur de révision OTP_1 est défini à partir de l'assignement fidèle normalisé $\leq_\varphi^{\text{OTP}_1}$ et de la définition 46. En particulier $Mod(\Phi \circ \varphi) = \min(Mod(\varphi), \leq_\varphi^{\text{OTP}_1})$.

Nous allons illustrer ce principe d' "effort minimal" sur un exemple: soit une base de connaissance $\varphi = (\neg a \vee \neg b) \wedge \neg c$.

La figure 3.4(a) représente le pré-ordre $\preceq_\varphi$. Puisque c'est un pré-ordre partiel, les flèches $I \leftarrow J$ dénotent les relations $I \prec_\varphi J$ (pour faciliter la lecture, la transitivité, la réflexivité et l'équivalence entre interprétations minimales ne sont pas représentées).

La figure 3.4(b) représente le pré-ordre $\leq_\varphi^{\text{OTP}_1}$ correspondant. Il est clair que si $I \prec_\varphi J$ alors $I <_\varphi^{\text{OTP}_1} J$. Donc la seule interprétation qui n'est pas directement placée est 110. L' "effort minimal" est illustré par la figure 3.4. Le premier niveau où peut être placé 110 est le second niveau, c'est donc le niveau choisi.

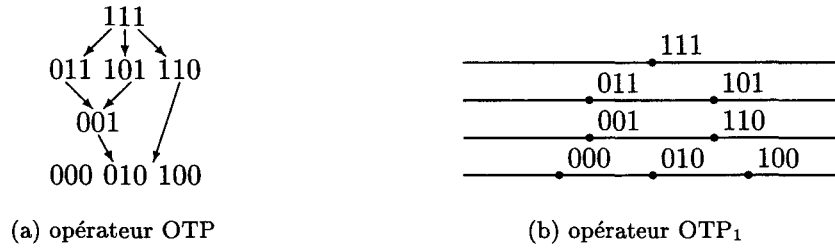


FIG. 3.4 – clôture du pré-ordre

Théorème 52 *L'opérateur de révision OTP_1 est un opérateur de révision à mémoire.*

La preuve de ce théorème est simple, il suffit de remarquer que la fonction qui, à chaque base de connaissance, associe le pré-ordre $\leq_\varphi^{\text{OTP}_1}$, est un assignement fidèle et donc, d'après la construction, l'opérateur obtenu est un opérateur de révision à mémoire.

Cardinalités

Une autre façon de définir un pré-ordre total à partir des pré-ordres partiels fournis par l'opérateur de Ryan est d'interpréter cette idée différemment. La motivation du pré-ordre $\preceq_\varphi$, qui définit l'opérateur de Ryan, est qu'une interprétation I est meilleure qu'une interprétation J pour une base de connaissance φ si I satisfait toutes les conséquences naturelles de φ que J satisfait. Autrement dit, I est meilleure que J si I satisfait plus de conséquences naturelles que J . En suivant cette idée on peut se concentrer uniquement sur le nombre de conséquences naturelles satisfaites et donc :

Définition 60 $I \leq_\varphi^{\text{OTP}_2} J$ si et seulement si

$$\text{card}(\{\mu \mid \varphi \vdash_{N\mu} \text{ et } J \models \varphi\}) \leq \text{card}(\{\mu \mid \varphi \vdash_{N\mu} \text{ et } I \models \varphi\})$$

On peut alors définir un opérateur de révision en utilisant ce pré-ordre et la définition 46.

Définition 61 Soient un état épistémique Φ et une formule φ , l'opérateur de révision OTP_2 est défini à partir de l'assignement fidèle normalisé $\leq_\varphi^{\text{OTP}_2}$ et de la définition 46. En particulier $\text{Mod}(\Phi \circ \varphi) = \min(\text{Mod}(\varphi), \leq_\varphi^{\text{OTP}_2})$.

Cette définition est également une "complétion" du pré-ordre $\preceq_\varphi$ puisque si $I \preceq_\varphi J$, alors $I \leq_\varphi^{\text{OTP}_2} J$.

Théorème 53 *L'opérateur de révision OTP_2 est un opérateur de révision à mémoire.*

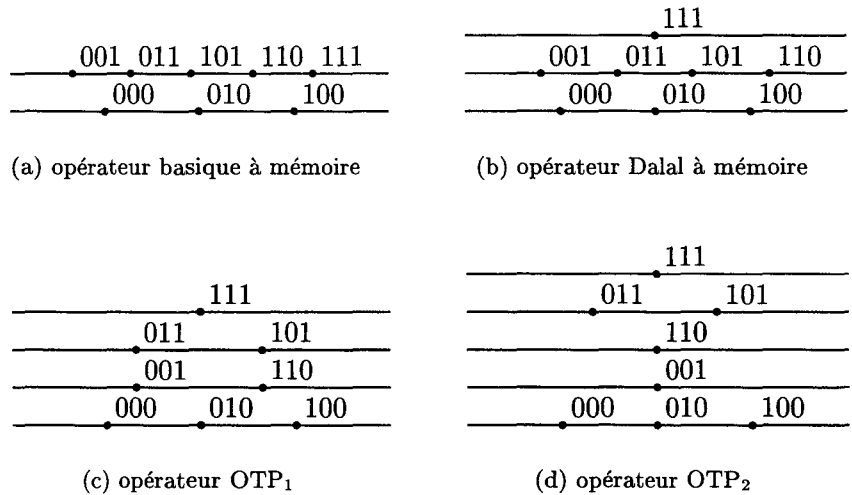


FIG. 3.5 – Distinction entre les opérateurs de révision à mémoire

3.4.4 Exemple

Nous allons montrer dans cette section que les opérateurs définis précédemment sont bien des opérateurs différents. Ce qui n'est pas évident *a priori* pour les opérateurs OTP_1 et OTP_2 .

Pour montrer cela, il suffit de montrer que les assignements fidèles correspondants à chacun de ces opérateurs sont différents. C'est ce que nous faisons ici en considérant la formule $\varphi = (\neg a \vee \neg b) \wedge c$. Nous donnons figure 3.5 les pré-ordres associés par les différents assignements.

Les quatre pré-ordres sont différents, les quatre opérateurs de révision à mémoire définis dans les sections précédentes sont donc différents.

3.5 Révision dynamique

Dans cette section nous allons définir et utiliser des Etats Epistémiques plus complexes que ceux utilisés dans les sections précédentes.

Les opérateurs de révision à mémoire définis dans ce chapitre donnent une priorité forte à la nouvelle information, permettant ainsi d'apporter un élément de réponse au problème de la révision itérée. C'est-à-dire qu'ils capturent un certain nombre de propriétés logiques en ce qui concerne le maintien de l'information conditionnelle.

La critique que l'on peut leur adresser est que, et bien que cela ait été intuitivement justifié, la stratégie de révision reste toujours la même.

C'est-à-dire que les connaissances de l'agent sont dynamiques mais sa façon de les remettre en question est statique. Il y a donc là, en quelque sorte, une information "ancienne", celle fournie par les pré-ordres *a priori* associés à chaque formule, qui ne peut être remise en question.

Il peut être intéressant de se demander s'il est possible d'adapter cette définition pour définir des opérateurs, en quelque sorte, plus dynamiques.

L'idée est la suivante: on part toujours de cet assignement fidèle donné *a priori*, qui code les croyances initiales de l'agent. C'est-à-dire qu'à la formule φ est associé un pré-ordre $\leq_\varphi$. Avec un opérateur à mémoire, à chaque fois que la nouvelle information φ devra être incorporée, c'est ce même pré-ordre qui lui sera associé. L'idée des opérateurs dynamiques est de changer ce pré-ordre en cours de révision.

De manière fonctionnelle, on peut décrire les opérateurs de révision à mémoire comme une fonction qui, à un état épistémique (un pré-ordre) et à une nouvelle information (une formule), associe un nouvel état épistémique :

$$(\leq_\Phi, \varphi) \xrightarrow{\circ} \leq_{\Phi'}$$

Mais en fait nous avons vu que chaque nouvelle information φ est associée à un pré-ordre $\leq_\varphi$ par un assignement f donné *a priori* et l'opérateur se décrit plutôt comme :

$$\begin{array}{c} \varphi \\ \downarrow f(\varphi) \\ (\leq_\Phi, \leq_\varphi) \end{array} \xrightarrow{\circ} \leq_{\Phi'}$$

La fonction f restant fixe. L'idée est donc de faire varier la fonction f en fonction des révisions, c'est-à-dire que la révision dynamique donnerait :

$$\begin{array}{c} \varphi \\ \downarrow f(\varphi) \\ (\leq_\Phi, \leq_\varphi) \end{array} \xrightarrow{\circ} (\leq_{\Phi'}, f')$$

où la nouvelle fonction f' est identique à f à l'exception faite de $f'(\varphi)$ qui est à présent égal au nouvel état épistémique obtenu. Plus formellement f' est définie comme:

$$f'(\varphi') = \begin{cases} f(\varphi') & \text{si } \text{Mod}(\varphi') \neq \min(\mathcal{W}, \leq_{\Phi'}) \\ \leq_{\Phi'} & \text{sinon} \end{cases}$$

Intuitivement, l'explication d'un tel comportement pour un opérateur de révision vient de l'idée que lorsque l'on rencontre une information à incorporer que l'on a déjà rencontrée, la répétition de cette évidence tend à nous faire penser que nos anciennes croyances étaient justes et à nous faire reconsidérer l'ancien état épistémique correspondant.

Nous allons à présent donner une définition plus formelle de ces opérateurs et étudier leurs propriétés.

Dans cette section on supposera qu'un état épistémique est constitué d'une base de connaissance, dénotant la connaissance actuelle, et d'un assignement fidèle sur les bases de connaissance, dénotant la stratégie de révision actuelle. Un état épistémique est donc un couple $\Phi = (\varphi, f)$ où φ est une base de connaissance et f est un assignement fidèle total, i.e. une fonction qui associe à chaque base de connaissance ψ un pré-ordre total $f(\psi)$ tel que les

modèles de la base de connaissance sont minimaux pour le pré-ordre $f(\psi)$.

Pour une base φ quelconque $f(\varphi)$ est donc un pré-ordre total. Pour faciliter la lecture on notera $I \leq_{f(\varphi)} J$ au lieu de $If(\varphi)J$. De même $I \simeq_{f(\varphi)} J$ est une notation pour $If(\varphi)J$ et $Jf(\varphi)I$; et $I <_{f(\varphi)} J$ dénote $If(\varphi)J$ et $\text{non}(Jf(\varphi)I)$.

On notera Ξ l'état épistémique initial, c'est-à-dire l'état épistémique correspondant à un historique vide. Cet état épistémique a donc comme base de connaissance associée $Bel(\Xi) = \top$ et l'assignement fidèle correspondant est quelconque et il dénote la stratégie initiale de révision. Ce peut donc être l'assignement fidèle basique (définition 47) ou l'assignement fidèle donné par la distance de Dalal (définition 50) par exemple.

L'opérateur de révision dynamique à mémoire est alors défini par :

Définition 62 Soit un état épistémique $\Phi = (\varphi, f)$ et soit une formule μ , on définit le nouvel état épistémique $\Phi \circ \mu$, issu de la révision dynamique à mémoire de Φ par la nouvelle information μ , par $\Phi \circ \mu = (\varphi', f')$ où φ' est une formule ayant comme modèles $\min(\text{Mod}(\mu), \leq_{f(\varphi)})$ et f' est une fonction identique à f pour toute base de connaissance φ différente de φ' et $\leq_{f'(\varphi')}$ est défini par :

$$I \leq_{f'(\varphi')} J \text{ ssi } I <_{f(\mu)} J, \text{ ou } I \simeq_{f(\mu)} J \text{ et } I \leq_{f(\varphi)} J$$

Il faut remarquer dans cette définition que l'on dispose en fait d'un assignement fidèle sur les états épistémiques induit par l'assignement fidèle sur les interprétations à chaque étape du processus de révision. Mais comme l'assignement fidèle sur les interprétations change à chaque étape, l'assignement fidèle sur les états épistémiques est également modifié à chaque itération. Cette modification est très progressive car la plupart des correspondances sont maintenues, seul le pré-ordre associé à la base de connaissance actuelle est modifié.

En fait, si l'on considère la définition d'un opérateur de révision au sens de Darwiche et Pearl (cf section 2.1.1), cela revient à considérer que l'on change d'opérateur de révision à chaque itération (puisque l'assignement fidèle correspondant change à chaque itération).

Théorème 54 Un opérateur de révision dynamique à mémoire satisfait les postulats (H1)-(H6).

Preuve : Ce théorème est une conséquence directe de la remarque précédente. Puisqu'à chaque itération on dispose d'un assignement fidèle sur les états épistémiques, l'opérateur satisfait donc les postulats (R*1)-(R*6).

Par conséquent il ne reste plus qu'à prouver (H4) qui est direct par définition des opérateurs de révision dynamique à mémoire. □

(H7) n'est pas vérifié par les opérateurs de révision dynamique. On peut illustrer cela sur l'exemple suivant :

Exemple 9 On considère un langage $\mathcal{L}$ constitué de 2 variables propositionnelles a et b . Soit un assignement fidèle initial qui assigne les pré-ordres suivants à a , b et $a \wedge b$:

Si l'on considère la séquence de révision suivante : $\Phi = a \circ b$ et $\Phi' = a \wedge b \circ \neg a$ on a alors $\Phi \circ \Phi' \leftrightarrow \neg a \wedge b$ alors que $\Phi \circ Bel(\Phi') \leftrightarrow \Phi \circ (\neg a \wedge \neg b) \leftrightarrow \neg a \wedge \neg b$ La différence de résultat

$$\begin{array}{lll}
\leq_a = & \begin{array}{cc} 00 & 01 \\ 10 & 11 \end{array} & \leq_b = \begin{array}{cc} 00 & 10 \\ 01 & 11 \end{array} & \leq_{a \wedge b} = \begin{array}{c} 01 \\ 10 \\ 00 \\ 11 \end{array}
\end{array}$$

vient du fait que dans le deuxième cas, $\Phi \circ \text{Bel}(\Phi')$ on utilise bien le pré-ordre initial pour $a \wedge b$, alors que dans le premier cas, $\Phi \circ \Phi'$, le pré-ordre associé à la base $a \wedge b$ est modifié par la séquence de révisions. Le pré-ordre utilisé est alors :

$$\begin{array}{c} 00 \\ \leq_{a \wedge b} = \begin{array}{c} 10 \\ 01 \\ 11 \end{array} \end{array}$$

Ces opérateurs sont des opérateurs de révision au sens de Darwiche et Pearl, puisqu'ils satisfont les postulats (R*1)-(R*6). En ce qui concerne les postulats d'itération de Darwiche et Pearl on a un résultat similaire aux opérateurs à mémoire :

Théorème 55 *Un opérateur de révision dynamique à mémoire satisfait les postulats (C1), (C3) et (C4).*

Preuve : Soit un état épistémique $\Phi = (\rho, f)$, la révision de cet état épistémique par α donne un nouvel état épistémique $\Phi \circ \alpha = (\rho_1, f_1)$ où $\rho_1 = \min(\text{Mod}(\alpha), \leq_{f(\rho)})$. De même la révision par μ donne $\Phi \circ \mu = (\rho_2, f_2)$ où $\rho_2 = \min(\text{Mod}(\mu), \leq_{f(\rho)})$ et la révision du résultat par α donne $\Phi \circ \mu \circ \alpha = (\rho_3, f_3)$ où $\rho_3 = \min(\text{Mod}(\alpha), \leq_{f_2(\rho_2)})$.

Pour prouver (C1), i.e. pour montrer que si $\alpha \vdash \mu$, alors $(\Phi \circ \mu) \circ \alpha \leftrightarrow \Phi \circ \alpha$, on suppose que $\alpha \vdash \mu$ et on veut montrer que $\rho_3 \leftrightarrow \rho_1$, i.e. $\min(\text{Mod}(\alpha), \leq_{f(\rho)}) = \min(\text{Mod}(\alpha), \leq_{f_2(\rho_2)})$. Par définition on sait que $f_2(\rho_2)$ est défini par

$$\begin{array}{c} I \leq_{f_2(\rho_2)} J \text{ ssi } I <_{f(\mu)} J, \text{ ou } \\ I \simeq_{f(\mu)} J \text{ et } I \leq_{f(\rho)} J \end{array}$$

Or $\forall I, J \models \alpha$, par hypothèse $I \leq_{f(\mu)} J$ et $J \leq_{f(\mu)} I$. Donc $\forall I, J \models \alpha$, $I \leq_{f_2(\rho_2)} J$ ssi $I \leq_{f(\rho)} J$. Donc $\min(\text{Mod}(\alpha), \leq_{f(\rho)}) = \min(\text{Mod}(\alpha), \leq_{f_2(\rho_2)})$.

Pour prouver (C3), i.e. pour montrer que si $\Phi \circ \alpha \vdash \mu$, alors $(\Phi \circ \mu) \circ \alpha \vdash \mu$, on suppose que $\rho_1 \vdash \mu$ et on veut montrer que $\rho_3 \vdash \mu$, i.e. $\min(\text{Mod}(\alpha), \leq_{f_2(\rho_2)}) \subseteq \text{Mod}(\mu)$. Si α est inconsistent le résultat est trivial, on suppose donc α consistant. Par l'absurde, supposons que $\exists I_1 \in \min(\text{Mod}(\alpha), \leq_{f_2(\rho_2)})$ et $I_1 \not\models \mu$. Or par hypothèse $(\Phi \circ \alpha) \wedge \mu \not\models \perp$, donc soit $I_2 \models (\Phi \circ \alpha) \wedge \mu$. Par définition $f_2(\rho_2)$ est défini par

$$\begin{array}{c} I \leq_{f_2(\rho_2)} J \text{ ssi } I <_{f(\mu)} J, \text{ ou } \\ I \simeq_{f(\mu)} J \text{ et } I \leq_{f(\rho)} J \end{array}$$

Or $I_2 <_{f(\mu)} I_1$ donc par définition $I_2 <_{f_2(\rho_2)} I_1$ donc $I_1 \notin \min(\text{Mod}(\alpha), \leq_{f_2(\rho_2)})$. Contradiction.

La preuve de (C4) est similaire à celle de (C3).

□

Alors qu'il était possible dans le cadre des opérateurs de révision à mémoire de trouver un opérateur vérifiant (C2), aucun opérateur de révision dynamique à mémoire ne vérifie ce postulat. Il est facile de trouver des exemples où $\Phi \circ \mu \circ \neg\mu \not\leftrightarrow \Phi \circ \neg\mu$. Cela arrive dès que les pré-ordres associés à μ et à $\neg\mu$ ne sont pas des pré-ordres à deux niveaux, ce qui arrive forcément au bout d'un certain historique (i.e. pour un certain Φ).

3.6 Révision par états épistémiques

Si l'on revient au "schéma fonctionnel" des opérateurs de révision à mémoire donné section précédente :

$$f'(\varphi') = \begin{cases} f(\varphi') & \text{si } \text{Mod}(\varphi') \neq \min(\mathcal{W}, \leq_{\Phi'}) \\ \leq_{\Phi'} & \text{sinon} \end{cases}$$

On s'aperçoit que la nouvelle information est associée systématiquement à un pré-ordre et que l'on a ensuite une fonction qui, à deux pré-ordres, associe un nouveau pré-ordre. Deux de ces pré-ordres sont des états épistémiques et le troisième, $f(\varphi)$, peut également être interprété comme un état épistémique induit par la nouvelle information.

La question est alors de savoir s'il n'est pas possible de généraliser cette approche pour pouvoir définir la révision d'un état épistémique par un autre état épistémique.

De telles opérations permettraient de réviser les croyances d'un agent par autre chose que des observations, par exemple par des informations conditionnelles. Par exemple les croyances d'un autre agent : "je pense actuellement que φ est vrai et je trouve également μ plus crédible que α ". Ce qui est impossible avec les opérateurs de révision usuels. Une telle généralisation est proposée dans [BKPP].

Un point de notation tout d'abord, puisque l'on travaille à présent avec deux états épistémiques, la conjonction entre états épistémiques $\Phi \wedge \Theta$, dénote la conjonction de leurs bases de connaissance associées i.e. $\Phi \wedge \Theta = \text{Bel}(\Phi) \wedge \text{Bel}(\Theta)$. De même $\Phi \vdash \Theta$ signifie $\text{Bel}(\Phi) \vdash \text{Bel}(\Theta)$; et $\neg\Theta$ représente la base de connaissance $\neg\text{Bel}(\Theta)$. On définit les opérateurs de révision avec mémoire opérant sur des états épistémiques par :

Définition 63 *Un opérateur $\circ$ qui associe à deux états épistémiques Φ et Θ un état épistémique $\Phi \circ \Theta$ est un opérateur de révision avec mémoire si et seulement si il satisfait les propriétés suivantes :*

- (H1*) $\Phi \circ \Theta \vdash \Theta$
- (H2*) Si $\Phi \wedge \Theta$ est consistant, alors $\Phi \circ \Theta \leftrightarrow \Phi \wedge \Theta$
- (H3*) Si Θ est consistant, alors $\Phi \circ \Theta$ est consistant
- (H4*) Si $\Phi_1 = \Phi_2$ et $\Theta_1 \leftrightarrow \Theta_2$, alors $\Phi_1 \circ \Theta_1 \leftrightarrow \Phi_2 \circ \Theta_2$
- (HI*) $(\Phi \circ \Theta) \circ \Gamma = \Phi \circ (\Theta \circ \Gamma)$

Remarquons que les postulats (H1*)-(H4*) et (HI*) sont la généralisation naturelle des postulats (H1)-(H4) et (H7) lorsque la nouvelle information est un état épistémique. De même que pour le cadre usuel, le postulat (HI*) est une généralisation de postulats qui correspondent à (H5) et (H6). Plus précisément nous avons :

Théorème 56 *Le postulat $(H\Gamma^*)$ implique, en présence de $(H1^*)$, $(H2^*)$, et $(H4^*)$ les deux postulats suivants :*

$$(H5^*) (\Phi \circ \Theta) \wedge \Gamma \vdash \Phi \circ (\Theta \wedge \Gamma)$$

$$(H6^*) \text{ Si } (\Phi \circ \Theta) \wedge \Gamma \text{ est consistant, alors } \Phi \circ (\Theta \wedge \alpha) \vdash \Phi \circ (\Theta \wedge \Gamma)$$

La preuve est similaire à celle du théorème 41. Et comme dans le cadre usuel, $(H1^*-H4^*)$ et $(H\Gamma^*)$ impliquent la propriété suivante :

$$(C^*) \text{ Si } \Theta \wedge \Gamma \text{ est satisfaisable, alors pour tout état épistémique } \Theta' \text{ tel que } \Theta' \leftrightarrow \Theta \wedge \Gamma \text{ on a } (\Phi \circ \Theta) \circ \Gamma \leftrightarrow \Phi \circ \Theta'$$

Ce postulat dit que si l'on révisé successivement par deux états épistémiques dont les ensembles de croyances sont cohérents entre eux, cela revient à réviser (au niveau des croyances) par un état épistémique dont l'ensemble des croyances est la conjonction des ensembles de croyances des états épistémiques.

A la différence des théorèmes de représentation bien connus pour les opérateurs de révision classiques (AGM) où l'on construit un ordre, le théorème de représentation doit nous donner une caractérisation précise de la façon dont on mélange les pré-ordres, l'ancien et le nouveau, pour obtenir un nouvel état épistémique.

Théorème 57 *Soit $\circ$ un opérateur de révision avec mémoire. On peut définir pour chaque état épistémique Φ un pré-ordre $\leq'_\Phi$ de la façon suivante :*

$$I \leq'_\Phi J \Leftrightarrow I \models \Phi \circ \Theta$$

avec $Mod(\Theta) = \{I, J\}$.

Cette proposition est en fait une des parties du futur théorème de représentation. Elle nous dit comment associer des pré-ordres à des états épistémiques. Notons que c'est grâce à $(H4^*)$ que $\leq'_\Phi$ est bien défini.

Définition 64 *Soient $\leq_1$ et $\leq_2$ deux pré-ordres sur $\mathcal{W}$. On définit le pré-ordre $\leq = \text{lex}(\leq_1, \leq_2)$ en posant $I \leq J$ ssi $I <_1 J$ ou $I =_1 J$ et $I \leq_2 J$.*

On définit la base de connaissance $Bel(\leq'_\Phi)$ comme une formule ayant comme modèles exactement l'ensemble $\min(\mathcal{W}, \leq'_\Phi)$.

Théorème 58 *Un opérateur $\circ$ est un opérateur de révision (par état épistémique) avec mémoire si et seulement si on peut associer à chaque état épistémique Φ consistant, un pré-ordre $\leq'_\Phi$ tel que :*

$$(i) \text{ } Bel(\Phi) = Bel(\leq'_\Phi)$$

$$(ii) \text{ } Bel(\Phi \circ \Theta) = Bel(\text{lex}(\leq'_\Theta, \leq'_\Phi))$$

Remarque 2 *Le premier aspect saillant de ce résultat est le passage des états épistémiques abstraits à des états épistémiques concrets : des pré-ordres. Le deuxième aspect important est que l'on se donne une procédure de calcul du nouvel état épistémique : l'ordre lex . Finalement, au niveau des croyances, la méthode abstraite et la méthode concrète coïncident.*

Ce que l'on peut reprocher à cette représentation est que l'on n'a pas nécessairement $\Phi \circ \Theta = \text{lex}(\leq'_\Theta, \leq'_\Phi)$ lorsque les états épistémiques sont des pré-ordres. Néanmoins lorsque l'on part d'états épistémiques qui sont des pré-ordres et si l'opérateur satisfait un minimum de rationalité (postulat (H0*) plus bas) le problème ne se pose plus.

Définition 65 *Considérons une représentation concrète des états épistémiques, i.e. chaque état épistémique Φ est identifié à un pré-ordre $\leq_\Phi$. On dit qu'un opérateur $\circ$ entre états épistémiques est rationnel s'il satisfait le postulat suivant :*

$$(H0^*) \text{ Mod}(\Phi \circ \Theta) = \min(\text{Mod}(\Theta), \leq_\Phi)$$

Remarquons que si un opérateur est rationnel alors il satisfait aussi les postulats (H1*), (H2*) et (H3*). Ceci, joint à la proposition 56, nous donne le résultat suivant :

Théorème 59 *Un opérateur avec mémoire est rationnel ssi il satisfait les postulats (H0*), (H4*) et (H1*).*

Pour ce type d'opérateurs, nous avons la caractérisation suivante :

Théorème 60 *Un opérateur $\circ$ est un opérateur de révision avec mémoire rationnel seulement si*

$$\leq_{\Phi \circ \Theta} = \text{lex}(\leq_\Theta, \leq_\Phi)$$

En particulier, il y a un seul opérateur de révision avec mémoire rationnel.

Nous voulons explorer le comportement de ces opérateurs concernant l'itération. Pour ce faire nous allons commencer par adapter les postulats de Darwiche et Pearl (DP) au cadre où la nouvelle information est aussi un état épistémique. On redéfinit les postulats DP lorsque la nouvelle information est un état épistémique.

Postulats DP pour la révision itérée d'états épistémiques :

- (C1*) Si $\Theta \vdash \Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \leftrightarrow \Phi \circ \Theta$
- (C2*) Si $\Theta \vdash \neg\Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \leftrightarrow \Phi \circ \Theta$
- (C3*) Si $\Phi \circ \Theta \vdash \Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \vdash \Gamma$
- (C4*) Si $\Phi \circ \Theta \not\vdash \neg\Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \not\vdash \neg\Gamma$.

On a un résultat similaire à celui du cadre usuel.

Théorème 61 *Un opérateur de révision avec mémoire satisfait les postulats (C1*), (C3*) et (C4*). Mais il ne satisfait pas (C2*).*

Le théorème 60 nous dit qu'il y a un unique opérateur de révision avec mémoire rationnel. Néanmoins, en posant des restrictions sur la nature des états épistémiques codant les nouvelles informations, nous avons des opérateurs avec d'autres propriétés.

Par exemple, supposons que l'on ne considère comme nouvelle information que des pré-ordres à deux niveaux. Dans ce cas on peut identifier la nouvelle information à une formule. Et l'on est en train de considérer l'opérateur de révision basique défini section 3.4.1.

Notons de même que si les nouvelles informations sont les états épistémiques déterminés par leurs connaissances (les formules préférées) et la distance de Dalal [Dal88a], alors l'opérateur obtenu est l'opérateur de Dalal avec mémoire défini section 3.4.2.

3.7 Semi-révision à mémoire

Le principe de primauté de la nouvelle information est souvent critiqué (cf section 1.3.4) puisque celle-ci ne peut pas être acceptée dans toutes circonstances. Quelquefois nous avons plus confiance en la base de connaissance actuelle qu'en cette nouvelle information. Quelques opérateurs de semi-révision, n'obéissant pas à ce principe, ont été proposés récemment [Han97, Mak98, FH99, Sch98].

Si l'on accepte ce point de vue, on peut tout de même utiliser l'approche que l'on vient de donner en inversant les arguments de la révision, c'est-à-dire en révisant la nouvelle information par l'état épistémique. Cela définit alors une famille d'opérateurs de semi-révision. Plus formellement :

Définition 66 *Un opérateur de semi-révision à mémoire $\square$ est défini comme :*

$$\mu \square \varphi = \varphi \circ \mu$$

où $\circ$ est un opérateur de révision à mémoire.

Cela nous donne donc une famille d'opérateurs qui donnent une forte priorité à la connaissance actuelle. La nouvelle information n'est donc pas acceptée dans le nouvel état épistémique mais sa crédibilité est augmentée. Par exemple, si deux mondes possibles sont aussi plausibles pour un état épistémique et si la nouvelle information n'est satisfaite que par un des deux mondes, alors ce monde sera plus crédible que l'autre dans le nouvel état épistémique.

Hansson donne un ensemble de postulats pouvant s'appliquer aux opérateurs de semi-révision sur des bases de connaissance closes déductivement [Han97]. Nous donnons ci-après les postulats vérifiés par les opérateurs que nous venons de définir :

- | | |
|--|------------------------------|
| (K□1) $K \square A$ est une théorie | (clôture) |
| (K□2) Si $\neg A \notin K$, alors $A \in K \square A$ | (succès faible) |
| (K□3) $K \square A \subseteq K + A$ | (inclusion) |
| (K□4) Si $\neg A \notin K$, alors $K + A \subseteq K \square A$ | (vacuité) |
| (K□5) Si K est consistant, alors $K \square A$ est consistant | (maintien de la consistance) |
| (K□6) Si $A \leftrightarrow B$, alors $K \square A = K \square B$ | (extensionnalité) |

Hansson donne également des alternatives pour certains postulats de cette liste.

Les opérateurs de révision à mémoire vérifient l'ensemble de ces postulats lorsqu'on les exprime en termes d'états épistémiques plutôt qu'en termes de bases de connaissance.

3.8 Remarques

Nous avons proposé dans ce chapitre une syntaxe pour les états épistémiques. Nous avons montré que cette syntaxe pouvait être interprétée par les sémantiques proposées dans les différents travaux sur la révision itérée. Cela permet, en particulier, de souligner la connexion entre le problème de révision itérée et celui de l'historique des révisions. Maintenir un historique des révisions est donc un moyen de garantir de bonnes propriétés pour l'itération.

Le fait de garantir une importance plus grande aux dernières informations incorporées est psychologiquement satisfaisant et étend de manière naturelle l'hypothèse de primauté de la

nouvelle information faite dans le cadre AGM usuel.

Les opérateurs de révision à mémoire utilisent une stratégie de révision fixée au début du processus de révision induite par un assignement fidèle sur les bases de connaissance, et c'est l'utilisation de cette stratégie tout au long du processus qui assure de bonnes propriétés logiques.

Lorsque l'on compare cette approche avec les autres, cela permet d'illustrer une fois de plus le caractère problématique du postulat (C2) de Darwiche et Pearl. C'est le seul postulat de Darwiche et Pearl non systématiquement vérifié par les opérateurs de révision à mémoire.

Le seul opérateur de révision à mémoire satisfaisant (C2) est l'opérateur à mémoire basique, qui est également un cas particulier des opérateurs de conditionnalisation de Spohn. Bien que le comportement de cet opérateur est critiquable, puisqu'il satisfait (C2), la simplicité de sa définition en fait tout de même un opérateur digne d'intérêt.

Les opérateurs définis à partir des OTP de Ryan sont également intéressants puisqu'ils ne sont pas définis sémantiquement à partir d'une distance mais syntaxiquement et d'une façon assez naturelle à partir des bases de connaissance elles-mêmes. L'opérateur que Ryan avait proposé dans [Rya94, Rya92] n'était pas un opérateur de révision car il ne satisfaisait pas les propriétés de base. Les deux opérateurs que nous avons définis et qui étendent ce dernier corrigent ce défaut.

Le fait que l'assignement fidèle sur les bases de connaissance, qui induit un opérateur de révision à mémoire, soit fixe tout au long du processus d'itération peut faire penser à une connaissance initiale statique. Une nouvelle information induit toujours une même préférence *a priori* mais la stratégie de révision évolue avec l'historique des états épistémiques. Cela peut se motiver par un souci de rationalité et s'illustre sur les exemples que nous avons donné : par exemple (avec l'opérateur Dalal à mémoire) si un agent apprend que l'information composée de trois faits atomiques $a \wedge b \wedge c$ est vraie, il trouvera que $a \wedge b \wedge \neg c$ est plus crédible que $\neg a \wedge \neg b \wedge \neg c$ puisque, dans le premier cas, deux des trois faits atomiques sont vrais, alors que, dans le second, aucun ne l'est. Et cela *quelle que soit sa connaissance actuelle*. En effet cette préférence induite par la nouvelle information n'a aucun rapport avec l'état épistémique actuel de l'agent, les croyances (et connaissances) issues de celui-ci n'interviennent qu'ensuite.

Néanmoins, nous avons critiqué l'existence d'une telle information statique pour les opérateurs de révision rangée de la section 2.3, et l'on pourrait donc nous retourner la critique. Les opérateurs de révision dynamique à mémoire sont alors une extension des derniers opérateurs permettant d'ajouter un aspect plus dynamique au changement. En effet, l'assignement fidèle sur les bases de connaissance est alors modifié (très peu) à chaque itération.

Les opérateurs précédents associent à un couple composé d'un état épistémique et d'une formule, un nouvel état épistémique résultat de la révision. Mais un examen de leur définition montre que la formule induit une information de même nature que les états épistémiques initiaux et finaux. Cela mène donc tout naturellement à généraliser les opérateurs précédents pour permettre la révision d'un état épistémique par un état épistémique. Cela permet, en particulier, de réviser par des informations contenant des informations conditionnelles. Mais

d'un point de vue sémantique, cela revient à agréger deux pré-ordres en donnant une priorité importante à l'un deux (ce que l'on nomme dictateur en théorie du choix social [Arr63, Kel78]). Cela peut donc être vu comme un cas particulier de fusion (pondérée).

Une dernière variation possible à partir de ces opérateurs à mémoire est la définition d'opérateurs de semi-révision. Lorsque l'on n'a pas une confiance supérieure en la nouvelle information qu'en la connaissance actuelle, il est possible d'inverser les opérands des opérateurs de révision à mémoire de façon, en quelque sorte, à réviser la nouvelle information par l'état épistémique actuel. Cela permet de prendre en compte la nouvelle information mais sans lui donner la primauté sur la connaissance actuelle. Les opérateurs ainsi définis vérifient la plupart des propriétés proposées par Hansson [Han97] pour les opérateurs de semi-révision pour les bases de connaissance closes déductivement.

Chapitre 4

Révision et chaînage avant

Un des problèmes de la révision est que c'est en général un processus d'une complexité importante [EG92, Neb96, Lib97]. Cela est dû au fait que les opérateurs de révision classiques manipulent des théories closes déductivement. Le calcul de l'ensemble des conséquences de la nouvelle connaissance pose des problèmes de représentation et de calcul. De ce fait, les opérateurs de révision sont difficilement implémentables de manière efficace.

Une solution est de travailler avec des bases de connaissances qui ne sont pas closes déductivement (cf section 1.3.2) et de ne calculer les conséquences logiques de ces bases que lorsque c'est nécessaire. Une autre solution est de s'appuyer sur des méthodes de recherche incomplètes [BGMS98, BGMS99]. Une troisième solution est de ne pas travailler avec la logique classique mais avec une logique plus faible et plus facile à implémenter.

Dans ce chapitre nous examinons des opérateurs de changement basés sur le chaînage avant [BJKP97, BJKP98]. L'utilisation du chaînage avant permet de calculer le résultat d'une révision de manière efficace. Nous combinons donc deux des approches exposées ci-dessus, c'est-à-dire que nous proposons des opérateurs travaillant sur des bases de connaissance non closes déductivement, en utilisant une logique plus faible que la logique classique.

De plus ces opérateurs sont facilement utilisables dans des systèmes d'informations basés sur ce type d'inférence, comme la plupart des systèmes experts ou les systèmes de diagnostic par exemple. Cela permet donc d'introduire facilement un raisonnement non monotone dans ces systèmes.

Quelques approches similaires ont déjà été proposées dans la littérature. Ces approches sont généralement des extensions de la programmation logique. On peut par exemple citer REVISE [DNP94] qui utilise la programmation logique étendue (négation explicite) et Revision Programming [MT95, MT94] qui utilise un langage à la DATALOG [Ull88, AV91] admettant les suppressions (voir aussi [ALP⁺98]). Notre approche est plus simple puisque nous n'utilisons que le chaînage avant sur les formules propositionnelles, nous n'utilisons donc pas la négation par l'échec, ni de prédicat marquant explicitement l'ajout ou la suppression d'un littéral.

Nous définissons cinq opérateurs de changement. Le premier, appelé *mise à jour factuelle*,

actualise un ensemble de faits par un autre, codant une nouvelle information, tout en respectant un ensemble de règles codant les contraintes d'intégrité du système. Le second, *révision hiérarchique*, est basé sur une hiérarchie sur les règles dénotant à quel point celles-ci sont exceptionnelles. Lorsqu'une nouvelle information arrive, on utilise les règles les moins exceptionnelles consistantes avec cette nouvelle information. Le troisième, nommé *révision gangue*, étend la révision hiérarchique et calcule les ensembles maximaux de règles consistants avec la nouvelle information. Le quatrième opérateur, *révision gangue étendue*, combine les approches de la révision gangue et de la mise à jour factuelle. Le dernier opérateur, *révision gangue sélective*, est en fait une famille d'opérateurs définis à partir de fonctions de sélection.

Nous travaillerons ici avec un langage¹ défini comme suit :

Définition 67 *Un littéral (ou fait) est un atome ou la négation d'un atome. L'ensemble des littéraux sera noté Lit .*

Définition 68 *Une règle est une formule de la forme $l_1 \wedge l_2 \wedge \dots \wedge l_n \rightarrow l_{n+1}$ où l_i est un littéral pour tout $i = 1, \dots, n+1$. Une telle règle sera notée $l_1, l_2, \dots, l_n \rightarrow l_{n+1}$.*

On admet des règles de la forme $\rightarrow l$ qui codent des faits.

Définition 69 *Soient un ensemble fini de règles R et un ensemble fini de littéraux L (ces deux ensembles pouvant être vides). Un programme P est l'ensemble $R \cup L$. On appellera les éléments de R les règles de P et les éléments des L les faits de P .*

L'ensemble des programmes sera noté $Prog$.

Soit un programme $P = R \cup L$, on définit l'ensemble des *conséquences* de P par chaînage avant, noté $C_{fc}(P)$, comme le plus petit ensemble de littéraux L' tel que :

1. $L \subseteq L'$.
2. Si $l_1, l_2, \dots, l_n \rightarrow l$ est dans R et $l_i \in L'$ pour $i = 1, \dots, n$ alors $l \in L'$.
3. Si L' contient deux littéraux opposés, alors $L' = Lit$.

Un programme P est dit *consistant* si et seulement si $C_{fc}(P)$ ne contient pas deux littéraux opposés, c'est-à-dire un atome et sa négation.

Soient un programme P et un ensemble fini de littéraux L , L est dit *P -consistant* si et seulement si $P \cup L$ est consistant.

4.1 Définition des opérateurs

Nous allons présenter dans cette section quatre opérateurs de changement de la connaissance basés sur le chaînage avant. Le premier, *mise à jour factuelle*, utilise un programme codant les contraintes d'intégrités du système. La base de connaissance et la nouvelle information sont des ensembles de faits.

1. On supposera pour simplifier l'exposé que le langage est propositionnel, mais les opérateurs peuvent s'étendre facilement au premier ordre.

Les trois autres opérateurs ne supposent pas l'existence de contrainte d'intégrité pour le système et permettent de réviser un programme par un autre codant la nouvelle information. Ces trois opérateurs, *révision hiérarchique*, *révision gangue* et *révision gangue étendue*, utilisent une hiérarchie sur les règles induite par la base de connaissance. La révision gangue et la révision gangue étendue sont des extensions de la révision hiérarchique permettant de capturer plus de transitivité. Mais, comme nous le verrons section 4.2, cela se fait au dépens des propriétés logiques de ces opérateurs.

La plupart de ces opérateurs utilisent la notion de flocks de programmes. Nous détaillons donc cette notion dans une première section.

4.1.1 Flock

La révision d'une base de connaissance syntaxique donne parfois un ensemble de bases de connaissance "candidates" pour le résultat. Cela peut être comparé à la notion d'extension en logique des défauts [Rei80, Eth88].

On peut avoir plusieurs attitudes envers cet ensemble, soit on adopte une attitude "aventureuse" et on choisit comme résultat une des bases (ou certaines des bases), soit on adopte une attitude prudente (ou sceptique) et on conserve l'ensemble de ces bases. Cette attitude prudente a été choisie par Fagin *et al.* dans le cadre de la mise à jour de bases de données. Ils ont nommé *flock* cet ensemble de bases de connaissance (cf [FKUV86]).

Nous utiliserons de tels flocks pour les opérateurs définis dans ce chapitre. Nous devons alors définir quelques opérations sur ces flocks.

On définit la concaténation de flocks ' $\cdot$ ' de manière naturelle :

$$\langle P_1, \dots, P_n \rangle \cdot \langle P'_1, \dots, P'_m \rangle = \langle P_1, \dots, P_n, P'_1, \dots, P'_m \rangle$$

De manière à pouvoir travailler avec des flocks, on doit définir ce que sont les conséquences d'un flock $\mathcal{F} = \langle P_1, \dots, P_n \rangle$. On définit les conséquences par chaînage avant d'un flock, notées $C_{fc}(\mathcal{F})$, par :

$$C_{fc}(\mathcal{F}) = \bigcap_{i=1}^n C_{fc}(P_i)$$

On aura également besoin de la notion de conséquences par chaînage avant (sous les contraintes P) d'un flock, notées $C_{fc}^P(\mathcal{F})$, par :

$$C_{fc}^P(\mathcal{F}) = \bigcap_{i=1}^n C_{fc}(P_i \cup P)$$

Soient deux flocks $\mathcal{F}$ et $\mathcal{F}'$. On dit que $\mathcal{F}$ implique $\mathcal{F}'$, noté $\mathcal{F} \vdash \mathcal{F}'$ si $C_{fc}(\mathcal{F}) \supseteq C_{fc}(\mathcal{F}')$. $\mathcal{F}$ et $\mathcal{F}'$ sont *équivalents*, noté $\mathcal{F} \leftrightarrow \mathcal{F}'$, si $\mathcal{F} \vdash \mathcal{F}'$ et $\mathcal{F}' \vdash \mathcal{F}$. Similairement, si P est un programme on définit $\mathcal{F} \vdash P$ comme $C_{fc}(\mathcal{F}) \supseteq C_{fc}(P)$.

Soient deux flocks $\mathcal{F} = \langle P_1, \dots, P_n \rangle$ et $\mathcal{F}' = \langle P'_1, \dots, P'_m \rangle$. Nous noterons $\mathcal{F} \wedge \mathcal{F}'$ le flock $\langle P''_{1,1}, \dots, P''_{n,m} \rangle$ où $P''_{j,k} = P_j \cup P'_k$. De même si P est un programme on définit le flock $\mathcal{F} \wedge P$ comme $\langle P_1 \cup P, \dots, P_n \cup P \rangle$.

Notons que si $\mathcal{F} = \langle P_1, \dots, P_n \rangle$ et $\mathcal{F}' = \mathcal{F} \cdot \langle P'_1, \dots, P'_m \rangle$ avec $\forall i P'_i$ inconsistant, alors $\mathcal{F} \leftrightarrow \mathcal{F}'$. Donc lorsque l'on considère un flock, il suffit de considérer ses bases de connaissance consistantes.

4.1.2 Mise à jour factuelle

Soit un programme P fixé qui représente les contraintes d'intégrité du système et soit un ensemble de littéraux L qui code notre connaissance du monde, on désire définir le changement opéré par un ensemble de littéraux L' , représentant une nouvelle information, sur nos connaissances actuelles. On définit donc cet opérateur de mise à jour factuelle (factual update) comme suit :

$$L \diamond_P L' = \begin{cases} Lit & \text{si } L \text{ ou } L' \text{ n'est pas } P\text{-consistant} \\ \langle L_1 \cup L', \dots, L_n \cup L' \rangle & \text{sinon} \end{cases}$$

où $\{L_1, \dots, L_n\}$ est l'ensemble des sous-ensembles maximaux de L qui sont $P \cup L'$ -consistants.

On voit qu'après la révision d'un ensemble de littéraux L , on obtient généralement comme résultat un ensemble d'ensembles (un flock) de littéraux $\langle L_1, \dots, L_n \rangle$. On définit donc l'opérateur de mise à jour factuelle sur des flocks par :

Soit un flock $\langle L_1, \dots, L_n \rangle$ et une nouvelle information L' , l'opérateur de *mise à jour factuelle* est défini par :

$$\langle L_1, \dots, L_n \rangle \diamond_P L' = \begin{cases} Lit & \text{si } L' \text{ ou tous les } L_i \text{ ne sont pas } P\text{-consistants} \\ (L_{i_1} \diamond_P L') \cdot (L_{i_2} \diamond_P L') \cdots (L_{i_k} \diamond_P L') & \text{sinon} \end{cases}$$

où $\{L_{i_1}, \dots, L_{i_k}\}$ est l'ensemble de tous les éléments de $\{L_1, \dots, L_n\}$ qui sont consistants avec P .

Les conséquences du flock, représentant la connaissance actuelle, sont celles issues du programme actuel sous les contraintes P . C'est à dire que l'on utilise C_{fc}^P comme opérateur de conséquence.

Exemple 10 Soient le programme $P = \{a, b \rightarrow c ; a, d \rightarrow c\}$ et l'ensemble de littéraux $L = \{a, b, d\}$. Si la nouvelle information est $L' = \{\neg c\}$ il est facile de voir que :

$$L \diamond_P L' = \langle \{a\} \cup \{\neg c\}, \{b, d\} \cup \{\neg c\} \rangle$$

Dans [BJKP98] un algorithme pour calculer l'opérateur de mise à jour factuelle est proposé. Il est basé sur le calcul d'ensembles recouvrants minimaux (hitting sets).

4.1.3 Révision hiérarchique

Dans le cas de la mise à jour factuelle le programme est fixé et on restreint la base de connaissance et la nouvelle information à des ensembles de faits. Les opérateurs définis ci-après lèvent cette contrainte et on envisage alors la révision d'un programme par un autre.

Les opérateurs définis ci après sont inspirés de la dualité existant entre révision et relations d'inférence (cf section 1.2.5). Ces opérateurs peuvent être vus comme une adaptation de la clôture rationnelle [LM92] (ou de manière équivalente du System Z de Pearl [Pea90]) à cette logique basée sur le chaînage avant.

Définition 70 Soit un programme P . Un ensemble de littéraux L est dit *exceptionnel* pour P si et seulement si L n'est pas P -consistant et une règle $L \rightarrow l$ de P est dite *exceptionnelle* pour P si et seulement si L est exceptionnel pour P .

Notons que lorsque le corps d'une règle est vide, cette règle sera exceptionnelle si et seulement si P n'est pas consistant. Mais dans ce cas toutes les règles sont exceptionnelles. Une définition similaire de l'exceptionnalité est donnée dans [LM92].

Soit $(P_i)_{i \in \mathbb{N}}$ la suite décroissante définie par : P_0 est P et P_{i+1} est l'ensemble de toutes les règles exceptionnelles de P_i . Puisque P est fini, il existe un plus petit entier n_0 tel que pour tout $m \geq n_0$, on a $P_m = P_{n_0}$.

Définition 71 Soit n_0 l'entier défini précédemment. Si $P_{n_0} \neq \emptyset$ on dit que $\langle P_0, \dots, P_{n_0}, \emptyset \rangle$ est la base de P . Si $P_{n_0} = \emptyset$ la base de P est alors $\langle P_0, \dots, P_{n_0} \rangle$.

Donc un programme P induit intrinsèquement une hiérarchie, sa base, pour laquelle plus n est élevé, plus l'information contenue dans P_n est exceptionnelle.

Définition 72 Soit un programme P et soit $\langle P_0, \dots, P_n \rangle$ la base de P . On définit $\rho_p : \text{Prog} \rightarrow \mathbb{N}$, la fonction de rang, comme suit : $\rho_p(P') = \min\{i \in \mathbb{N} : P' \text{ est } P_i\text{-consistant}\}$ si P est consistant, sinon $\rho_p(P') = n$.

Notons que si $P' \subseteq P''$ alors $\rho_p(P') \leq \rho_p(P'')$.

Définition 73 Soient deux programmes P et P' . On définit la révision hiérarchique (*ranked revision*) de P par P' , notée $P \circ_{rk} P'$, par :

$$P \circ_{rk} P' = P_{\rho_p(P')} \cup P'$$

C'est-à-dire que ce sont les règles les moins exceptionnelles de P qui s'accordent avec la nouvelle information qui sont conservées dans le nouveau programme.

4.1.4 Révision gangue

Cet opérateur est une extension du précédent. Sa définition a pour but de capturer plus de transitivité (voir [BP96, BMP97]).

Soit $I_P(P')$ l'ensemble des sous-ensembles maximaux de P consistants avec P' et contenant $P_{\rho_P(P')}$. Lorsqu'il n'y a pas d'ambiguïté $I_P(P')$ est noté $I(P')$. On définit $h : \text{Prog} \rightarrow \mathcal{P}(P)$ par $h(P') = \bigcap I(P')$

Définition 74 La révision gangue (*hull revision*) d'un programme P par un programme P' , notée $P \circ_h P'$ est définie par :

$$P \circ_h P' = h(P') \cup P'$$

Remarque 3 Par définition il est facile de voir que

$$C_{fc}(P \circ_{rk} P') \subseteq C_{fc}(P \circ_h P')$$

L'opérateur de révision gangue $\circ_h$ peut donc être considéré comme une extension de l'opérateur de révision hiérarchique $\circ_{rk}$.

L'opérateur de révision gangue permet de garder des règles n'ayant rien à voir avec les contradictions induites par l'ajout de la nouvelle information mais qui auraient été supprimées par l'opérateur de révision hiérarchique car elles appartenaient à une couche plus exceptionnelle de la base de l'ancienne information.

4.1.5 Révision gangue étendue

Rappelons que, dans la définition de la révision gangue de P par P' , on calcule tout d'abord l'ensemble $I(P')$ des sous-ensembles de P maximaux consistants avec P' et contenant $P_{\rho_P(P')}$. Nous avons alors une approche très drastique en posant $P \circ_h P' = (\bigcap I(P')) \cup P'$. Ce que l'on souhaite donc ici est d'avoir une solution un peu moins sceptique, on considérera donc $I(P')$ comme un flock. L'idée de cet opérateur est donc assez proche de celle de l'opérateur de mise à jour factuelle.

Soient deux programmes P et P' , on définit $P \circ_{eh} P'$:

$$P \circ_{eh} P' = \begin{cases} \langle H_1 \cup P', \dots, H_n \cup P' \rangle & \text{si } I(P') = \{H_1, \dots, H_n\} \\ P' & \text{si } I(P') = \emptyset \end{cases}$$

On peut noter que le résultat est un flock de programmes. Supposons donc à présent, plus généralement, que $\mathcal{F}$ est un flock de programmes : $\mathcal{F} = \langle Q_1, \dots, Q_n \rangle$. On définit alors la *révision gangue étendue* (extended hull révision) $\mathcal{F} \circ_{eh} P$ par :

$$\mathcal{F} \circ_{eh} P = (Q_1 \circ_{eh} P) \cdot (Q_2 \circ_{eh} P) \cdots (Q_n \circ_{eh} P)$$

où $\cdot$ est la concaténation de flocks.

Notons qu'avec cette définition l'opérateur $\circ_{eh}$ est une extension de l'opérateur $\circ_h$, i.e. $C_{fc}(P \circ_h P') \subseteq C_{fc}(P \circ_{eh} P')$. C'est pour cette raison que l'opérateur $\circ_{eh}$ est appelé *révision gangue étendue*.

4.1.6 Exemples

Dans cette section on donne des exemples illustrant les différences de comportement entre les opérateurs de révision hiérarchique, révision gangue et révision gangue étendue.

Exemple 11 Soit $P = \{b \rightarrow f;$
 $b \rightarrow w;$
 $o \rightarrow b;$
 $o \rightarrow \neg f\}$

où b, o, f, w signifient respectivement oiseaux, autruches, volent et ont des ailes.

Il est facile de voir que la base de P est $\langle P_0, P_1, P_2 \rangle$ où

$$\begin{aligned} P_0 &= P, \\ P_1 &= \{o \rightarrow b; o \rightarrow \neg f\} \text{ et} \\ P_2 &= \emptyset. \end{aligned}$$

Notons que $\rho_p(b) = 0$ d'où $I(b) = P_0 = P$ et $P \circ_{rk} \{b\} = P \circ_h \{b\} = P \circ_{eh} \{b\} = P \cup \{b\}$,
 $C_{fc}(P \circ_{rk} \{b\}) = \{b, f, w\}$.

Pour le même P , on a $\rho_p(o) = 1$. Puisque l'ensemble $\{o \rightarrow b; o \rightarrow \neg f; b \rightarrow w\}$ est la seule extension de P_1 consistante avec $\{o\}$ on a $h(o) = \{o \rightarrow b; o \rightarrow \neg f; b \rightarrow w\}$ d'où

$$P \circ_h \{o\} = P \circ_{eh} \{o\} = \{o \rightarrow b; o \rightarrow \neg f; b \rightarrow w\} \cup \{o\}$$

Puisque $\rho_p(o) = 1$ on a $P \circ_{rk} \{o\} = \{o \rightarrow b; o \rightarrow \neg f\} \cup \{o\}$. D'où

$$C_{fc}(P \circ_h \{o\}) = C_{fc}(P \circ_{eh} \{o\}) = \{b, o, \neg f, w\}$$

mais

$$C_{fc}(P \circ_{rk} \{o\}) = \{b, o, \neg f\}$$

On voit sur cet exemple que les opérateurs de révision gangue et révision gangue étendue permettent de conserver plus d'information que l'opérateur de révision hiérarchique. Un autre exemple classique est :

Exemple 12 Soit $P = \{m \rightarrow s;$
 $c \rightarrow m;$
 $c \rightarrow \neg s;$
 $n \rightarrow c;$
 $n \rightarrow s\}$

où m, s, c, n signifient respectivement mollusques, coquillages, céphalopodes et nautilus (on ne détaille pas les calculs).

La base est $\langle P_0, P_1, P_2, P_3 \rangle$ où

$$\begin{aligned} P_0 &= P, \\ P_1 &= \{c \rightarrow m; c \rightarrow \neg s; n \rightarrow c; n \rightarrow s\}, \\ P_2 &= \{n \rightarrow c; n \rightarrow s\} \text{ et} \\ P_3 &= \emptyset. \end{aligned}$$

On a $C_{fc}(P \circ_h \{n\}) = \{n, c, s, m\} = C_{fc}(P \circ_{eh} \{n\})$ et $C_{fc}(P \circ_{rk} \{n\}) = \{n, c, s\}$. Cela montre que les révisions gangues permettent plus d'inférences que la révision hiérarchique.

Dans d'autres cas les trois opérateurs coïncident, par exemples $C_{fc}(P \circ_h \{c, \neg n\}) = \{c, \neg n, m, \neg s\} = C_{fc}(P \circ_{rk} \{c, \neg n\})$.

À présent, un dernier exemple montre qu'en général l'opérateur de révision gangue étendue permet plus d'inférences que l'opérateur de révision gangue.

Exemple 13 Soit le programme $P = \{r_1, \dots, r_5\}$ où $r_1 = a, b \rightarrow c$,

$$r_2 = a, \neg c \rightarrow d,$$

$$r_3 = b, \neg c \rightarrow d,$$

$$r_4 = a,$$

$$r_5 = b.$$

Soit $P' = \{\neg c\}$. La base de P est $\langle P_0, P_1, P_2 \rangle$ avec

$$P_0 = P,$$

$$P_1 = \{r_2, r_3\} \text{ et}$$

$$P_3 = \emptyset.$$

On trouve facilement $\rho_p(P') = 1$ et $I(P') = \{P_1 \cup \{r_1, r_4\}, P_1 \cup \{r_1, r_5\}, P_1 \cup \{r_4, r_5\}\}$. Donc $h(P') = P_1$ et alors $d \notin C_{fc}(P \circ_h P') = C_{fc}(P_1 \cup \neg c)$. Alors que $d \in C_{fc}(P \circ_{eh} P')$ car pour tout $Q \in I(P')$ on a $d \in C_{fc}(Q \cup \{\neg c\})$.

4.1.7 Coder l'opérateur de révision gangue

On montre dans cette section comment, via un codage simple, calculer l'opérateur de révision gangue en utilisant l'opérateur de mise à jour factuelle défini section 4.1.2.

La base $\langle P_0, \dots, P_n \rangle$ de P se calcule facilement. Soient $\ell' : P \longrightarrow \{r_1, \dots, r_m\}$ et $\ell'' : P' \longrightarrow \{q_1, \dots, q_k\}$, deux bijections où les r_i et les q_j sont de nouveaux atomes.

On définit $\ell^{-1} : \{r_1, \dots, r_m\} \cup \{q_1, \dots, q_k\} \longrightarrow P \cup P'$ par $\ell^{-1}(r) = \ell'^{-1}(r)$ si $r \in P$ et $\ell^{-1}(r') = \ell''^{-1}(r')$ si $r' \in P'$.

Soit $M(P)$, la modification de P définie comme suit : chaque règle $L \rightarrow l$ de P est remplacée par la règle $r, L \rightarrow l$ où $r = \ell'(L \rightarrow l)$. De manière analogue, soit $M(P')$, la modification de P' définie par : chaque règle $L \rightarrow l$ de P' est remplacée par la règle $r, L \rightarrow l$ où $r = \ell''(L \rightarrow l)$.

Remarquons que les sous-ensembles maximaux de P qui sont consistants avec P' et qui contiennent $P_{\rho_p(P')}$ sont alors ceux qui correspondent aux sous-ensembles maximaux de la base de faits $\ell'(P)$ calculée comme la mise à jour par $\ell''(P') \cup \ell'(P_{\rho_p(P')})$ sous les contraintes $M(P) \cup M(P')$.

On a plus précisément :

$$P \circ_h P' = \ell^{-1} \left\{ \bigcap [\ell'(P) \diamond_{M(P) \cup M(P')} (\ell''(P') \cup \ell'(P_{\rho_p(P')}))] \right\}$$

Et, de manière analogue :

$$P \circ_{eh} P' = \ell^{-1} [\ell'(P) \diamond_{M(P) \cup M(P')} (\ell''(P') \cup \ell'(P_{\rho_p(P')}))]$$

De façon à illustrer ce codage, considérons l'exemple suivant :

Exemple 14 $P = \{b \rightarrow f,$
 $b \rightarrow w,$
 $o \rightarrow b,$
 $o \rightarrow \neg f\}$

Soit $P' = \{o\}$, on définit $\ell : P \cup P' \longrightarrow \{1,2,3,4,5\}$ tel que : $M(P) = \{1, b \rightarrow f,$
 $2, b \rightarrow w,$
 $3, o \rightarrow b,$
 $4, o \rightarrow \neg f\}$

et $M(P') = \{5 \rightarrow o\}$

Nous avons vu dans les exemples précédents que $P_{\rho_P(P')} = \{o \rightarrow b, o \rightarrow \neg f\}$ donc $\ell(R_{\rho_P(L)}) = \{3,4\}$. D' où

$$\ell(P) \diamond_{M(P) \cup M(P')} (\ell(P') \cup \ell(P_{\rho_P(P')})) = \{1,2,3,4\} \diamond_{M(P) \cup M(P')} \{5\} \cup \{3,4\} = \langle \{2,3,4,5\} \rangle$$

$$\text{Et alors } C_{fc}(P \circ_h P') = C_{fc}(\ell^{-1}(\{2,3,4,5\})) = C_{fc}(\{o \rightarrow b, o \rightarrow \neg f, b \rightarrow w, o\}) = \{o, b, \neg f, w\}.$$

4.2 Propriétés logiques

Dans cette section nous étudions les propriétés de rationalité des opérateurs définis à la section précédente.

Théorème 62 *L'opérateur de mise à jour factuelle $\diamond_P$ est un opérateur de mise à jour syntaxique, c'est-à-dire qu'il vérifie les postulats (U1), (U2), (U3), (U5), (U6) et (U8).*

Preuve :

(U1) On veut montrer que $C_{fc}^P(L') \subseteq C_{fc}^P(\mathcal{F} \diamond_P L')$, i.e. que $C_{fc}^P(L')$ est un sous-ensemble de $C_{fc}^P(\mathcal{F} \diamond_P L')$. Si L' est P -inconsistant ou si tous les éléments de $\mathcal{F}$ sont P -inconsistants, alors le résultat est direct.

Supposons alors que $\mathcal{F} \diamond_P L' = \langle L_1, \dots, L_n \rangle$. Par définition de $\diamond_P$, on a $L' \subseteq L_i$ pour tout $i = 1, \dots, n$. De ce fait $L' \subseteq \bigcap_{i=1}^n C_{fc}^P(L_i) = C_{fc}^P(\mathcal{F} \diamond_P L')$. La relation de conséquence utilisée étant idempotente et monotone on obtient l'inclusion recherchée.

(U2) Supposons que $C_{fc}^P(L') \subseteq C_{fc}^P(\mathcal{F})$. On veut montrer que $C_{fc}^P(\mathcal{F}) = C_{fc}^P(\mathcal{F} \diamond_P L')$. Si $C_{fc}^P(\mathcal{F}) = Lit$ ou $C_{fc}^P(L') = Lit$ alors le résultat suit directement de la définition. Supposons donc que $C_{fc}^P(\mathcal{F}) \neq Lit$ et $C_{fc}^P(L') \neq Lit$. Par hypothèse, on peut supposer que $\mathcal{F} = \langle L_1, \dots, L_n \rangle$ et $\mathcal{F} \diamond_P L' = (L_{i_1} \diamond_P L') \cdots (L_{i_K} \diamond_P L')$ où l'ensemble $\{L_{i_1} \dots L_{i_K}\}$ est le sous-ensemble maximal de $\{L_1 \dots L_n\}$ tel que chaque L_{i_j} est P -consistant. Puisque par hypothèse $C_{fc}^P(L') \subseteq C_{fc}^P(L_{i_j})$ on a que $L_{i_j} \diamond_P L' = L_{i_j} \cup L'$ et comme la relation de conséquence est idempotente et monotone $C_{fc}^P(L_{i_j} \cup L') = C_{fc}^P(L_{i_j})$. D'où

$$C_{fc}^P(\mathcal{F} \diamond_P L') = C_{fc}^P(L_{i_1} \cup L', \dots, L_{i_K} \cup L') = \bigcap_{j=1}^K C_{fc}^P(L_{i_j}) = \bigcap_{i=1}^n C_{fc}^P(L_i) = C_{fc}^P(\mathcal{F})$$

où l'avant dernière égalité est due au fait que si L_i n'est égal à aucun L_{i_j} alors $C_{fc}^P(L_i) = Lit$.

(U3) Direct par définition de $\diamond_P$.

(U5) On veut montrer $C_{fc}^P(\mathcal{F} \diamond_P (L \wedge L')) \subseteq C_{fc}^P((\mathcal{F} \diamond_P L) \wedge L')$. On prouve d'abord ce résultat lorsque le flock $\mathcal{F}$ n'est composé que d'un seul élément : $\mathcal{F} = \langle H \rangle$. Si H est P -inconsistant ou $L \cup L'$ est P -inconsistant le résultat est direct. Supposons donc que H et $P \cup P'$ sont P -consistants. On a par définition

$$\begin{aligned} H \diamond_P (L \wedge L') &= \langle L_1 \cup L \cup L', \dots, L_n \cup L \cup L' \rangle \\ (H \diamond_P L) \wedge L' &= \langle K_1 \cup L, \dots, K_p \cup L \rangle \wedge L' \\ &= \langle K_1 \cup L \cup L', \dots, K_p \cup L \cup L' \rangle \end{aligned}$$

où L_i est un sous-ensemble maximal de H tel que $L_i \cup L \cup L'$ est P -consistant, pour tout $i = 1, \dots, n$, et K_j est un sous-ensemble maximal de H tel que $K_j \cup L$ est P -consistant, pour tout $j = 1, \dots, p$. Notons que soit $C_{fc}^P(K_j \cup L \cup L') = Lit$, soit il existe un m tel que $K_j = L_m$. De ce fait si $l \in \bigcap_{i=1}^n C_{fc}^P(P \cup L_i \cup L' \cup L'')$ alors $l \in \bigcap_{j=1}^p C_{fc}^P(P \cup K_j \cup L' \cup L'')$.

On peut à présent prouver le cas général. Soit $\mathcal{F} = \langle L_1, \dots, L_n \rangle$, si $\mathcal{F}$ est P -inconsistant ou $L \cup L'$ est P -inconsistant le résultat est trivial. Supposons donc que $\mathcal{F}$ et $P \cup P'$ sont P -consistants. Alors

$$\begin{aligned} \mathcal{F} \diamond_P (L \wedge L') &= (L_{i_1} \diamond_P (L \cup L')) \cdots (L_{i_k} \diamond_P (L \cup L')) \\ (\mathcal{F} \diamond_P L) \wedge L' &= ((L_{i_1} \diamond_P L \cdots L_{i_k} \diamond_P L) \wedge L') \\ &= ((L_{i_1} \diamond_P L) \wedge L') \cdots ((L_{i_k} \diamond_P L) \wedge L') \end{aligned}$$

où $\{L_{i_1}, \dots, L_{i_k}\}$ est le sous-ensemble maximal de $\{L_1, \dots, L_n\}$ tel que chaque élément est P -consistant. D'après le premier cas on a

$$C_{fc}^P((L_{i_j} \diamond_P (L \wedge L')) \subseteq C_{fc}^P((L_{i_j} \diamond_P L) \wedge L') \quad \text{pour } j = 1, \dots, k$$

Donc $C_{fc}^P((\mathcal{F} \diamond_P (L \wedge L')) \subseteq C_{fc}^P((\mathcal{F} \diamond_P L) \wedge L')$.

(U6) Supposons $C_{fc}^P(L_2) \subseteq C_{fc}^P(\mathcal{F} \diamond_P L_1)$ et $C_{fc}^P(L_1) \subseteq C_{fc}^P(\mathcal{F} \diamond_P L_2)$. On veut montrer que $C_{fc}^P(\mathcal{F} \diamond_P L_1) = C_{fc}^P(\mathcal{F} \diamond_P L_2)$. Si $\mathcal{F}$, L_1 ou L_2 est P -inconsistant le résultat est direct. Supposons donc que les trois sont P -consistants.

Supposons d'abord que $\mathcal{F} = \langle L \rangle$. Par hypothèse il est clair que $L_2 \subseteq C_{fc}^P(L \diamond_P L_1)$ et $L_1 \subseteq C_{fc}^P(L \diamond_P L_2)$. Posons

$$\begin{aligned} L \diamond_P L_1 &= \langle L_1^1, \dots, L_{n_1}^1 \rangle \wedge L_1 \\ L \diamond_P L_2 &= \langle L_1^2, \dots, L_{n_2}^2 \rangle \wedge L_2 \end{aligned}$$

où L_i^1 est un sous-ensemble maximal de L tel que $L_i^1 \cup L_1$ est P -consistant pour tout $i = 1, \dots, n_1$ et L_j^2 est un sous-ensemble maximal de L tel que $L_j^2 \cup L_2$ est P -consistant pour tout $j = 1, \dots, n_2$. Par hypothèse il est facile de voir que :

(a) $L_1 \cup L_j^2$ est P -consistant pour tout $j = 1, \dots, n_2$ et

(b) $L_2 \cup L_i^1$ est P -consistant pour tout $i = 1, \dots, n_1$.

De (b) on déduit $\forall i \in \{1, \dots, n_1\} \exists j \in \{1, \dots, n_2\}$ tel que $L_i^1 \subseteq L_j^2$ et de (a) on a $\forall j \in \{1, \dots, n_2\} \exists i \in \{1, \dots, n_1\}$ tel que $L_j^2 \subseteq L_i^1$. Mais cela implique, par la maximalité des ensembles L_i^1 et L_j^2 , que $n_1 = n_2$ et il y a une permutation σ de $\{1, \dots, n_1\}$ telle que

$L_i^1 = L_{\sigma(i)}^2$. Sans perte de généralité on peut supposer que σ est l'identité. Remarquons finalement que

$$\begin{aligned} C_{fc}^P(\langle L_1^1, \dots, L_{n_1}^1 \rangle \wedge L_1) &= C_{fc}^P(\langle L_1^1, \dots, L_{n_1}^1 \rangle \wedge (L_1 \cup L_2)) \\ &= C_{fc}^P(\langle L_1^2, \dots, L_{n_1}^2 \rangle \wedge (L_2 \cup L_1)) \\ &= C_{fc}^P(\langle L_1^2, \dots, L_{n_1}^2 \rangle \wedge L_2) \end{aligned}$$

Les première et troisième égalités viennent de l'hypothèse et du fait que la relation de conséquence est idempotente et monotone.

On prouve à présent le cas général. Soit $\mathcal{F} = \langle H_1, \dots, H_n \rangle$, supposons que $C_{fc}^P(L_2) \subseteq C_{fc}^P(\mathcal{F} \diamond_P L_1)$ et $C_{fc}^P(L_1) \subseteq C_{fc}^P(\mathcal{F} \diamond_P L_2)$. Alors

$$\mathcal{F} \diamond_P L_i = (H_{j_1} \diamond_P L_i) \cdots (H_{j_k} \diamond_P L_i) \quad i = 1, 2$$

Mais il est facile de voir que $L_1 \subseteq C_{fc}^P(H_{j_m} \diamond_P L_2)$ et $L_2 \subseteq C_{fc}^P(H_{j_m} \diamond_P L_1)$ pour tout $m = 1, \dots, k$. Alors, à l'aide du premier cas, on déduit $C_{fc}^P(H_{j_m} \diamond_P L_1) = C_{fc}^P(H_{j_m} \diamond_P L_2)$ et alors $C_{fc}^P(\mathcal{F} \diamond_P L_1) = C_{fc}^P(\mathcal{F} \diamond_P L_2)$.

(U8) Cette propriété est vérifiée trivialement par définition de l'opérateur.

□

On peut souligner que le postulat (U7) n'est pas traduisible ici puisque nous n'avons pas de disjonction sur les faits. L'opérateur de mise à jour factuelle ne vérifie bien sûr pas le postulat (U4) d'indépendance de syntaxe puisque c'est un opérateur syntaxique. De la même façon les trois autres opérateurs ne satisferont pas (R4). On montre cela au théorème 66.

Théorème 63 *L'opérateur de révision hiérarchique $\circ_{rk}$ est un opérateur de révision syntaxique, c'est-à-dire qu'il satisfait les propriétés (R1), (R2), (R3), (R5) et (R6).*

Preuve :

- (R1) On veut montrer $C_{fc}(P') \subseteq C_{fc}(P \circ_{rk} P')$. C'est clairement vrai car $P \circ_{rk} P' = P_{\rho_p(P')} \cup P'$.
- (R2) Supposons que $P \wedge P'$ est consistant, i.e. $C_{fc}(P \cup P') \neq Lit$. On veut montrer $P \circ_{rk} P' = P \cup P'$. Ce que l'on a directement puisque $C_{fc}(P \cup P') \neq Lit$ implique $\rho_p(P') = 0$.
- (R3) Supposons que P' est consistant, i.e. $C_{fc}(P') \neq Lit$. On veut montrer que $P \circ_{rk} P'$ est également consistant, i.e. $C_{fc}(P \circ_{rk} P') \neq Lit$. Cela est vérifié car $P \circ_{rk} P' = P_{\rho_p(P')} \cup P'$ et $P_{\rho_p(P')}$ est par définition consistant avec P' .
- (R5) et (R6) Supposons que $C_{fc}((P \circ_{rk} P') \wedge P'') \neq Lit$, i.e. $(P \circ_{rk} P') \cup P''$ est consistant (sinon (R5) est trivial). On veut montrer $C_{fc}((P \circ_{rk} P') \wedge P'') = C_{fc}(P \circ_{rk} (P' \wedge P''))$. Pour cela il est suffisant de montrer $P \circ_{rk} (P' \cup P'') = (P \circ_{rk} P') \cup P''$. Par hypothèse $(P_{\rho_p(P')} \cup P') \cup P''$ est consistant. Alors $\rho_p(P' \cup P'') \leq \rho_p(P')$ et donc $\rho_p(P' \cup P'') = \rho_p(P')$. On conclut alors facilement.

□

Théorème 64 *L'opérateur de révision gangue $\circ_h$ satisfait les postulats (R1), (R2) et (R3). Il ne satisfait pas (R5) et (R6).*

Preuve :

- (R1) On veut montrer $C_{fc}(P') \subseteq C_{fc}(P \circ_h P')$. C'est clairement vrai car $P \circ_h P' = h(P') \cup P'$.
 (R2) Supposons que $P \wedge P'$ est consistant, i.e. $C_{fc}(P \cup P') \neq Lit$. On veut montrer $P \circ_h P' = P \cup P'$. Ce que l'on a directement puisque $C_{fc}(P \cup P') \neq Lit$ implique $\rho_p(P') = 0$.
 (R3) Supposons que P' est consistant, i.e. $C_{fc}(P') \neq Lit$. On veut montrer que $P \circ_h P'$ est également consistant, i.e. $C_{fc}(P \circ_h P') \neq Lit$. Cela est vérifié car $P \circ_h P' = h(P') \cup P'$ et $h(P')$ est par définition contenu dans un ensemble consistant avec P' .

Considérons le contre-exemple suivant pour (R5) : Soit le programme $P = \{b \rightarrow w,$
 $w \rightarrow w',$
 $w' \rightarrow f,$
 $o \rightarrow b,$
 $o \rightarrow \neg f\}$

La base de P est $\langle P_0, P_1, P_2 \rangle$ avec

$$\begin{aligned} P_0 &= P, \\ P_1 &= \{o \rightarrow b, o \rightarrow \neg f\} \text{ et} \\ P_2 &= \emptyset \end{aligned}$$

Posons $P' = \{o\}$ et $P'' = \{w'\}$. Il n'est pas dur de voir que $h(P') = P_1$ et $h(P' \cup P'') = P_1 \cup \{b \rightarrow w, w \rightarrow w'\}$. Donc $C_{fc}((P \circ_h P') \cup P'') = \{o, w', b, \neg f\}$ et $C_{fc}(P \circ_h (P' \cup P'')) = \{o, w', b, \neg f, w\}$. Alors $C_{fc}(P \circ_h (P' \cup P'')) \not\subseteq C_{fc}((P \circ_h P') \cup P'')$, c'est-à-dire que (R5) n'est pas vérifié.

Pour (R6) considérons le contre-exemple suivant : Soit $P = \{r_0, r_1, r_2\}$ où $r_0 = a \rightarrow c,$
 $r_1 = e \rightarrow \neg c,$
 $r_2 = b \rightarrow \neg c$

Soit $P' = \{a, e\}$ et $P'' = \{b\}$. La base de P est $\langle P, \emptyset \rangle$. Il est alors facile de voir que $P_{\rho_p(P')} = P_{\rho_p(P' \cup P'')} = \emptyset$ et

$$\begin{aligned} I(P') &= \{\{r_0, r_2\}, \{r_1, r_2\}\} \text{ et} \\ I(P' \cup P'') &= \{\{r_0\}, \{r_1, r_2\}\} \end{aligned}$$

Donc $h(P') = \{r_2\}$ et $h(P' \cup P'') = \emptyset$. Donc $\neg c \in C_{fc}((P \circ_h P') \cup P'')$ et $\neg c \notin C_{fc}(P \circ_h (P' \cup P''))$, c'est-à-dire que (R6) n'est pas vérifié. □

Théorème 65 *L'opérateur de révision gangue étendue $\circ_{eh}$ satisfait (R1), (R3) et (R5). Il ne vérifie pas (R2), (U2), (U6) et (R6). Il satisfait une version faible de (R2) : si P est consistant avec tous les éléments de $\mathcal{F}$, alors $\mathcal{F} \circ_{eh} P = \mathcal{F} \wedge P$.*

Preuve :

On prouve d'abord les postulats vérifiés :

- (R1) ce postulat est prouvé de manière similaire à (U1) pour l'opérateur $\diamond_P$ (théorème 62).
 (R3) est satisfait directement par définition.

(R2w) La version faible de (R2) est également satisfaite trivialement.

(R5) Considérons d'abord le cas $\mathcal{F} = P$. On veut montrer que $C_{fc}(P \circ_{eh} (P' \wedge P'')) \subseteq C_{fc}((P \circ_{eh} P') \wedge P'')$. Supposons que

$$\begin{aligned} P \circ_{eh} P' &= \langle Q_1 \cup P', \dots, Q_n \cup P' \rangle \\ P \circ_{eh} (P' \wedge P'') &= P \circ_{eh} (P' \cup P'') = \langle H_1 \cup P' \cup P'', \dots, H_k \cup P' \cup P'' \rangle \end{aligned}$$

Alors $(P \circ_{eh} P') \wedge P'' = \langle Q_1 \cup P' \cup P'', \dots, Q_n \cup P' \cup P'' \rangle$. Si $C_{fc}((P \circ_{eh} P') \wedge P'') = Lit$ on a directement le résultat. Sinon il existe un Q_i tel que $Q_i \cup P' \cup P''$ est consistant. Mais puisque Q_i est un sous-ensemble de P maximal consistant avec P' et contenant $P_{\rho_p(P')}$ nécessairement $\rho_p(P') = \rho_p(P' \cup P'')$. En fait pour tout $i = 1, \dots, n$ soit $Q_i \cup P' \cup P''$ est inconsistent, soit il existe un $j \leq k$ tel que $Q_i = H_j$. Pour montrer cela supposons que $Q_i \cup P' \cup P''$ est consistant, alors Q_i est un sous-ensemble maximal consistant avec $P' \cup P''$ contenant $P_{\rho_p(P')} = P_{\rho_p(P' \cup P'')}$, i.e. $Q_i = H_j$ pour un j , par définition des ensembles H_q pour $q = 1, \dots, k$. Finalement, de ce fait, on obtient facilement que $\bigcap_{j=1}^k C_{fc}(H_j \cup P' \cup P'') \subseteq \bigcap_{i=1}^n C_{fc}(Q_i \cup P' \cup P'')$, i.e. $C_{fc}(P \circ_{eh} (P' \wedge P'')) \subseteq C_{fc}((P \circ_{eh} P') \wedge P'')$.

Le cas général, lorsque $\mathcal{F} = \langle P_1, \dots, P_n \rangle$, se déduit du cas précédent par définition de $\circ_{eh}$ et en utilisant le même argument que pour la preuve de (U5) dans le théorème 62.

On donne un contre-exemple pour chacun des postulats non vérifiés.

(R2) Prenons $\mathcal{F} = \langle \{a\}, \{\neg b\} \rangle$ et $P = \{b\}$. Alors

$$C_{fc}(\mathcal{F} \wedge P) = C_{fc}(\langle \{a, b\}, \{\neg b, b\} \rangle) = \{a, b\}$$

donc $\mathcal{F} \wedge P$ est consistant. Mais $C_{fc}(\mathcal{F} \circ_{eh} P) = C_{fc}(\langle \{a, b\}, \{b\} \rangle) = \{b\}$. (R2) n'est donc pas vérifié.

(U2) Soient $P = \{a\}$ et $P' = \{a \rightarrow b\}$. Alors $\emptyset = C_{fc}(P') \subseteq C_{fc}(P) = \{a\}$. Mais $C_{fc}(P \circ_{eh} P') = C_{fc}(P \cup P') = \{a, b\}$. Donc (U2) n'est pas satisfait.

(U6) Prenons $P_1 = \{a \rightarrow b\}$, $P_2 = \{c \rightarrow d\}$ et $Q = \{a\}$. On a clairement $C_{fc}(P_1) = C_{fc}(P_2) = \emptyset$, donc l'hypothèse de (U6) est vérifiée. Pourtant $C_{fc}(Q \circ_{eh} P_1) = C_{fc}(Q \cup P_1) = \{a, b\} \neq \{a\} = C_{fc}(Q \cup P_2) = C_{fc}(Q \circ_{eh} P_2)$.

(R6) On utilise le même contre-exemple que dans le théorème 64.

□

Théorème 66 *Tous les opérateurs définis précédemment sont sensibles à la syntaxe, i.e. (R4) n'est pas satisfait par les opérateurs $\diamond_P$, $\circ_{rk}$, $\circ_h$ et $\circ_{eh}$.*

Preuve :

Nous donnons d'abord un contre-exemple pour $\diamond_P$. Soit $P = \{a \rightarrow b\}$. Prenons $L_1 = \{a, b\}$, $L_2 = \{a\}$ et $L' = \{\neg a\}$. On a clairement $C_{fc}^P(L_1) = C_{fc}^P(L_2) = \{a, b\}$. Il est facile de voir que $L_1 \diamond_P L' = \{b, \neg a\}$. Donc $C_{fc}^P(L_1 \diamond_P L') = \{b, \neg a\}$. Mais on a également facilement $L_2 \diamond_P L' = \{\neg a\}$ donc $C_{fc}^P(L_2 \diamond_P L') = \{\neg a\}$. Finalement $C_{fc}^P(L_1 \diamond_P L') \neq C_{fc}^P(L_2 \diamond_P L')$.

On donne à présent un même contre-exemple pour $\circ_{rk}$, $\circ_h$ et $\circ_{eh}$. Soient $P_1 = \{a \rightarrow b\}$, $P_2 = \{a \rightarrow c\}$ et $P' = \{a\}$. Alors $C_{fc}(P_1) = \emptyset = C_{fc}(P_2)$. Mais $C_{fc}(P_1 \circ P') = C_{fc}(P_1 \cup P') = \{a, b\} \neq \{a\} = C_{fc}(P_2 \cup P') = C_{fc}(L_2 \circ P')$ quel que soit l'opérateur $\circ \in \{\circ_{rk}, \circ_h, \circ_{eh}\}$.

□

4.3 Révision gangue sélective

Dans la section précédente on a montré que les deux tentatives “sceptiques” d’extension de l’opérateur de révision hiérarchique ne satisfont que très peu de propriétés logiques. Dans cette section nous définissons une nouvelle extension de la révision gangue étendue, plus aventureuse, utilisant une fonction de sélection. Cet opérateur est un opérateur de révision syntaxique.

Soit une fonction S associant, à un ensemble de programmes, un programme i.e. $S : \mathcal{P}(Prog) \rightarrow Prog$. On dit que S est une fonction de sélection ssi elle vérifie les propriétés suivantes :

- (i) $S(\emptyset) = \emptyset$
- (ii) Si $D \neq \emptyset$, alors $S(D) \in D$

Ce type de fonctions de sélection est nommé fonctions de sélection à choix maximal (cf section 1.2.1).

Soient deux programmes P et P' , $I_P(P')$ est défini comme dans la section 4.1.3, i.e. l’ensemble des sous-ensembles de P maximaux consistants avec P' et contenant $P_{\rho_p(P')}$. Soit une fonction de sélection S . On définit l’opérateur de *révision gangue sélective* (selection hull revision), noté $\circ_{sh}$, par :

$$P \circ_{sh} P' = S(I_P(P')) \cup P'$$

La définition suivante nous présente une classe de fonctions de sélection pour lesquelles l’opérateur $\circ_{sh}$ est un opérateur de sélection syntaxique.

Définition 75 Une fonction de sélection S est dite sensée si et seulement si elle vérifie la propriété suivante : pour tous programmes P , P' et P''

$$I_P(P') \cap I_P(P' \cup P'') \neq \emptyset \Rightarrow S(I_P(P')) = S(I_P(P' \cup P''))$$

La propriété exigée de S par la définition 75 est très intuitive : si vous devez choisir parmi les éléments de $\{Q\} \cup D_1$ et que votre choix est Q , cela signifie que vous considérez que Q est l’élément qui correspond le plus à (le plus “proche” de) P , alors si vous devez choisir parmi les éléments de $\{Q\} \cup D_2$ alors que les éléments de D_2 sont contenus dans les éléments de D_1 , vous devez alors choisir Q .

Théorème 67 Si S est une fonction de sélection sensée, alors l’opérateur de révision gangue sélective $\circ_{sh}$ ainsi défini est un opérateur de révision syntaxique, i.e. il satisfait (R1), (R2), (R3), (R5) et (R6).

Preuve :

(R1),(R2) et (R3) sont obtenues directement par définition de $\circ_{sh}$.

Pour prouver (R5) et (R6) on va montrer que si $(P \circ_{sh} P') \cup P''$ est consistant, alors $(P \circ_{sh} P') \cup P'' = P \circ_{sh} (P' \cup P'')$. Ce qui est clairement plus fort que (R5) et (R6).

Supposons que $I(P') = \{Q\} \cup D_1$ et $S(I(P')) = Q$. Par hypothèse $Q \cup P' \cup P''$ est consistant, donc $\rho_p(P') = \rho_p(P' \cup P'')$ et $Q \in I(P' \cup P'')$. Alors $I(P' \cup P'') = \{Q\} \cup D_2$. Puisque $\rho_p(P') = \rho_p(P' \cup P'')$ et $P' \subseteq P' \cup P''$ on a pour tout R de D_2 qu'il existe R' de D_1 tel que $R \subseteq R'$. Et des propriétés de S on a $S(I(P' \cup P'')) = Q$, on conclut alors facilement.

□

On peut remarquer que la définition de l'opérateur $\circ_{sh}$ est une extension de $\circ_{eh}$, i.e. $C_{fc}(P \circ_{eh} P') \subseteq C_{fc}(P \circ_{sh} P')$. On a donc les inclusions de la figure 4.1.

$$C_{fc}(P \circ_{rk} P') \subseteq C_{fc}(P \circ_h P') \subseteq C_{fc}(P \circ_{eh} P') \subseteq C_{fc}(P \circ_{sh} P')$$

FIG. 4.1 – Relations entre les opérateurs de révision chaînage avant

Nous donnons à présent un exemple de fonction de sélection sensée, montrant qu'il est possible pour une fonction de sélection de satisfaire la propriété demandée au théorème 67.

On note $|Q|$ la cardinalité de l'ensemble Q . Soit $\{r_1, \dots, r_n\}$, une énumération sans répétition de règles et de faits. Soit π une fonction (pondération) $\pi : \{r_1, \dots, r_n\} \rightarrow \mathbb{N}$ telle que $\pi(r_i) = 2^i$. On étend additivement π à $\mathcal{P}(Prog)$ en posant

$$\pi(\{r_{i_1}, \dots, r_{i_k}\}) = \sum_{j=1}^k 2^{i_j}$$

Notons que si $P, P' \in Prog$ et $P \neq P'$, alors $\pi(P) \neq \pi(P')$.

Soit $<_\ell$, l'ordre lexicographique sur $\mathbb{N}^2$. On définit alors $S : \mathcal{P}(Prog) \rightarrow Prog$ par $S(\emptyset) = \emptyset$ et si D est non vide

$$S(D) = Q \text{ ssi } Q \in D \text{ et } \forall R \in D (R \neq Q \Rightarrow (|R|, \pi(R)) <_\ell (|Q|, \pi(Q)))$$

C'est-à-dire que S choisit parmi les ensembles de plus grande cardinalité, celui de poids le plus important. La pondération des ensembles est induite par la pondération des règles. Ces règles sont ordonnées selon un ordre strict. Cela assure donc que l'on a bien un seul ensemble sélectionné.

Il est facile de voir, à partir de la définition de π , que S est une fonction de sélection sensée.

Opérateur	Propriétés vérifiées	Propriétés non vérifiées
$\diamond_P$	U1+U2+U3+U5+U6+U8	R4
$\circ_{rk}$	R1+R2+R3+R5+R6	R4
$\circ_h$	R1+R2+R3	R4+R5+R6
$\circ_{eh}$	R1+R3+R5	R2+U2+ R4+R6+U6
$\circ_{sh}$	R1+R2+R3+R5+R6	R4

FIG. 4.2 – Propriétés des opérateurs syntaxiques

4.4 Conclusion

Nous avons proposé dans ce chapitre cinq opérateurs de changements basés sur le chaînage avant. Les idées menant à la définition de ces opérateurs sont très simples et rappellent, dans la plupart des cas, des méthodes bien connues [FKUV86, BKMS92, MT95].

L'intérêt de cette proposition est la définition d'une fonction de rang pour les opérateurs de révision et le fait que la relation de conséquence et la logique considérées pour tous ces opérateurs est le chaînage avant. Cela donne des opérateurs sensibles à la syntaxe, mais plus simples (en terme de complexité) que les opérateurs basés sur la logique classique et donc plus facilement implémentables. Nous montrons alors que la plupart de ces opérateurs satisfont suffisamment de conditions logiques de rationalité. Le tableau 4.2 résume les propriétés vérifiées.

On constate, en particulier que l'opérateur de mise à jour factuelle est un opérateur de mise à jour syntaxique. De plus l'opérateur de révision hiérarchique et celui de révision gangue sélective sont des opérateurs de révision syntaxique. Curieusement, les opérateurs de révision gangue et de révision gangue étendue, qui étendent de façon naturelle l'opérateur de révision hiérarchique ne vérifient que très peu de propriétés logiques.

L'originalité de ce travail est de proposer une définition syntaxique d'opérateurs de révision accompagnée de l'étude de leurs propriétés logiques. Cette étude des opérateurs, bien que très utile car permettant de caractériser leur comportement, est souvent absente des travaux proposant des méthodes pratiques de révision.

Tous ces opérateurs, basés sur le chaînage avant, sont facilement calculables. On peut noter en particulier que l'opérateur de révision hiérarchique est polynomial. Les autres opérateurs sont NP-complets mais leur complexité reste abordable.

Deuxième partie

Fusion

Chapitre 5

Introduction

Le tout est plus grand que la somme des parties.
Confucius

Le problème de la fusion de bases de connaissance est un problème central dans plusieurs domaines. En bases de données, un des enjeux majeurs des futures bases de données sera l'utilisation et l'intégration de données hétérogènes et réparties [SSU91]. Un travail important est fait en ce qui concerne le traitement de données hétérogènes [SAB⁺, CMH⁺94, MIR94]. Mais rares sont les travaux traitant du problème de l'inconsistance de données contradictoires, la plupart des approches supposent l'existence d'un oracle permettant de trancher en cas de conflit (voir par exemple [Sub94]). Une autre façon d'éviter le problème est de supposer l'existence d'une information sur la crédibilité relative des bases [Cho93, Cho95, Cho98]. Un tel oracle est difficilement concevable et on ne peut pas toujours disposer de cette information sur la crédibilité relative des bases, il est donc nécessaire de définir des méthodes pouvant traiter ces contradictions de manière adéquate.

En intelligence artificielle, il peut également être nécessaire de déterminer la connaissance du système lorsque les données fournies au système proviennent de sources différentes ou de déterminer la connaissance globale d'un système quand l'information de celui-ci est distribuée. Par exemple, si l'on désire concevoir un système expert codant la connaissance d'un groupe d'experts humains, il est alors sensé de coder la connaissance de chaque expert dans une base de connaissance et d'agréger ensuite ces connaissances individuelles pour définir la connaissance globale.

Très peu de travaux concernent la rationalité de la fusion de bases de connaissance. Nous proposons dans cette partie une caractérisation logique de ces opérateurs. Dans ce chapitre nous donnons quelques notations et définitions propres à cette partie et rappelons les différentes propositions de la littérature.

5.1 Notations

Définition 76 *Un ensemble de connaissance Ψ est un multi-ensemble de bases de connaissance.*

Si $\Psi = \{\varphi_1, \dots, \varphi_n\}$, on note $\bigwedge \Psi$ la conjonction des bases de connaissance de Ψ , i.e.

$\bigwedge \Psi = \varphi_1 \wedge \dots \wedge \varphi_n$. Et, si φ est une base de connaissance, on note $\Psi \wedge \varphi$ la base de connaissance $\bigwedge \Psi \wedge \varphi$. De même, on note $\bigvee \Psi$ la disjonction des bases de connaissance de Ψ , i.e. $\bigvee \Psi = \varphi_1 \vee \dots \vee \varphi_n$.

L'union sur les multi-ensembles sera notée $\sqcup$.

Remarque 4 Une base de connaissance inconsistante n'apportant pas d'information pour la fusion, nous supposons que toutes les bases de connaissance d'un ensemble de connaissance sont consistantes.

On notera $\mathcal{B}$ l'ensemble des bases de connaissance consistantes et $\mathcal{S}$ l'ensemble des ensembles de connaissance.

Un ensemble de connaissance Ψ est *consistant* si et seulement si $\bigwedge \Psi$ est consistant. De même, si φ est une base de connaissance, on dira que Ψ est consistant avec φ si $\bigwedge \Psi \wedge \varphi$ est consistant. On appellera modèles d'un ensemble de connaissance, les modèles de la conjonction de ses bases de connaissance et on notera $Mod(\Psi) = Mod(\bigwedge \Psi)$ et $I \models \Psi$ pour $I \in Mod(\Psi)$.

Par abus, si φ est une base de connaissance et Ψ un ensemble de connaissance on notera $\Psi \sqcup \varphi$ au lieu de $\Psi \sqcup \{\varphi\}$ et $\varphi \sqcup \varphi'$ au lieu de $\{\varphi\} \sqcup \{\varphi'\}$.

Pour tout entier naturel non nul n , on notera Ψ^n l'ensemble de connaissance $\underbrace{\Psi \sqcup \dots \sqcup \Psi}_n$.

Définition 77 Soient Ψ et Ψ' , deux ensembles de connaissance. On dit que Ψ et Ψ' sont équivalents, noté $\Psi \leftrightarrow \Psi'$, si et seulement si il existe une bijection f de $\Psi = \{\varphi_1, \dots, \varphi_n\}$ vers $\Psi' = \{\varphi'_1, \dots, \varphi'_n\}$ telle que $f(\varphi) \leftrightarrow \varphi$.

Soient deux bases de connaissance φ et μ , un ensemble de connaissance Ψ , et un opérateur Δ . On définit la séquence $\langle \Delta_\mu^n(\Psi, \varphi) \rangle_{n \geq 1}$ comme suit :

$$\begin{aligned} \Delta_\mu^1(\Psi, \varphi) &= \Delta_\mu(\Psi \sqcup \varphi) \\ \text{et } \Delta_\mu^{n+1}(\Psi, \varphi) &= \Delta_\mu(\Delta_\mu^n(\Psi, \varphi) \sqcup \varphi) \end{aligned}$$

5.2 Autres travaux

Ce que nous allons présenter dans les prochains chapitres est une étude et une caractérisation logique des opérateurs de fusion de bases de connaissance.

Très peu de travaux portent sur l'étude de ces opérateurs et nous allons les détailler ici. Cela permettra ensuite de souligner les différences et les similitudes entre notre proposition et les approches existantes. Revesz [Rev93] fut le premier à proposer l'étude des opérateurs de fusion dans le cadre de la révision AGM. Il proposa donc une caractérisation logique et un théorème de représentation pour ces opérateurs. Mais, comme nous le verrons par la suite, cette caractérisation s'avère un peu trop restrictive.

5.2.1 Revesz

Peter Revesz définit dans [Rev93, Rev97] des opérateurs de model fitting (*adéquation sémantique*). Ces opérateurs sont très proches des opérateurs de fusion contrainte que nous introduirons au chapitre suivant.

Une différence importante entre son approche et la nôtre est la notion d'ensemble de connaissance et d'équivalence entre ces ensembles. Pour Revesz, un ensemble de connaissance est un ensemble de bases de connaissance et sa notion d'équivalence est la suivante :

Définition 78 Soient Ψ_1 et Ψ_2 deux ensembles de connaissance. $\Psi_1 \Rightarrow \Psi_2$ ssi $\forall \varphi_2 \in \Psi_2 \exists \varphi_1 \in \Psi_1$ t.q. $\varphi_2 \leftrightarrow \varphi_1$. $\Psi_1 \Leftrightarrow \Psi_2$ ssi $\Psi_1 \Rightarrow \Psi_2$ et $\Psi_2 \Rightarrow \Psi_1$.

Revesz considère des opérateurs $\triangleright$ qui, à un ensemble de connaissance Ψ et une base de connaissance μ , associent une base de connaissance $\Psi \triangleright \mu$. L'opérateur $\triangleright$ est un opérateur de *model fitting* s'il satisfait les propriétés suivantes :

- (M1) $\Psi \triangleright \mu$ implique μ .
- (M2) Si $\Psi \wedge \mu$ est consistant, alors $\Psi \triangleright \mu \leftrightarrow \Psi \wedge \mu$
- (M3) Si μ est consistant, alors $\Psi \triangleright \mu$ est consistant
- (M4) Si $\Psi_1 \Leftrightarrow \Psi_2$ et $\mu \leftrightarrow \phi$, alors $\Psi_1 \triangleright \mu \leftrightarrow \Psi_2 \triangleright \phi$
- (M5) $(\Psi \triangleright \mu) \wedge \phi$ implique $\Psi \triangleright (\mu \wedge \phi)$
- (M6) Si $(\Psi \triangleright \mu) \wedge \phi$ est consistant alors $\Psi \triangleright (\mu \wedge \phi)$ implique $(\Psi \triangleright \mu) \wedge \phi$
- (M7) $(\Psi_1 \triangleright \mu) \wedge (\Psi_2 \triangleright \mu)$ implique $(\Psi_1 \cup \Psi_2) \triangleright \mu$

Revesz donne également un théorème de représentation pour ces opérateurs.

Définition 79 Un assignement loyal est une fonction qui associe à chaque ensemble de connaissance Ψ un pré-ordre $\leq_\Psi$ sur les interprétations tel que :

1. Si $I \models \Psi$ et $J \models \Psi$, alors $I \simeq_\Psi J$
2. Si $I \models \Psi$ et $J \not\models \Psi$, alors $I <_\Psi J$
3. Si $\Psi_1 \Leftrightarrow \Psi_2$, alors $\leq_{\Psi_1} = \leq_{\Psi_2}$
4. Si $I \leq_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors $I \leq_{\Psi_1 \cup \Psi_2} J$

Théorème 68 Un opérateur $\triangleright$ vérifie les propriétés (M1)-(M7) si et seulement si il existe un assignement loyal qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$ tel que

$$\text{Mod}(\Psi \triangleright \mu) = \min(\text{Mod}(\mu), \leq_\Psi)$$

Revesz [Rev93] définit alors l'*arbitrage* entre 2 bases de connaissance par le model fitting par $\top$:

$$\varphi \triangle \varphi' = \{\varphi, \varphi'\} \triangleright \top$$

Ces opérateurs de model fitting (et d'arbitrage) sont des opérateurs égalitaristes, c'est-à-dire qu'ils tentent de minimiser l'insatisfaction individuelle.

Avec cette définition il n'est pas possible de définir des opérateurs utilitaristes, c'est-à-dire des opérateurs majoritaires. Revesz définit alors des opérateurs de model fitting pondéré qui permettent un tel comportement.

Définition 80 *Un ensemble de connaissance pondéré est une fonction de l'ensemble des ensembles d'interprétations vers les nombres réels positifs : $\tilde{\Psi} : 2^{\mathcal{W}} \rightarrow \mathbb{R}^+$*

Ces nombres réels dénotent d'après Revesz le degré d'importance de chaque ensemble de modèles pour l'ensemble de connaissance pondéré. Il présente cette notion comme une généralisation de la définition d'ensemble de connaissance puisqu'un ensemble de connaissance usuel peut-être traduit en un ensemble de connaissance pondéré en posant pour $\Psi = \{\varphi_1, \dots, \varphi_n\}$, $\tilde{\Psi}(M) = 1$ si $\exists i M = \text{Mod}(\varphi_i)$ et $\tilde{\Psi}(M) = 0$ sinon.

L'union entre ensembles de connaissance est alors une union pondérée, $\tilde{\Psi}_1 \uplus \tilde{\Psi}_2$ est égal à l'ensemble de connaissance pondéré $\tilde{\Psi}$ défini par $\tilde{\Psi}(M) = \tilde{\Psi}_1(M) + \tilde{\Psi}_2(M)$ pour chaque ensemble d'interprétations M .

On utilise une fonction *form* pour obtenir un ensemble de connaissance à partir d'un ensemble de connaissance pondéré : $\text{form}(\tilde{\Psi}) = \{\text{form}(M_i) : \tilde{\Psi}(M_i) \neq 0\}$.

Soit un opérateur $\triangleright$ qui, à un ensemble de connaissance pondéré $\tilde{\Psi}$ et à une base de connaissance μ , associe une base de connaissance $\tilde{\Psi} \triangleright \mu$. L'opérateur $\triangleright$ est un opérateur de *model fitting pondéré* s'il satisfait les propriétés suivantes :

- (W1) $\tilde{\Psi} \triangleright \mu$ implique μ .
- (W2) Si $\text{form}(\tilde{\Psi}) \wedge \mu$ est consistant, alors $\tilde{\Psi} \triangleright \mu \leftrightarrow \text{form}(\tilde{\Psi}) \wedge \mu$
- (W3) Si μ est consistant, alors $\tilde{\Psi} \triangleright \mu$ est consistant
- (W4) Si $\mu \leftrightarrow \phi$, alors $\tilde{\Psi} \triangleright \mu \leftrightarrow \tilde{\Psi} \triangleright \phi$
- (W5) $(\tilde{\Psi} \triangleright \mu) \wedge \phi$ implique $\tilde{\Psi} \triangleright (\mu \wedge \phi)$
- (W6) Si $(\tilde{\Psi} \triangleright \mu) \wedge \phi$ est consistant alors $\tilde{\Psi} \triangleright (\mu \wedge \phi)$ implique $(\tilde{\Psi} \triangleright \mu) \wedge \phi$
- (W7) $(\tilde{\Psi}_1 \triangleright \mu) \wedge (\tilde{\Psi}_2 \triangleright \mu)$ implique $(\tilde{\Psi}_1 \uplus \tilde{\Psi}_2) \triangleright \mu$
- (W8) Si $(\tilde{\Psi}_1 \triangleright \mu) \wedge (\tilde{\Psi}_2 \triangleright \mu)$ est consistant, alors $(\tilde{\Psi}_1 \uplus \tilde{\Psi}_2) \triangleright \mu$ implique $(\tilde{\Psi}_1 \triangleright \mu) \wedge (\tilde{\Psi}_2 \triangleright \mu)$

Ces postulats sont principalement une généralisation des postulats de model fitting pour des bases pondérées. Seul le postulat (W8) est nouveau et renforce les conditions pour l'agrégation. Néanmoins, on peut remarquer qu'aucun postulat ne parle de l'importance des pondérations.

Il existe également un théorème de représentation pour ces opérateurs :

Définition 81 *Un assignement loyal pondéré est une fonction qui associe à chaque ensemble de connaissance pondéré $\tilde{\Psi}$ un pré-ordre $\leq_{\tilde{\Psi}}$ sur les interprétations tel que :*

1. Si $I \models \tilde{\Psi}$ et $J \models \tilde{\Psi}$, alors $I \simeq_{\tilde{\Psi}} J$
2. Si $I \models \tilde{\Psi}$ et $J \not\models \tilde{\Psi}$, alors $I <_{\tilde{\Psi}} J$
3. Si $\tilde{\Psi}_1 \Leftrightarrow \tilde{\Psi}_2$, alors $\leq_{\tilde{\Psi}_1} = \leq_{\tilde{\Psi}_2}$
4. Si $I \leq_{\tilde{\Psi}_1} J$ et $I \leq_{\tilde{\Psi}_2} J$, alors $I \leq_{\tilde{\Psi}_1 \uplus \tilde{\Psi}_2} J$
5. Si $I <_{\tilde{\Psi}_1} J$ et $I \leq_{\tilde{\Psi}_2} J$, alors $I <_{\tilde{\Psi}_1 \uplus \tilde{\Psi}_2} J$

Théorème 69 *Un opérateur vérifie les propriétés (W1)-(W8) si et seulement si il existe un assignement loyal pondéré qui associe à chaque ensemble de connaissance pondéré $\tilde{\Psi}$ un pré-ordre total $\leq_{\tilde{\Psi}}$ tel que*

$$Mod(\tilde{\Psi} \triangleright \mu) = \min(Mod(\mu), \leq_{\tilde{\Psi}})$$

Le principal problème lié à cette définition est que, nulle part dans la caractérisation logique, on ne parle de ces poids attachés aux ensembles d'interprétations. C'est-à-dire que la définition d'un opérateur de model fitting pondéré permet d'introduire une mesure de confiance pour chacune des bases de connaissance mais de tels opérateurs n'utilisent pas forcément cette information sur l'importance relative des informations.

Un autre problème d'interprétation est que Revesz semble faire une identification entre pondération et majorité [Rev97].

“Note that in the case of weighted model fitting the instructor tries to satisfy the majority of the class, instead of trying to satisfy each member to the best degree possible”

Mais, comme nous le verrons section 6.1, la notion d'opérateurs majoritaires n'est pas liée à l'existence d'une pondération.

5.2.2 Liberatore et Schaerf

Liberatore et Schaerf [LS95, LS98] définissent des opérateurs de fusion qu'ils nomment indifféremment opérateurs d'arbitrage ou opérateurs de révision commutative. Nous appellerons ces opérateurs opérateurs de révision commutative, d'une part parce que le terme “opérateur d'arbitrage” à un sens particulier dans notre approche et, d'autre part, ce nom illustre bien l'idée sous-jacente de ces opérateurs.

Ils considèrent des opérateurs $\diamond$ qui associent à un couple de bases de connaissance une nouvelle base de connaissance résultat de la fusion des deux bases initiales. Les opérateurs de *révision commutative* sont les opérateurs qui vérifient les propriétés suivantes :

- (LS1) $\varphi \diamond \mu \leftrightarrow \mu \diamond \varphi$
- (LS2) $\varphi \wedge \mu$ implique $\varphi \diamond \mu$
- (LS3) Si $\varphi \wedge \mu$ est consistant, alors $\varphi \diamond \mu$ implique $\varphi \wedge \mu$
- (LS4) $\varphi \diamond \mu$ est inconsistant ssi φ et μ sont inconsistants
- (LS5) Si $\varphi_1 \leftrightarrow \varphi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\varphi_1 \diamond \mu_1 \leftrightarrow \varphi_2 \diamond \mu_2$
- (LS6) $\varphi \diamond (\mu \vee \theta) = \begin{cases} \varphi \diamond \mu & \text{ou} \\ \varphi \diamond \theta & \text{ou} \\ (\varphi \diamond \mu) \vee (\varphi \diamond \theta) \end{cases}$
- (LS7) $\varphi \diamond \mu$ implique $\varphi \vee \mu$
- (LS8) If φ est consistant alors $\varphi \wedge (\varphi \diamond \mu)$ est consistant

Liberatore et Schaerf appellent postulats de base les postulats (LS1)-(LS5) et (LS6)-(LS8) les postulats spécifiques de la révision commutative.

Ces postulats sont une généralisation des postulats de la révision lorsque l'on contraint l'opérateur à être commutatif. D'après les auteurs, le but est d'obtenir des opérateurs de fusion dont le résultat est compris entre la disjonction et la conjonction des deux bases.

Ces opérateurs portent bien leur nom de révision commutative puisque ces opérateurs peuvent s'écrire :

$$\varphi \diamond \mu = (\varphi \circ \mu) \vee (\mu \circ \varphi)$$

où $\circ$ est un opérateur de révision AGM.

Cette idée est l'idée sous jacente dans la construction des opérateurs de Liberatore et Schaerf. Nous donnons section 8.3 un théorème de représentation pour ces opérateurs qui caractérisera la famille d'opérateurs de révision correspondant aux opérateurs de révision commutative via cette égalité.

Les postulats (LS7) et (LS8) impliquent la propriété suivante [LS98] :

(LS78) Si φ et μ sont deux formules complètes, alors $\varphi \diamond \mu = \varphi \vee \mu$

Théorème 70 *Un opérateur qui satisfait (LS1)-(LS6) satisfait (LS78) si et seulement si il satisfait (LS7) et (LS8).*

Cette propriété paraît très contraignante. Nous avons argumenté dans [KP98] que la propriété (LS7) était discutable en tant que propriété générale pour un opérateur de fusion :

Exemple 15 *Supposons que je veuille jouer en bourse. Je demande l'avis de deux experts financiers à propos de quatre actions A, B, C, D . On note 1 si l'action monte et 0 si elle descend (on suppose que sa valeur ne peut rester stable). Ces deux personnes sont aussi compétentes l'une que l'autre et j'ai donc une confiance égale en ces deux avis. Le premier expert m'annonce que les quatre actions vont monter $\varphi_1 = \{(1,1,1,1)\}$, le second que toutes les actions vont descendre $\varphi_2 = \{(0,0,0,0)\}$. Les opérateurs de Liberatore et Schaerf donnent alors le résultat suivant à partir de ces deux opinions : $R = \{(0,0,0,0), (1,1,1,1)\}$. Cela signifie donc que soit l'expert φ_1 a totalement faux, soit c'est φ_2 qui s'est entièrement trompé. Mais intuitivement, si les deux experts sont aussi fiables l'un que l'autre, il n'y a aucune raison que l'un se soit trompé plus que l'autre : il doivent tout deux être à la même "distance" de la vérité. Ils ont alors certainement faux à propos de deux actions et le résultat doit être quelque chose comme $R' = \{(0,0,1,1), (0,1,0,1), (0,1,1,0), (1,0,0,1), (1,0,1,0), (1,1,0,0)\}$. C'est-à-dire que deux actions vont monter et deux vont baisser mais on ne sait pas lesquelles.*

Cette propriété (LS7) n'est donc pas anodine et restreint les opérateurs de Liberatore et Schaerf à des applications où le résultat de la fusion doit être une alternative proposée par un des intervenants.

Liberatore et Schaerf proposent un théorème de représentation pour ces opérateurs en termes de familles de pré-ordres sur les ensembles d'interprétations.

Une fonction $F \mapsto \leq_F$ qui, à un ensemble d'interprétations, associe un pré-ordre sur les ensembles d'interprétation est appelée un *assignement commutatif* si elle satisfait les propriétés suivantes :

- (L1) Si $A \leq_F B$ et $B \leq_F C$, alors $A \leq_F C$
- (L2) Si $A \subseteq B$, alors $B \leq_F A$
- (L3) $A \leq_F A \cup B$ ou $B \leq_F A \cup B$
- (L4) $B \leq_F C$ pour chaque C ssi $F \cap B \neq \emptyset$

$$(L5) \quad A \leq_{C \cup D} B \Leftrightarrow \begin{cases} C \leq_{A \cup B} D \text{ et } A \leq_C B & \text{ou} \\ D \leq_{A \cup B} C \text{ et } A \leq_D B \end{cases}$$

Soit A , un ensemble d'interprétations, $\hat{A}$ dénote l'ensemble $\{\{I\} | I \in A\}$. Leur théorème de représentation est alors le suivant :

Théorème 71 *Un opérateur $\diamond$ est un opérateur de révision commutative si et seulement si il existe un assignement commutatif qui associe à chaque formule μ un pré-ordre total sur les ensembles d'interprétations tel que*

$$Mod(\varphi \diamond \mu) = \{I | \{I\} \in \min(\widehat{Mod(\varphi)}, \leq_{Mod(\mu)}) \cup \min(\widehat{Mod(\mu)}, \leq_{Mod(\varphi)})\}$$

Ce théorème diffère des théorèmes de représentation usuels car les pré-ordres portent sur des ensembles d'interprétations et non pas sur des interprétations.

Néanmoins, une analyse de la preuve de ce théorème fait apparaître qu'un opérateur de révision commutative est défini à partir d'un opérateur de révision AGM ayant certaines propriétés. Ce qui permet de définir un théorème de représentation pour les opérateurs de révision commutative en termes de familles de pré-ordres sur les interprétations. Ce résultat sera donné section 8.3 lorsque l'on examinera les rapports entre les opérateurs de Liberatore and Schaerf et les opérateurs de fusion contrainte.

Un problème de ces opérateurs est qu'ils ne permettent de fusionner que 2 bases de connaissance. En effet, les opérateurs de Liberatore et Schaerf étant des opérateurs binaires, ils devraient donc être associatifs pour pouvoir opérer sur des ensembles de connaissance de taille quelconque. Malheureusement, si l'on exige l'associativité, on tombe sur des résultats de trivialité [LS98]. Nous verrons section 8.3 une solution pour généraliser ces opérateurs à plus de deux bases de connaissance.

5.2.3 Lin et Mendelzon

Lin et Mendelzon ont défini des opérateurs qu'ils nomment opérateurs de fusion majoritaire [LM99, LM98, Lin95].

Soit un opérateur $\blacktriangle$ qui, à un ensemble de connaissance Ψ , associe une base de connaissance $\blacktriangle(\Psi)$. L'opérateur $\blacktriangle$ est un opérateur de *fusion majoritaire* s'il vérifie les propriétés suivantes :

(LM1) $\blacktriangle(\Psi)$ est consistant

(LM2) Si $\bigwedge \Psi$ est consistant, alors $\blacktriangle(\Psi) = \bigwedge \Psi$

(LM3) Si $\Psi \leftrightarrow \Psi'$, alors $\blacktriangle(\Psi) \leftrightarrow \blacktriangle(\Psi')$

(LM4) Soit un littéral l , si $|\{\varphi_i \in \Psi : \varphi_i \models l\}| > |\{\varphi_i \in \Psi : \varphi_i \models \neg l\}| + |\{\varphi_i \in \Psi : \varphi_i \triangleright \neg l\}|$, alors $\blacktriangle(\Psi)$ implique l .

Où $\varphi_i \triangleright \neg l$ signifie que la base de connaissance φ_i *supporte implicitement* (*partially support*) $\neg l$, c'est-à-dire qu'il existe β , qui ne mentionne aucun des atomes de $\neg l$, tel que $\varphi_i \models \neg l \vee \beta$ mais $\varphi_i \not\models \neg l$ et $\varphi_i \not\models \beta$. Alors que $\varphi_i \models \neg l$ est appelé un *support explicite* (*full support*) de $\neg l$.

Le postulat (LM4) exprime simplement l'idée d'un vote pour ou contre l , c'est-à-dire que l est "élu" s'il y a plus de supports explicites pour l que de supports explicites et implicites pour $\neg l$. Il est nécessaire de prendre en compte les supports implicites pour éviter des réponses

incohérentes (il est par exemple sinon possible d'avoir l et $\neg l$ simultanément).

L'exemple d'opérateur que donnent Lin et Mendelzon est un opérateur bien connu (cf section 7.4).

L'originalité de cette définition réside dans le postulat (LM4) qui tente de capturer le comportement majoritaire en considérant la fusion de bases de connaissance comme une élection des différentes variables propositionnelles. Lin et Mendelzon proposent également une méthode syntaxique de calcul du résultat de la fusion lorsque les bases de connaissance sont exprimées en forme normale disjonctive (DNF).

Chapitre 6

Fusion contrainte

Ce qui est contraire est utile et c'est de ce qui est en lutte que naît la plus belle harmonie ; tout se fait par discorde.
Héraclite

Dans ce chapitre nous allons définir un cadre général pour la fusion qui étendra les idées de Revesz, de Liberatore et Schaerf et de Lin et Mendelzon.

Nous donnons une caractérisation logique des opérateurs de fusion contrainte. Ces opérateurs réalisent la fusion d'un ensemble de connaissance lorsque le résultat de la fusion doit obéir à un ensemble de contraintes d'intégrité.

Nous définissons également deux sous-classes d'opérateurs de fusion contrainte : les opérateurs majoritaires et les opérateurs d'arbitrage. Les opérateurs majoritaires tiennent compte de la majorité pour réaliser l'arbitrage, alors que les opérateurs d'arbitrage ont un comportement plus égalitaire, tentant de satisfaire au mieux chacune des bases de connaissance.

Nous examinons également plusieurs autres propriétés logiques que l'on pourrait trouver intéressantes pour la fusion, comme l'associativité ou la monotonie par exemple.

Enfin nous donnons des théorèmes de représentation pour les différents types d'opérateurs définis. Ces théorèmes disent qu'un opérateur de fusion contrainte correspond à une famille de pré-ordre sur les interprétations.

6.1 Caractérisation logique

Un opérateur de fusion Δ est une fonction qui associe, à un ensemble de connaissance et à une base de connaissance, une base de connaissance, i.e.

$$\Delta : \begin{array}{l} \mathcal{S} \times \mathcal{B} \quad \mapsto \mathcal{B} \\ (\Psi, \mu) \longrightarrow \varphi \end{array}$$

où Ψ est l'ensemble de bases de connaissance à fusionner, μ représente les contraintes d'intégrité de la fusion et φ est la base de connaissance résultat de la fusion.

Dans la suite nous noterons $\Delta_\mu(\Psi)$ à la place de $\Delta(\Psi, \mu)$.

Dans ce cadre, le résultat de la fusion doit obéir à un ensemble de contraintes d'intégrité. On rassemble ces contraintes dans une base de connaissance μ . Ces contraintes portent sur le résultat de la fusion et non pas sur les bases de connaissance intervenant dans le processus. Ainsi, il est possible que certaines des bases de connaissance de l'ensemble de connaissance à fusionner n'obéissent pas à ces contraintes d'intégrité. Nous exigeons que les contraintes d'intégrité soient vraies dans la base de connaissance résultat de la fusion et pas simplement constantes avec le résultat, contrairement à ce qui est habituellement demandé dans le cadre des bases de données.

Nous définissons à présent les opérateurs de fusion contrainte [KP99b, KP99a]:

Définition 82 Δ est un opérateur de fusion contrainte si et seulement si il satisfait les propriétés suivantes :

- (IC0) $\Delta_\mu(\Psi) \vdash \mu$
- (IC1) Si μ est consistant, alors $\Delta_\mu(\Psi)$ est consistant
- (IC2) Si Ψ est consistant avec μ , alors $\Delta_\mu(\Psi) = \bigwedge \Psi \wedge \mu$
- (IC3) Si $\Psi_1 \leftrightarrow \Psi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\Delta_{\mu_1}(\Psi_1) \leftrightarrow \Delta_{\mu_2}(\Psi_2)$
- (IC4) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\Delta_\mu(\varphi \sqcup \varphi') \wedge \varphi \not\vdash \perp \Rightarrow \Delta_\mu(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$
- (IC5) $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2) \vdash \Delta_\mu(\Psi_1 \sqcup \Psi_2)$
- (IC6) Si $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant, alors $\Delta_\mu(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$
- (IC7) $\Delta_{\mu_1}(\Psi) \wedge \mu_2 \vdash \Delta_{\mu_1 \wedge \mu_2}(\Psi)$
- (IC8) Si $\Delta_{\mu_1}(\Psi) \wedge \mu_2$ est consistant, alors $\Delta_{\mu_1 \wedge \mu_2}(\Psi) \vdash \Delta_{\mu_1}(\Psi) \wedge \mu_2$

La signification intuitive de ces propriétés est la suivante: (IC0) assure que le résultat de la fusion satisfait les contraintes d'intégrité. (IC1) dit que si les contraintes d'intégrité sont constantes alors le résultat de la fusion sera consistant. (IC2) demande que, lorsque c'est possible, le résultat de la fusion soit simplement la conjonction des bases de connaissance et des contraintes d'intégrité. (IC3) est le principe d'indépendance de syntaxe, c'est-à-dire que lorsque deux ensembles de connaissance sont équivalents et deux bases de connaissance sont équivalentes alors les bases de connaissance résultant des deux fusions seront équivalentes. (IC4) est la propriété d'équité. Elle assure que lorsque l'on fusionne deux bases de connaissance, l'opérateur ne peut pas donner de préférences à l'une d'elles. (IC5) exprime l'idée suivante: si un groupe Ψ_1 se met d'accord sur un ensemble d'alternatives qui contient A , et si un autre groupe Ψ_2 se met d'accord sur un autre ensemble d'alternatives qui contient également A , alors si l'on joint les deux groupes A fera encore partie des alternatives acceptables. Et (IC5) et (IC6) expriment le fait que, dès que l'on peut trouver deux sous-groupes qui s'accordent sur au moins une alternative, alors le résultat de la fusion sera exactement l'ensemble des alternatives sur lesquelles ces deux groupes s'accordent. (IC7) et (IC8) sont une généralisation directe des postulats (R5) et (R6) de la révision. Ils expriment des conditions sur les conjonctions de contraintes d'intégrité et s'assurent de ce fait que la notion de *proximité* est bien fondée.

(IC6) peut paraître un peu fort. Nous pensons que ce n'est pas le cas, mais il existe des opérateurs ayant un comportement convenable qui satisfont tous les postulats sauf (IC6). Il

peut donc être intéressant d'affaiblir (IC6) pour caractériser également les opérateurs n'ayant pas toutes les qualités requises pour être des opérateurs de fusion contrainte.

(IC6') Si $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant, alors $\Delta_\mu(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\mu(\Psi_1) \vee \Delta_\mu(\Psi_2)$

(IC6') demande que s'il existe une alternative commune lors de la fusion de deux sous-groupes, alors les alternatives résultats de la concertation générale sont incluses dans les alternatives convenables pour l'un des sous-groupes.

Nous appellerons les opérateurs vérifiant les propriétés (IC1)-(IC5), (IC6'), (IC7)-(IC8) des opérateurs de *quasi-fusion*.

En plus de ces propriétés de base, on peut également envisager les propriétés suivantes :

Tout d'abord on peut exiger une propriété d'*itération*. Cette propriété donne un caractère "topologique" aux opérateurs de fusion.

(IC_{it}) Si $\varphi \vdash \mu$ alors $\exists n \Delta_\mu^n(\Psi, \varphi) \vdash \varphi$ (itération)

L'idée intuitive de cette propriété est que, puisque les opérateurs de fusion donnent, en quelque sorte, la connaissance moyenne d'un ensemble de connaissance, si l'on itère le processus de fusion avec une même base de connaissance, on doit atteindre cette base après un certain nombre d'itérations.

Nous allons à présent définir deux sous classes d'opérateurs de fusion, les opérateurs de fusion majoritaires et les opérateurs d'arbitrage.

Un opérateur de *fusion majoritaire* est un opérateur de fusion contrainte qui satisfait la propriété suivante :

(Maj) $\exists n \Delta_\mu(\Psi_1 \sqcup \Psi_2^n) \vdash \Delta_\mu(\Psi_2)$ (majorité)

Ce postulat exprime, entres autres, le fait que si une opinion a une large audience, ce sera alors l'opinion du groupe. Les opérateurs de fusion majoritaire tentent donc de satisfaire au mieux le groupe dans son ensemble. D'un autre côté, les opérateurs d'arbitrage tentent de satisfaire chacun des éléments du groupe pris individuellement du mieux possible. Un opérateur d'*arbitrage* est un opérateur de fusion contrainte qui satisfait la propriété suivante :

(Arb)
$$\left. \begin{array}{l} \Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2) \\ \Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2) \\ \mu_1 \not\vdash \mu_2 \\ \mu_2 \not\vdash \mu_1 \end{array} \right\} \Rightarrow \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow \Delta_{\mu_1}(\varphi_1) \quad \text{(arbitrage)}$$

Ce postulat dit que si un ensemble d'alternatives préférées sous un ensemble de contraintes d'intégrité μ_1 pour une base de connaissance φ_1 correspond à l'ensemble des alternatives préférées par la base φ_2 sous les contraintes μ_2 , et si les alternatives qui n'appartiennent qu'à un des deux ensembles de contraintes d'intégrité sont toutes aussi crédibles pour le groupe ($\varphi_1 \sqcup \varphi_2$), alors les alternatives préférées pour le groupe parmi la disjonction des deux ensembles de contraintes seront celles préférées par chacune des bases sous leur contraintes

respectives.

Cette propriété est bien plus naturelle quand elle est exprimée en terme de conditions sur les pré-ordres (voir condition 8 des *assignements synchrétiques justes*, définition 83), montrant que ce sont les choix possibles médians qui sont préférés. Illustrons cela sur un exemple :

Exemple 16 *Tom et David ont raté le match de football d'hier entre les rouges et les jaunes. Ils ne connaissent donc pas le résultat du match. Tom a entendu ce matin à la radio que les rouges ont fait un très bon match. Il pense donc qu'une victoire des rouges est plus crédible qu'un match nul, et qu'un match nul est plus crédible qu'une victoire des jaunes. Un ami a dit à David qu'après ce match, les jaunes ont à présent toutes les chances de remporter le championnat. Il déduit de cette information que les jaunes ont très certainement gagné le match, ou sinon, au moins obtenu un match nul. En confrontant leurs points de vue, Tom et David se mettent d'accord sur le fait que les deux équipes sont de la même force et qu'elles avaient donc les mêmes chances de remporter le match. Ce que demande la propriété d'arbitrage est qu'avec ces informations Tom et David doivent se mettre d'accord sur le fait qu'un match nul est le résultat le plus crédible.*

Une autre propriété, opposée à la propriété de majorité, que nous pouvons également mentionner est celle d'indépendance de la majorité :

$$(MI) \quad \forall n \quad \Delta_\mu(\Psi_1 \sqcup \Psi_2^n) \leftrightarrow \Delta_\mu(\Psi_1 \sqcup \Psi_2) \quad (\text{indépendance de la majorité})$$

Cette propriété très forte impose que le résultat de la fusion soit totalement indépendant de la popularité des opinions, mais prenne simplement en compte les différentes opinions exprimées.

Une conséquence de cette propriété est que pour les opérateurs satisfaisant (MI) les ensembles de connaissance peuvent être considérés comme de simples ensembles, c'est-à-dire que l'on considère pour ces opérateurs qu'il n'y a pas de répétitions dans les multi-ensembles.

Mais cette propriété n'est malheureusement pas compatible avec celles d'un opérateur de fusion contrainte :

Théorème 72 *Il n'y a pas d'opérateur de fusion contrainte satisfaisant (MI).*

Preuve : Cette preuve est due à P. Liberatore (communication personnelle): Soient $\Psi_1 = \{\varphi, \neg\varphi\}$ et $\Psi_2 = \{\varphi\}$ deux ensembles de connaissance. Par (MI) on trouve que $\Delta_\top(\Psi_1 \sqcup \Psi_2) = \Delta_\top(\Psi_1)$. Par (IC4) nous avons également $\Delta_\top(\Psi_1) \not\models \varphi$ et $\Delta_\top(\Psi_1) \not\models \neg\varphi$. De plus par (IC2) on déduit $\Delta_\top(\Psi_2) = \varphi$. Alors $\Delta_\top(\Psi_1) \wedge \Delta_\top(\Psi_2)$ est consistant et par (IC6) on obtient $\Delta_\top(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\top(\Psi_1) \wedge \Delta_\top(\Psi_2)$, i.e. $\Delta_\top(\Psi_1) \vdash \Delta_\top(\Psi_1) \wedge \varphi$. Alors $\Delta_\top(\Psi_1) \vdash \varphi$, ce qui contredit (IC4). □

Néanmoins, une forme affaiblie de cette dernière propriété, appelée indépendance faible de la majorité, est compatible avec les opérateurs de fusion contrainte :

$$(WMI) \quad \forall \varphi' \exists \varphi \varphi' \not\models \varphi \quad \forall n \quad \Delta_\mu(\varphi' \sqcup \varphi^n) = \Delta_\mu(\varphi' \sqcup \varphi) \quad (\text{indépendance faible de la majorité})$$

Cette forme affaiblie de l'indépendance de la majorité demande simplement l'existence de cas où le résultat de la fusion est indépendant de la popularité des avis.

Dans la suite de cette section nous donnons quelques résultats sur les rapports entre les postulats.

Théorème 73 *Si un opérateur satisfait (IC1) et (IC2), alors il ne peut vérifier à la fois (MI) et (Maj).*

Preuve : De (MI) et (Maj) pour tout Ψ on a $\forall \varphi \Delta_{\top}(\Psi \sqcup \varphi) \leftrightarrow \Delta_{\top}(\Psi \sqcup \varphi^n) \vdash \Delta_{\top}(\varphi)$ et de (IC2) on déduit que

$$\forall \varphi \Delta_{\top}(\Psi \sqcup \varphi) \vdash \varphi \quad (*)$$

Prenons φ' tel que $\varphi \wedge \varphi' \vdash \perp$. Alors avec $\Psi = \varphi'$, par (*), on a que $\Delta_{\top}(\varphi' \sqcup \varphi) \vdash \varphi$. De même $\Delta_{\top}(\varphi \sqcup \varphi') \vdash \varphi'$ donc $\Delta_{\top}(\varphi \sqcup \varphi') \vdash \varphi \wedge \varphi'$. D'où $\Delta_{\top}(\varphi \sqcup \varphi') \vdash \perp$ ce qui contredit (IC1). $\square$

Théorème 74 *Si un opérateur satisfait (IC1), (IC2) et (IC4), alors il ne peut satisfaire à la fois (WMI) et (Maj).*

Preuve : A partir de (WMI) et (Maj) on déduit facilement $\forall \varphi' \exists \varphi \varphi' \not\vdash \varphi \Delta_{\top}(\varphi' \sqcup \varphi) \vdash \varphi$. Ce qui par (IC4) donne $\forall \varphi' \exists \varphi \varphi' \not\vdash \varphi \Delta_{\top}(\varphi' \sqcup \varphi) \vdash \varphi \wedge \varphi'$. Or si $\varphi' = \varphi_I = \text{form}(I)$, alors $\varphi_I \not\vdash \varphi$ peut se réécrire $\varphi_I \wedge \varphi \vdash \perp$ et donc $\Delta_{\top}(\varphi_I \sqcup \varphi) \vdash \perp$. Ce qui contredit (IC1). $\square$

Intéressons nous à présent aux propriétés d'associativité et de monotonie. Nous pensons que ces propriétés ne sont pas souhaitables pour des opérateurs de fusion et nous montrons que les opérateurs de fusion contrainte ne les satisfont pas.

Donnons tout d'abord la définition formelle de ces propriétés.

$$(\text{Ass}) \quad \Delta_{\mu}(\Psi_1 \sqcup \Delta_{\mu}(\Psi_2)) = \Delta_{\mu}(\Psi_1 \sqcup \Psi_2) \quad (\text{associativité})$$

L'associativité est une propriété enviable car elle permet de définir des opérateurs autorisant des sous-fusions dans l'ensemble de connaissance, ce qui peut-être très intéressant du point de vue calculatoire, en particulier dans des systèmes distribués. Néanmoins, cette propriété va à l'encontre de l'idée de *connaissance moyenne* que représente le résultat d'une fusion.

$$(\text{Mon}) \text{ Si } \varphi_1 \vdash \varphi'_1, \dots, \varphi_n \vdash \varphi'_n \text{ alors } \Delta_{\mu}(\varphi_1 \sqcup \dots \sqcup \varphi_n) \vdash \Delta_{\mu}(\varphi'_1 \sqcup \dots \sqcup \varphi'_n) \quad (\text{monotonie})$$

La propriété de monotonie exprime le fait que si un ensemble de connaissance est plus "fort" qu'un autre, alors le résultat de la fusion du premier doit être une base de connaissance logiquement plus forte que le résultat de la fusion du second.

Théorème 75 *Si un opérateur satisfait (IC2) et (IC4), alors il ne satisfait pas (Mon).*

Preuve : Soient I et J , deux interprétations distinctes. Soient $\varphi_1 = \varphi'_1 = \text{form}(I)$, $\varphi_2 = \text{form}(J)$, et $\varphi'_2 = \text{form}(I, J)$, nous avons donc $\varphi_1 \vdash \varphi'_1$ et $\varphi_2 \vdash \varphi'_2$. De (IC2) on déduit $\Delta_{\top}(\varphi'_1 \sqcup \varphi'_2) = \text{form}(I)$ et de (IC4) on a $\Delta_{\top}(\varphi_1 \sqcup \varphi_2) \neq \text{form}(I)$. On a donc $\Delta_{\top}(\varphi_1 \sqcup \varphi_2) \not\vdash \Delta_{\top}(\varphi'_1 \sqcup \varphi'_2)$.

□

Il est donc clair que la propriété de monotonie n'est pas vérifiée par les opérateurs de fusion contrainte. Ce n'est pas exactement la même chose pour l'associativité, nous montrons qu'elle n'est pas satisfaite par les opérateurs majoritaires et qu'elle n'est pas compatible avec la propriété d'itération.

Théorème 76 *Si un opérateur satisfait (IC1), (IC2), (IC4) et (Maj), alors il ne vérifie pas (Ass).*

Preuve : Prenons deux formules complètes distinctes φ_I et φ_J . Par (Maj), (IC1) et (IC2) on sait que $\exists n \Delta_{\top}(\varphi_I \sqcup \varphi_J^n) = \varphi_J$. Par (Ass) on a $\Delta_{\top}(\varphi_I \sqcup \varphi_J^n) = \Delta_{\top}(\varphi_I \sqcup \Delta_{\top}(\varphi_J^n))$. Mais par (IC2) on a $\Delta_{\top}(\varphi_J^n) = \varphi_J$, on obtient donc $\Delta_{\top}(\varphi_I \sqcup \varphi_J) = \varphi_J$. Ce qui contredit (IC4). □

Théorème 77 *Si un opérateur vérifie (IC2) et (IC4), alors il ne peut satisfaire à la fois (IC_{it}) et (Ass).*

Preuve : (IC_{it}) et (IC2) donnent $\exists n \Delta_{\top}^n(\Psi, \varphi) \vdash \varphi$, mais d'après (Ass) on déduit que $\Delta_{\top}^n(\Psi, \varphi) = \Delta_{\top}(\Psi \sqcup \varphi^n) = \Delta_{\top}(\Psi \sqcup \Delta_{\top}(\varphi^n))$ et par (IC2) on a $\Delta_{\top}(\Psi \sqcup \Delta_{\top}(\varphi^n)) = \Delta_{\top}(\Psi \sqcup \varphi)$. On a donc $\Delta_{\top}(\Psi \sqcup \varphi) \vdash \varphi$, ce qui, en prenant $\Psi = \varphi'$ avec $\varphi' \wedge \varphi \vdash \perp$, contredit (IC4). □

Remarquons que même si l'on considère la version faible de la propriété d'associativité suivante :

$$(WA) \quad \Delta_{\mu}(\Psi_1 \sqcup \Psi_2) \vdash \Delta_{\mu}(\Psi_1 \sqcup \Delta_{\mu}(\Psi_2)) \quad (\text{associativité faible})$$

on n'échappe pas à ces résultats négatifs.

Nous terminons cette section en donnant deux propriétés supplémentaires des opérateurs de fusion contrainte.

Théorème 78 $\Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2) \vdash \Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2)$

Preuve : Soit $I \models \Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$. Supposons que $I \not\models \Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2)$. Comme $I \models \Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$, alors par (IC8) $\forall J \models \mu_1 \wedge \mu_2 \quad I \models \Delta_{\varphi_{\{I, J\}}}(\Psi_1)$ et $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi_2)$, alors par (IC5) $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi_1 \sqcup \Psi_2)$. Mais si $I \not\models \Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2)$, alors par (IC1) il existe $J \models \Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2)$ et par (IC7) et (IC8) on a que $\Delta_{\varphi_{\{I, J\}}}(\Psi_1 \sqcup \Psi_2) = \varphi_{\{J\}}$. Contradiction. □

Théorème 79 *Si $\Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$ est consistant alors $\Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2) \vdash \Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$*

Preuve : Supposons que $\Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$ est consistant et que $I \models \Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2)$. Supposons que $I \not\models \Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$. C'est-à-dire, sans perte de généralité, supposons que $I \not\models \Delta_{\mu_1}(\Psi_1)$. Comme $\Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$ est consistant alors il existe $J \models \Delta_{\mu_1}(\Psi_1) \wedge \Delta_{\mu_2}(\Psi_2)$. Donc par (IC7) et (IC8) $\Delta_{\varphi_{\{I, J\}}}(\Psi_1) \wedge \Delta_{\varphi_{\{I, J\}}}(\Psi_2)$ est consistant, et par (IC5) et (IC6) on

conclut que $\Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2) = \varphi_{\{J\}}$. Mais d'après l'hypothèse on a $I \models \Delta_{\mu_1 \wedge \mu_2}(\Psi_1 \sqcup \Psi_2)$ alors $\forall J \models \mu_1 \wedge \mu_2$ par (IC7) on a que $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2)$. Contradiction. $\square$

6.2 Discussion sur les postulats

La plupart des propriétés logiques que doivent satisfaire les opérateurs de fusion contrainte sont une généralisation assez naturelle des propriétés existantes pour la révision. Leur justification est donc assez aisée. En revanche, quelques propriétés sont nouvelles et assez différentes de celles déjà proposées dans la littérature. Nous allons tenter de justifier un peu plus ici le bien fondé de ces propriétés. Nous donnerons également d'autres propriétés envisageables et motiverons nos choix.

6.2.1 Équité

Ce que l'on entend par équité pour la fusion est qu'un opérateur ne doit pas donner une préférence arbitraire à l'un des intervenants du groupe. Cette notion est interprétée par le postulat suivant dans notre caractérisation :

(IC4) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\Delta_\mu(\varphi \sqcup \varphi') \wedge \varphi \not\vdash \perp \Rightarrow \Delta_\mu(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$ (équité)

Mais on peut imaginer d'autres postulats d'"équité". Nous en donnons quelques uns ci-dessous :

(Eq1) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \not\vdash \varphi \Rightarrow \Delta_\mu(\varphi' \sqcup \varphi) \neq \varphi'$

(Eq2) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \vdash \varphi \Leftrightarrow \Delta_\mu(\varphi' \sqcup \varphi) = \varphi'$

(Eq3) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \wedge \varphi \vdash \perp \Rightarrow \Delta_\mu(\varphi' \sqcup \varphi) \neq \varphi'$

(Eq4) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \wedge \varphi \vdash \perp \Rightarrow \Delta_\mu(\varphi' \sqcup \varphi) \not\vdash \varphi'$

(Eq5) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \wedge \varphi \not\vdash \perp \Leftrightarrow \Delta_\mu(\varphi' \sqcup \varphi) \vdash \varphi'$

Ces propriétés sont très proches les unes des autres. En fait, on a différentes implications qui sont données figure 6.1.

Les flèches $(A \rightarrow B)$ signifient que A implique B, et * signifie que l'implication est vrai si Δ satisfait (IC2).

Les preuves que (IC4) implique (Eq4) et que (Eq1) implique (Eq3) sont directes. La preuve que (Eq1) est équivalent à (Eq2) est similaire à celle de l'équivalence entre (Eq4) et (Eq5).

Théorème 80 *En présence d'(IC2), (Eq4) implique (Eq1).*

Preuve : Supposons que $\varphi' \not\vdash \varphi$ et considérons les deux cas suivants :

1. $\varphi' \wedge \varphi \not\vdash \perp$ donc par (IC2) on conclut que $\Delta_\mu(\varphi' \sqcup \varphi) = \varphi' \wedge \varphi$, or comme $\varphi' \not\vdash \varphi$ on a $\Delta_\mu(\varphi' \sqcup \varphi) \neq \varphi'$.

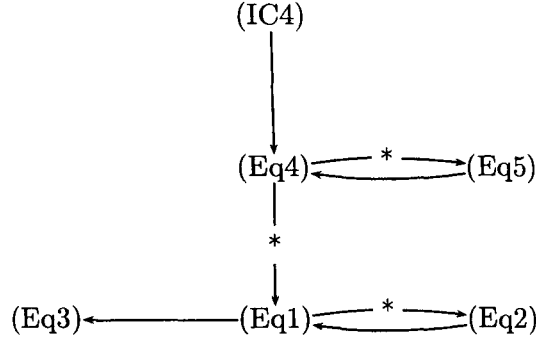


FIG. 6.1 – Les différentes expressions de la propriété d'équité

2. $\varphi' \wedge \varphi \vdash \perp$ donc par (Eq4) on a $\Delta_\mu(\varphi' \sqcup \varphi) \not\vdash \varphi'$, donc $\Delta_\mu(\varphi' \sqcup \varphi) \neq \varphi'$.

□

Théorème 81 *En présence d'(IC2), (Eq5) est équivalent à (Eq4).*

Preuve : (Eq5) implique directement (Eq4). (Eq4) implique (Eq5) car le sens *seulement si* de (Eq5) est donné par (IC2) et le sens *si* est la contraposée de (Eq4).

□

Nous avons donc choisi la propriété la plus forte pour exprimer l'équité. Elle est la seule à induire une certaine symétrie entre les intervenants qui semble naturelle lorsque l'on considère le résultat d'une fusion comme une moyenne entre les différentes bases de connaissance.

On ne peut pas aller beaucoup plus loin dans le renforcement de cette propriété. Il n'est par exemple pas possible de remplacer les bases de connaissance par des ensembles de connaissance car cela conduit alors à des comportements contre-intuitifs et les postulats ainsi obtenus sont inconsistants avec les autres postulats des opérateurs de fusion.

6.2.2 Majorité

L'idée du comportement majoritaire est très intuitive. Mais il y a différentes façons d'exprimer la majorité, on peut par exemple faire une distinction entre majorité relative et majorité absolue. Plus généralement la majorité induit l'existence d'un quorum requis pour que l'opinion en question l'emporte.

Nous avons décidé d'être très large sur cette notion de majorité. Nous n'avons donc pas fixé de quorum mais simplement supposé son existence.

$$(Maj) \quad \exists n \quad \Delta_\mu(\Psi_1 \sqcup \Psi_2^n) \vdash \Delta_\mu(\Psi_2) \quad \textbf{(majorité)}$$

Une autre expression de la majorité, proposée dans [KP98], est la suivante :

$$(Ma2) \quad \exists n \quad \Delta_\mu(\Psi \sqcup \varphi^n) \vdash \Delta_\mu(\varphi)$$

C'est-à-dire que si un avis est suffisamment représenté, alors le choix du groupe se conformera à l'avis majoritaire. Le postulat (Maj) n'est qu'une généralisation de cette propriété à un groupe.

Les autres définitions de la propriété de majorité que l'on pourrait donner seraient basées sur la définition d'un quorum et elles impliqueraient donc (Maj).

6.2.3 Arbitrage

Le comportement désiré pour les opérateurs d'arbitrage est de contenter au mieux tous les intervenants, c'est-à-dire de contenter au mieux l'intervenant le plus défavorisé. Nous avons exprimé cela au travers de la propriété suivante :

$$(Arb) \quad \left. \begin{array}{l} \Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2) \\ \Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2) \\ \mu_1 \not\vdash \mu_2 \\ \mu_2 \not\vdash \mu_1 \end{array} \right\} \Rightarrow \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow \Delta_{\mu_1}(\varphi_1) \quad (\text{arbitrage})$$

Les propriétés proposées dans [KP98] pour modéliser cette notion étaient celles d'indépendance de la majorité (MI) et celle d'indépendance faible de la majorité (WMI) présentées dans la section précédente. Comme il a été souligné alors, (MI) n'est pas consistant avec les postulats de la fusion contrainte. D'un autre côté, (WMI) ne traduit pas directement un comportement d'arbitrage mais traduit simplement une propriété de non-majorité. Soulignons finalement qu'il n'y a pas de lien logique immédiat entre (Arb) et les deux propriétés (WMI) et (MI).

Un autre moyen d'exprimer la notion d'arbitrage différemment est de considérer que les agents ont un droit de veto. Les opérateurs d'arbitrage sont alors les opérateurs qui permettent aux protagonistes d'exprimer leur droit de veto sur certaines alternatives. C'est-à-dire que si une alternative est trop loin des connaissances de l'agent, alors elle ne sera pas, dans la mesure du possible, retenue dans le résultat de la fusion.

Un premier essai de formalisation de cette idée est la propriété :

$$(Ar3) \quad \left. \begin{array}{l} \forall \mu_i \exists \mu_i \forall \mu \Delta_{\mu_i \vee \mu}(\varphi_i) \vdash \mu \\ \forall i \mu' \wedge \mu_i \vdash \perp \end{array} \right\} \Rightarrow \Delta_{\mu' \vee \bigvee \mu_i}(\sqcup \varphi_i) \vdash \mu'$$

Une autre formulation assez proche mais un peu moins exigeante est :

$$(Ar4) \quad \left. \begin{array}{l} \forall \mu_i \exists \mu_i \forall \mu \Delta_{\mu_i \vee \mu}(\varphi_i) \vdash \mu \\ \forall i \mu' \wedge \mu_i \vdash \perp \end{array} \right\} \Rightarrow \Delta_{\top}(\sqcup \varphi_i) \wedge \bigvee \mu_i \vdash \perp$$

Dans les deux propriétés précédentes, les μ_i représentent intuitivement les formules qui contiennent les mondes possibles les "pires" pour la base de connaissance φ_i correspondante. C'est-à-dire que si l'on donne un autre choix μ à la base, le résultat de la fusion de cette base seule impliquera μ . La propriété (Ar3) dit alors que s'il existe une contrainte d'intégrité μ' différente de tout ces μ_i , alors la fusion de l'ensemble de connaissance sous les contraintes $\mu' \vee \bigvee \mu_i$ impliquera μ' . La propriété (Ar4) relâche un peu cette condition en demandant simplement alors que la fusion sans contrainte (par $\top$) donne un résultat différent des μ_i .

Ces propriétés tentent donc d'assurer le fait que si une formule est "la plus mauvaise" pour une base de connaissance et qu'il est possible d'en trouver une qui ne soit pas aussi mauvaise pour une autre des bases de l'ensemble, alors ce sera cette dernière qui sera préférée pour le résultat de la fusion.

6.3 Théorèmes de représentation

A présent que nous disposons d'une définition logique des opérateurs de fusion contrainte, nous allons donner des théorèmes de représentation pour chaque type d'opérateur.

Ces théorèmes fournissent une sémantique à ces opérateurs et en donnent une définition plus constructive.

Plus précisément, nous allons montrer que chaque opérateur de fusion contrainte correspond à une famille de pré-ordres sur les interprétations. Nous devons d'abord définir ce qu'est un assignement synchrétique¹ :

Définition 83 *Un assignement synchrétique est une fonction qui associe à chaque ensemble de connaissance Ψ un pré-ordre $\leq_\Psi$ sur les interprétations telle que pour tous ensembles de connaissance Ψ, Ψ_1, Ψ_2 et pour toutes bases de connaissance φ, φ' les conditions suivantes sont satisfaites :*

1. Si $I \models \Psi$ et $J \models \Psi$, alors $I \simeq_\Psi J$
2. Si $I \models \Psi$ et $J \not\models \Psi$, alors $I <_\Psi J$
3. Si $\Psi_1 \leftrightarrow \Psi_2$, alors $\leq_{\Psi_1} = \leq_{\Psi_2}$
4. $\forall I \models \varphi \exists J \models \varphi' J \leq_{\varphi \sqcup \varphi'} I$
5. Si $I \leq_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors $I \leq_{\Psi_1 \sqcup \Psi_2} J$
6. Si $I <_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors $I <_{\Psi_1 \sqcup \Psi_2} J$

La condition 1 dit que deux modèles de l'ensemble de connaissance sont équivalents pour le pré-ordre associé et la condition 2 dit qu'un modèle de l'ensemble de connaissance est toujours préféré à un contre-modèle. La condition 3 dit que si deux ensembles de connaissance sont équivalents alors les deux pré-ordres associés sont équivalents. Ces trois conditions sont une généralisation des conditions de l'assignement fidèle (cf section 1.2.4). La condition 4 dit que pour le pré-ordre associé à un ensemble de connaissance composé de deux bases de connaissance, pour chaque modèle de l'une, il existe un modèle de l'autre qui est au moins aussi bon. On peut également faire la remarque suivante :

Remarque 5 *Lorsque les pré-ordres générés par l'assignement synchrétique sont totaux, la condition 4 est équivalente à la condition suivante :*

$$4'. \exists J \models \varphi' \forall I \models \varphi J \leq_{\varphi \sqcup \varphi'} I$$

Preuve : La condition 4' implique la condition 4 directement. Pour montrer que la condition 4 implique la condition 4' il suffit de noter que $\leq_{\varphi \sqcup \varphi'}$ est un pré-ordre total, donc si on choisit $J \in \min(\text{Mod}(\varphi'), \leq_{\varphi \sqcup \varphi'})$ alors par transitivité $\forall I \models \varphi J \leq_{\varphi \sqcup \varphi'} I$, ce qui est la condition 4'.

¹ synchrétisme : n. m. Système philosophique ou religieux qui tend à fondre plusieurs doctrines différentes. (Larousse)

□

La condition 5 dit que si une interprétation est au moins aussi bonne qu'une autre pour un ensemble de connaissance Ψ_1 , et que cette interprétation est également au moins aussi bonne pour un ensemble de connaissance Ψ_2 , alors elle sera au moins aussi bonne que l'autre pour la réunion des deux ensembles de connaissance.

La condition 6 renforce un peu ce résultat en exigeant que si une interprétation est strictement meilleure qu'une autre pour un ensemble de connaissance Ψ_1 , et que cette interprétation est au moins aussi bonne pour un ensemble de connaissance Ψ_2 , alors cette interprétation doit être strictement meilleure que l'autre pour la réunion des deux ensembles de connaissance.

Nous devons également définir des assignements particuliers :

Définition 84 *Un assignement quasi-synchrétique est un assignement qui satisfait les conditions 1-5 et la condition 6' suivante :*

$$6'. \text{ Si } I <_{\Psi_1} J \text{ et } I <_{\Psi_2} J, \text{ alors } I <_{\Psi_1 \sqcup \Psi_2} J$$

Cette condition 6' est clairement plus faible que la condition 6 énoncée précédemment. Cette condition exige simplement que si une interprétation est strictement meilleure qu'une autre pour un ensemble de connaissance Ψ_1 et si cette interprétation est strictement meilleure que l'autre pour un ensemble de connaissance Ψ_2 , alors cette interprétation doit être strictement meilleure que l'autre pour la réunion des deux ensembles de connaissance.

Définition 85 *Un assignement synchrétique majoritaire est un assignement synchrétique qui satisfait la condition suivante :*

$$7. \text{ Si } I <_{\Psi_2} J, \text{ alors } \exists n \text{ } I <_{\Psi_1 \sqcup \Psi_2^n} J$$

Cette condition dit que si l'on répète un groupe Ψ_2 suffisamment de fois alors les préférences strictes de ce groupe seront respectées.

Définition 86 *Un assignement synchrétique juste est un assignement synchrétique qui satisfait la condition suivante :*

$$8. \left. \begin{array}{l} I <_{\varphi_1} J \\ I <_{\varphi_2} J' \\ J \simeq_{\varphi_1 \sqcup \varphi_2} J' \end{array} \right\} \Rightarrow I <_{\varphi_1 \sqcup \varphi_2} J$$

Cette condition dit que ce sont les choix médians qui sont préférés pour le groupe. Ce comportement est illustré figure 6.2.

Nous pouvons à présent énoncer le théorème de représentation pour les opérateurs de fusion contrainte :

Théorème 82 *Un opérateur Δ est un opérateur de fusion contrainte si et seulement si il existe un assignement synchrétique qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_{\Psi}$ tel que*

$$Mod(\Delta_{\mu}(\Psi)) = \min(Mod(\mu), \leq_{\Psi})$$

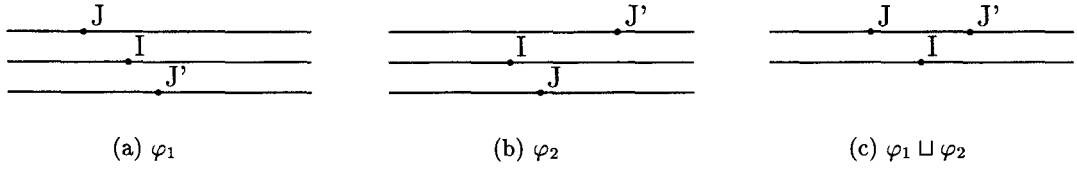


FIG. 6.2 – Arbitrage

Preuve : (Seulement si) Soit Δ un opérateur satisfaisant les propriétés (IC0-IC8). On définit un assignement de la façon suivante: pour chaque ensemble de connaissance Ψ on définit une relation $\leq_\Psi$ par $\forall I, J \in \Psi, I \leq_\Psi J$ si et seulement si $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi)$. Nous allons montrer que la relation $\leq_\Psi$ est un pré-ordre total et que l'assignement ainsi défini est un assignement synchrétique.

Nous montrons d'abord que $\leq_\Psi$ est un pré-ordre total :

Total : $\forall I, J \in \Psi$, de (IC1) $\Delta_{\varphi_{\{I, J\}}}(\Psi) \neq \emptyset$ et de (IC0) $\Delta_{\varphi_{\{I, J\}}}(\Psi) \vdash \varphi_{\{I, J\}}$, d'où $I \leq_\Psi J$ ou $J \leq_\Psi I$.

Réflexif : De (IC0) et (IC1) on obtient que $\Delta_{\varphi_{\{I\}}}(\Psi) = \varphi_{\{I\}}$. Donc $I \leq_\Psi I$.

Transitif : Supposons que $I \leq_\Psi J$ et $J \leq_\Psi L$. Par l'absurde, supposons que $I \not\leq_\Psi L$. Donc par définition et de (IC0) et (IC1) $\Delta_{\varphi_{\{I, L\}}}(\Psi) = \varphi_{\{L\}}$. Par (IC7) on peut trouver $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{I, L\}} \vdash \Delta_{\varphi_{\{I, L\}}}(\Psi)$. On considère deux cas :

Cas 1 : $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{I, L\}}$ est consistant, alors $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{I, L\}} \leftrightarrow \varphi_{\{L\}}$. On a donc $I \not\models \Delta_{\varphi_{\{I, J, L\}}}(\Psi)$. Mais par (IC1) $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \neq \emptyset$ donc par (IC0) $Mod(\Delta_{\varphi_{\{I, J, L\}}}(\Psi)) = \{J, L\}$ ou $Mod(\Delta_{\varphi_{\{I, J, L\}}}(\Psi)) = \{L\}$. Dans le premier cas par (IC7) et (IC8) on conclut que $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{I, J\}} \leftrightarrow \Delta_{\varphi_{\{I, J\}}}(\Psi)$ et alors $I \not\models \Delta_{\varphi_{\{I, J\}}}(\Psi)$. Contradiction. Dans le second cas par (IC7) et (IC8) $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{J, L\}} \leftrightarrow \Delta_{\varphi_{\{J, L\}}}(\Psi)$ mais $J \not\models \Delta_{\varphi_{\{I, J, L\}}}(\Psi)$ donc $J \not\models \Delta_{\varphi_{\{J, L\}}}(\Psi)$. Contradiction.

Cas 2 : $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{I, L\}}$ n'est pas consistant, donc $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) = \varphi_{\{J\}}$. Alors $\Delta_{\varphi_{\{I, J, L\}}}(\Psi) \wedge \varphi_{\{I, J\}} = \varphi_{\{J\}}$. Par (IC7) et (IC8) il s'ensuit que $\Delta_{\varphi_{\{I, J\}}}(\Psi) = \varphi_{\{J\}}$, ce qui par définition est $J <_\Psi I$. Contradiction.

Nous montrons à présent que $Mod(\Delta_\mu(\Psi)) = \min(Mod(\mu), \leq_\Psi)$. Premièrement pour l'inclusion $Mod(\Delta_\mu(\Psi)) \subseteq \min(Mod(\mu), \leq_\Psi)$ supposons que $I \models \Delta_\mu(\Psi)$ et par l'absurde, supposons que $I \notin \min(Mod(\mu), \leq_\Psi)$. Donc on peut trouver un $J \models \mu$ tel que $J <_\Psi I$, alors $I \not\models \Delta_{\varphi_{\{I, J\}}}(\Psi)$. Comme $\Delta_\mu(\Psi) \wedge \varphi_{\{I, J\}}$ est consistant, de (IC7) et (IC8) on a $\Delta_\mu(\Psi) \wedge \varphi_{\{I, J\}} \leftrightarrow \Delta_{\varphi_{\{I, J\}}}(\Psi)$. Mais $I \not\models \Delta_{\varphi_{\{I, J\}}}(\Psi)$ donc $I \not\models \Delta_\mu(\Psi)$. Contradiction.

Pour l'autre inclusion $Mod(\Delta_\mu(\Psi)) \supseteq \min(Mod(\mu), \leq_\Psi)$ supposons que $I \in \min(Mod(\mu), \leq_\Psi)$. On désire montrer que $I \models \Delta_\mu(\Psi)$. Comme $I \in \min(Mod(\mu), \leq_\Psi) \forall J \models \mu, I \leq_\Psi J$ et donc $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi)$. Comme $\Delta_\mu(\Psi) \wedge \varphi_{\{I, J\}}$ est consistant, de (IC7) et (IC8) on déduit $\Delta_\mu(\Psi) \wedge \varphi_{\{I, J\}} \leftrightarrow \Delta_{\varphi_{\{I, J\}}}(\Psi)$. Mais $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi)$ donc $I \models \Delta_\mu(\Psi)$.

Il reste à vérifier les conditions de l'assignement synchrétique :

1. Si $I \models \Psi$ et $J \models \Psi$, alors par (IC2) on a $\Delta_{\varphi_{\{I, J\}}}(\Psi) = \varphi_{\{I, J\}}$, donc $I \leq_\Psi J$ et $J \leq_\Psi I$

par définition et alors $I \simeq_{\Psi} J$.

2. Si $I \models \Psi$ et $J \not\models \Psi$, alors par (IC2) $\Delta_{\varphi_{\{I,J\}}}(\Psi) = \varphi_{\{I\}}$, donc $I \leq_{\Psi} J$ et $J \not\leq_{\Psi} I$, i.e. $I <_{\Psi} J$.
3. Directement de (IC3)
4. On désire montrer $\forall I \models \varphi \exists J \models \varphi' J \leq_{\varphi \sqcup \varphi'} I$. On montre d'abord que $\exists J \models \Delta_{\varphi \sqcup \varphi'}(\varphi \sqcup \varphi') \wedge \varphi'$. Si ce n'est pas le cas on a $\Delta_{\varphi \sqcup \varphi'}(\varphi \sqcup \varphi') \wedge \varphi' \vdash \perp$, de (IC0) et (IC1) on a alors $\Delta_{\varphi \sqcup \varphi'}(\varphi \sqcup \varphi') \vdash \varphi$, à présent par (IC4) on obtient que $\Delta_{\varphi \sqcup \varphi'}(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$. Contradiction.
Soit I un modèle de φ et prenons J tel que $J \models \Delta_{\varphi \sqcup \varphi'}(\varphi \sqcup \varphi') \wedge \varphi'$. On obtient de (IC7) et (IC8) que $J \models \Delta_{\varphi_{\{I,J\}}}(\varphi \sqcup \varphi')$. Donc $J \leq_{\varphi \sqcup \varphi'} I$.
5. Si $I \leq_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi_1) \wedge \Delta_{\varphi_{\{I,J\}}}(\Psi_2)$. Donc de (IC5) $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2)$ et par définition $I \leq_{\Psi_1 \sqcup \Psi_2} J$.
6. Supposons que $I <_{\Psi_1} J$ et $I \leq_{\Psi_2} J$. On veut montrer $I <_{\Psi_1 \sqcup \Psi_2} J$. Par hypothèse $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi_1) \wedge \Delta_{\varphi_{\{I,J\}}}(\Psi_2)$ et $J \not\models \Delta_{\varphi_{\{I,J\}}}(\Psi_1) \wedge \Delta_{\varphi_{\{I,J\}}}(\Psi_2)$. Donc de (IC5) et (IC6) $\Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2) = \varphi_{\{I\}}$. Alors $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2)$ et $J \not\models \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2)$ et par définition $I <_{\Psi_1 \sqcup \Psi_2} J$.

(Si) Soit un assignement syncrétique qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_{\Psi}$. On définit l'opérateur Δ en posant $Mod(\Delta_{\mu}(\Psi)) = \min(Mod(\mu), \leq_{\Psi})$. On montre que Δ satisfait les postulats (IC0-IC8).

- (IC0) Par définition $Mod(\Delta_{\mu}(\Psi)) \subseteq Mod(\mu)$.
- (IC1) Si μ est consistant, alors $Mod(\mu) \neq \emptyset$ et, comme il n'y a qu'un nombre fini d'interprétations, il n'y a pas de chaîne infinie descendante d'inégalités strictes, donc $\min(Mod(\mu), \leq_{\Psi}) \neq \emptyset$. Donc $\Delta_{\mu}(\Psi)$ est consistant.
- (IC2) Supposons que $\Psi \wedge \mu$ est consistant. On veut montrer que $\min(Mod(\mu), \leq_{\Psi}) = Mod(\Psi \wedge \mu)$. On peut noter d'abord que si $I \models \Psi$ alors des conditions 1 et 2, $I \in \min(W, \leq_{\Psi})$. Donc si $I \models \Psi \wedge \mu$ alors $I \in \min(Mod(\mu), \leq_{\Psi})$. Donc $\min(Mod(\mu), \leq_{\Psi}) \supseteq Mod(\Psi \wedge \mu)$. Pour l'autre inclusion on considère $I \in \min(Mod(\mu), \leq_{\Psi})$. Par l'absurde, supposons $I \not\models \Psi \wedge \mu$. Puisque $I \not\models \Psi$ par la condition 2 on a que $\forall J \models \Psi J <_{\Psi} I$. En particulier $\forall J \models \Psi \wedge \mu J <_{\Psi} I$. Donc $I \notin \min(Mod(\mu), \leq_{\Psi})$. Contradiction.
- (IC3) Directement par la condition 3 et la définition de Δ .
- (IC4) Supposons que $\varphi \vdash \mu$, $\varphi' \vdash \mu$, et $\Delta_{\mu}(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$, on veut montrer que $\Delta_{\mu}(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$. Considérons $I \models \Delta_{\mu}(\varphi \sqcup \varphi') \wedge \varphi$. Alors $\forall I' \models \mu I \leq_{\varphi \sqcup \varphi'} I'$. Mais par la condition 4 on a que $\exists J \models \varphi'$ tel que $J \leq_{\varphi \sqcup \varphi'} I$. Alors $\forall I' \models \mu J \leq_{\varphi \sqcup \varphi'} I'$. Alors $J \models \Delta_{\mu}(\varphi \sqcup \varphi')$ et donc $\Delta_{\mu}(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$.
- (IC5) Si $I \models \Delta_{\mu}(\Psi_1) \wedge \Delta_{\mu}(\Psi_2)$ alors $I \in \min(Mod(\mu), \leq_{\Psi_1})$ et donc $\forall J \models \mu I \leq_{\Psi_1} J$. De la même façon on a $\forall J \models \mu I \leq_{\Psi_2} J$. Donc par la condition 5 on obtient $\forall J \models \mu I \leq_{\Psi_1 \sqcup \Psi_2} J$. Donc $I \in \min(Mod(\mu), \leq_{\Psi_1 \sqcup \Psi_2})$. Donc par définition $I \models \Delta_{\mu}(\Psi_1 \sqcup \Psi_2)$.
- (IC6) Supposons que $\Delta_{\mu}(\Psi_1) \wedge \Delta_{\mu}(\Psi_2)$ est consistant. On veut montrer que $\Delta_{\mu}(\Psi_1 \sqcup \Psi_2) \vdash \Delta_{\mu}(\Psi_1) \wedge \Delta_{\mu}(\Psi_2)$. Prenons $I \models \Delta_{\mu}(\Psi_1 \sqcup \Psi_2)$, donc $\forall J \models \mu I \leq_{\Psi_1 \sqcup \Psi_2} J$. Vers une contradiction, supposons que $I \not\models \Delta_{\mu}(\Psi_1) \wedge \Delta_{\mu}(\Psi_2)$. Donc $I \not\models \Delta_{\mu}(\Psi_1)$ ou $I \not\models \Delta_{\mu}(\Psi_2)$. Supposons que $I \not\models \Delta_{\mu}(\Psi_1)$ (l'autre cas est symétrique). Comme $\Delta_{\mu}(\Psi_1) \wedge \Delta_{\mu}(\Psi_2)$ est

- consistant $\exists J \models \Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$. Donc $J <_{\Psi_1} I$ et $J \leq_{\Psi_2} I$ alors par la condition 6 $J <_{\Psi_1 \sqcup \Psi_2} I$ et donc $I \not\models \Delta_\mu(\Psi_1 \sqcup \Psi_2)$. Contradiction.
- (IC7) Considérons $I \models \Delta_{\mu_1}(\Psi) \wedge \mu_2$. On a $\forall J \models \mu_1 \ I \leq_\Psi J$. Donc $\forall J \models \mu_1 \wedge \mu_2 \ I \leq_\Psi J$, alors $I \models \Delta_{\mu_1 \wedge \mu_2}(\Psi)$.
- (IC8) Supposons que $\Delta_{\mu_1}(\Psi) \wedge \mu_2$ est consistant, alors $\exists J \models \Delta_{\mu_1}(\Psi) \wedge \mu_2$. Considérons $I \models \Delta_{\mu_1 \wedge \mu_2}(\Psi)$ et supposons que $I \not\models \Delta_{\mu_1}(\Psi)$. Donc $J <_\Psi I$ et comme $J \models \mu_1 \wedge \mu_2$ alors $I \notin \min(\text{Mod}(\mu_1 \wedge \mu_2), \leq_\Psi)$. Et $I \not\models \Delta_{\mu_1 \wedge \mu_2}(\Psi)$. Contradiction.

□

Ce théorème exprime la correspondance entre opérateurs de fusion contrainte et assignements synchroniques. C'est-à-dire que les conditions logiques (IC0-IC8) données pour les opérateurs de fusion contrainte correspondent aux conditions 1-6 sur les pré-ordres associés aux ensembles de connaissance.

Lorsque ce théorème est vérifié, on dit que l'assignement synchronique correspond à l'opérateur de fusion contrainte.

Avant de donner un théorème similaire pour les opérateurs de quasi-fusion, on peut souligner que l'on peut trouver des correspondances plus faibles que celle donnée ci-dessus, par exemple :

Corollaire 83 *Un opérateur satisfait les conditions (IC0-IC5), (IC7) et (IC8) si et seulement si il correspond à un assignement vérifiant les conditions 1-5.*
Un opérateur satisfait les conditions (IC0-IC3), (IC5-IC8) si et seulement si il correspond à un assignement vérifiant les conditions 1-3, 5 and 6.

La justification de ce résultat vient d'une analyse de la preuve du théorème précédent. Il suffit de noter que le postulat (IC6) n'intervient que dans la preuve de la condition 6 et, symétriquement, la condition 6 n'est utilisée que pour la preuve du postulat (IC6). De la même manière (IC4) correspond à la condition 4 sur l'assignement.

On peut donc à présent donner le théorème de représentation pour les opérateurs de quasi-fusion :

Théorème 84 *Un opérateur Δ est un opérateur de quasi-fusion si et seulement si il existe un assignement quasi-synchronique qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$ tel que*

$$\text{Mod}(\Delta_\mu(\Psi)) = \min(\text{Mod}(\mu), \leq_\Psi)$$

Preuve : (Seulement si) Soit Δ un opérateur vérifiant les postulats (IC0-IC5), (IC6'), (IC7) et (IC8). On définit un assignement de la façon suivante : pour chaque ensemble de connaissance Ψ on définit une relation $\leq_\Psi$ par $\forall I, J \in I \leq_\Psi J$ si et seulement si $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi)$.

Par le corollaire 83 cet assignement correspondant à Δ satisfait les conditions 1-5. Il reste simplement à prouver la condition 6'. Supposons que $I <_{\Psi_1} J$ et $I <_{\Psi_2} J$. On veut montrer que $I <_{\Psi_1 \sqcup \Psi_2} J$. Par hypothèse $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi_1) \wedge \Delta_{\varphi_{\{I, J\}}}(\Psi_2)$ et $J \not\models \Delta_{\varphi_{\{I, J\}}}(\Psi_1) \vee \Delta_{\varphi_{\{I, J\}}}(\Psi_2)$.

Donc de (IC6') $\Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2) = \varphi_{\{I\}}$. Alors $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2)$ et $J \not\models \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2)$ et par définition $I <_{\Psi_1 \sqcup \Psi_2} J$.

(Si) Soit un assignement vérifiant les conditions 1-5 et 6' qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$. On définit l'opérateur Δ en posant $Mod(\Delta_\mu(\Psi)) = \min(Mod(\mu), \leq_\Psi)$. Par le corollaire 83 nous savons que Δ vérifie (IC0-IC5), (IC7) et (IC8). Il reste à prouver (IC6').

On veut montrer que si $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant, alors $\Delta_\mu(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\mu(\Psi_1) \vee \Delta_\mu(\Psi_2)$. Supposons que $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant et prenons $I \models \Delta_\mu(\Psi_1 \sqcup \Psi_2)$, donc $\forall J \models \mu \ I \leq_{\Psi_1 \sqcup \Psi_2} J$. La condition 6' est équivalente à

$$I \leq_{\Psi_1 \sqcup \Psi_2} J \Rightarrow I \leq_{\Psi_1} J \text{ et } I \leq_{\Psi_2} J$$

Par l'absurde, supposons que $I \not\models \Delta_\mu(\Psi_1) \vee \Delta_\mu(\Psi_2)$, c'est-à-dire $I \not\models \Delta_\mu(\Psi_1)$ et $I \not\models \Delta_\mu(\Psi_2)$. Ce qui peut être réécrit en $\exists J_1 \models \mu \ J_1 <_{\Psi_1} I$ et $\exists J_2 \models \mu \ J_2 <_{\Psi_2} I$. Mais $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant, donc $\exists J_3 \models \mu$ tel que $\forall I' \models \mu \ J_3 \leq_{\Psi_1} I'$ et $J_3 \leq_{\Psi_2} I'$. En particulier on a $J_3 \leq_{\Psi_1} J_1$ et $J_3 \leq_{\Psi_2} J_2$. Par transitivité on trouve $J_3 <_{\Psi_1} I$ et $J_3 <_{\Psi_2} I$, et par la condition 6' on conclut $J_3 <_{\Psi_1 \sqcup \Psi_2} I$. Donc $I \not\models \Delta_\mu(\Psi_1 \sqcup \Psi_2)$. Contradiction.

□

Nous disposons également d'un théorème de représentation pour les opérateurs de fusion majoritaire et pour les opérateurs d'arbitrage. Ce sont les deux théorèmes suivants.

Théorème 85 *Un opérateur Δ est un opérateur de fusion majoritaire si et seulement si il existe un assignement synchrétique majoritaire qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$ tel que*

$$Mod(\Delta_\mu(\Psi)) = \min(Mod(\mu), \leq_\Psi)$$

Preuve : (Seulement si) Soit Δ un opérateur vérifiant les postulats (IC0-IC8) et (Maj). On définit un assignement de la façon suivante : pour chaque ensemble de connaissance Ψ on définit une relation $\leq_\Psi$ par $\forall I, J \in I \leq_\Psi J$ si et seulement si $I \models \Delta_{\varphi_{\{I,J\}}}(\Psi)$.

Par le théorème 82 c'est un assignement synchrétique. Il reste simplement à prouver la condition 7.

Supposons que $I <_{\Psi_2} J$. Alors $\Delta_{\varphi_{\{I,J\}}}(\Psi_2) = \varphi_{\{I\}}$. De (Maj) on a que $\exists n$ tel que $\Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2^n) \vdash \Delta_{\varphi_{\{I,J\}}}(\Psi_2)$, donc $\exists n \Delta_{\varphi_{\{I,J\}}}(\Psi_1 \sqcup \Psi_2^n) = \varphi_{\{I\}}$, i.e. $\exists n \ I <_{\Psi_1 \sqcup \Psi_2^n} J$.

(Si) Soit un assignement synchrétique majoritaire qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$. On définit l'opérateur Δ en posant $Mod(\Delta_\mu(\Psi)) = \min(Mod(\mu), \leq_\Psi)$.

Par le théorème 82 on sait que Δ vérifie (IC0-IC8). Il reste à montrer (Maj).

Des conditions 6 et 7 on déduit directement la condition suivante :

$$I <_{\Psi_2} J \Rightarrow \exists n_0 \forall n \geq n_0 \ I <_{\Psi_1 \sqcup \Psi_2^n} J$$

Puisque pour chaque Ψ , $\leq_\Psi$ est total la contraposée de cette condition peut s'écrire :

$$\forall n_0 \exists n \geq n_0 \ I \leq_{\Psi_1 \sqcup \Psi_2^n} J \Rightarrow I \leq_{\Psi_2} J \quad (*)$$

A présent, par l'absurde, supposons que $\forall n \Delta_\mu(\Psi_1 \sqcup \Psi_2^n) \not\vdash \Delta_\mu(\Psi_2)$. Par cette hypothèse on obtient que $\forall n \exists I \models \mu \forall J \models \mu I \leq_{\Psi_1 \sqcup \Psi_2^n} J$ et $\exists J' \models \mu J' <_{\Psi_2} I$. Puisque le nombre d'interprétations est fini, il existe I tel que $I \leq_{\Psi_1 \sqcup \Psi_2^n} J$ pour tout $J \models \mu$ et une infinité d'entiers n^2 . Cela est exactement les prémisses de la condition (*), on a donc $I \leq_{\Psi_2} J$ pour tout $J \models \mu$, ce qui contredit me fait $J' <_{\Psi_2} I$. $\square$

Théorème 86 *Un opérateur Δ est un opérateur d'arbitrage si et seulement si il existe un assignement syncrétique juste qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$ tel que*

$$Mod(\Delta_\mu(\Psi)) = \min(Mod(\mu), \leq_\Psi)$$

Preuve : (Seulement si) Soit Δ un opérateur vérifiant les postulats (IC0-IC8) et (Arb). On définit un assignement de la façon suivante : pour chaque ensemble de connaissance Ψ on définit une relation $\leq_\Psi$ par $\forall I, J \in \Psi I \leq_\Psi J$ si et seulement si $I \models \Delta_{\varphi_{\{I, J\}}}(\Psi)$.

Par le théorème 82 on sait que cet assignement est un assignement syncrétique, il ne reste donc plus qu'à montrer la condition 8. Supposons que $J <_{\varphi_1} I$, $J <_{\varphi_2} J'$ et $I \simeq_{\varphi_1 \sqcup \varphi_2} J'$. D'abord si $I = J'$ alors $J <_{\varphi_1 \sqcup \varphi_2} I$ se déduit de la condition 6. A présent si $I \neq J'$, alors $\Delta_{\varphi_{\{I, J\}}}(\varphi_1) \leftrightarrow \Delta_{\varphi_{\{J, J'\}}}(\varphi_2) = \varphi_{\{J\}}$. De plus $\Delta_{\varphi_{\{I, J'\}}}(\varphi_1 \sqcup \varphi_2) = \varphi_{\{I, J'\}}$, et $\varphi_{\{I, J\}} \wedge \neg \varphi_{\{I, J'\}}$ et $\varphi_{\{I, J'\}} \wedge \neg \varphi_{\{I, J\}}$ sont tous deux consistants. Alors par (Arb) on a que $\Delta_{\varphi_{\{I, J, J'\}}}(\varphi_1 \sqcup \varphi_2) = \varphi_{\{J\}}$. Et par (IC7) et (IC8) on en conclut $\Delta_{\varphi_{\{I, J\}}}(\varphi_1 \sqcup \varphi_2) = \varphi_{\{J\}}$, c'est-à-dire $J <_{\varphi_1 \sqcup \varphi_2} I$.

(Si) Soit un assignement syncrétique juste qui associe à chaque ensemble de connaissance Ψ un pré-ordre total $\leq_\Psi$. On définit l'opérateur Δ en posant $Mod(\Delta_\mu(\Psi)) = \min(Mod(\mu), \leq_\Psi)$.

Nous savons grâce au théorème 82 que Δ satisfait (IC0-IC8), il est donc suffisant de montrer (Arb).

Supposons que $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$, $\Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2)$, $\mu_1 \wedge \neg \mu_2 \not\vdash \perp$ et $\mu_2 \wedge \neg \mu_1 \not\vdash \perp$. On veut montrer que $\Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow \Delta_{\mu_1}(\varphi_1)$.

On montre d'abord $\Delta_{\mu_1}(\varphi_1) \vdash \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Soit $I \models \Delta_{\mu_1}(\varphi_1)$, supposons vers une contradiction que $I \not\models \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Alors $\exists J \models \mu_1 \vee \mu_2 J <_{\varphi_1 \sqcup \varphi_2} I$.

Nous considérons 3 cas : $J \models \mu_1 \wedge \mu_2$, $J \models \mu_1 \wedge \neg \mu_2$ ou $J \models \neg \mu_1 \wedge \mu_2$.

cas 1 : $J \models \mu_1 \wedge \mu_2$. Puisque $I \models \Delta_{\mu_1}(\varphi_1)$, $I \leq_{\varphi_1} J$. Par hypothèse $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$. Donc $I \models \Delta_{\mu_2}(\varphi_2)$ et alors $I \leq_{\varphi_2} J$. Alors par la condition 5 on a $I \leq_{\varphi_1 \sqcup \varphi_2} J$. Contradiction.

cas 2 : $J \models \mu_1 \wedge \neg \mu_2$ (le cas 3, $J \models \neg \mu_1 \wedge \mu_2$, est symétrique). Puisque $J \not\models \varphi_2$ et $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$ on a $J \not\models \Delta_{\mu_1}(\varphi_1)$, donc $I <_{\varphi_1} J$. Par hypothèse on peut trouver un $J' \models \mu_2 \wedge \neg \mu_1$ et avec un argument analogue $I <_{\varphi_2} J'$. On sait également que $\Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2)$, cela implique $J \simeq_{\varphi_1 \sqcup \varphi_2} J'$. Et par la condition 8 on obtient $I <_{\varphi_1 \sqcup \varphi_2} J$. Contradiction.

A présent on montre que $\Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \vdash \Delta_{\mu_1}(\varphi_1)$. Supposons que $I \models \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$ et, par l'absurde, supposons que $I \not\models \Delta_{\mu_1}(\varphi_1)$. On considère les trois cas suivants :

2. Nous appliquons ici le principe de Dirichlet ou "principe du tiroir", i.e. si l'on remplit une infinité de tiroirs avec des nombres compris entre 0 et n , à raison d'un nombre par tiroir, alors il y a une infinité de tiroirs contenant un même nombre.

cas 1 : $I \models \mu_1 \wedge \mu_2$ alors $\exists J \models \Delta_{\mu_1}(\varphi_1)$, donc $J <_{\varphi_1} I$. Et, comme $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$, $J <_{\varphi_2} I$. Donc par la condition 6 on a que $J <_{\varphi_1 \sqcup \varphi_2} I$, donc $I \not\models \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Contradiction.

cas 2 : $I \models \mu_1 \wedge \neg \mu_2$ (le cas 3, où $I \models \neg \mu_1 \wedge \mu_2$, est symétrique). Par hypothèse on sait que $\exists I' \models \neg \mu_1 \wedge \mu_2$. Puisque $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$ $\exists J \models \Delta_{\mu_1}(\varphi_1)$ tel que $J <_{\varphi_1} I$ et $J <_{\varphi_2} I'$. On déduit également de $\Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2)$ que $I \simeq_{\varphi_1 \sqcup \varphi_2} I'$, donc par la condition 8 on a que $J <_{\varphi_1 \sqcup \varphi_2} I$. Donc $I \not\models \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Contradiction.

□

Chapitre 7

Quelques opérateurs de fusion

Dans ce chapitre nous allons donner des exemples d'opérateurs de fusion.

Nous donnons tout d'abord des opérateurs simples qui ne sont pas des opérateurs de fusion contrainte. Ceci montre que les postulats sont suffisamment forts pour éliminer des opérateurs trop "basiques".

Nous présentons ensuite trois familles d'opérateurs définis à partir d'une distance entre interprétations. La première famille d'opérateurs présentée définit des opérateurs de quasi-fusion. La seconde donne des opérateurs de fusion majoritaire et la troisième des opérateurs d'arbitrage.

Finalement nous montrons qu'il est possible d'être à la fois un opérateur majoritaire et un opérateur d'arbitrage.

7.1 Opérateurs simples

Le premier opérateur que l'on peut considérer est un opérateur inspiré de la révision par intersection totale, c'est-à-dire :

Définition 87 *L'opérateur de fusion par intersection totale est défini par :*

$$\Delta_{\mu}^{fm}(\Psi) = \begin{cases} \Psi \wedge \mu & \text{si consistant} \\ \mu & \text{sinon} \end{cases}$$

Mais cet opérateur n'est pas un opérateur de fusion contrainte car il ne satisfait pas (IC6).

Théorème 87 *L'opérateur de fusion par intersection totale satisfait les postulats (IC0)-(IC5), (IC6'), (IC7) et (IC8). Mais il ne satisfait pas (IC6).*

Preuve : Les preuves pour les postulats (IC0)-(IC5), (IC6'), (IC7) et (IC8) sont directes. Pour montrer que (IC6) n'est pas satisfait il suffit de choisir Ψ_1 , Ψ_2 et μ tels que $\bigwedge \Psi_1 \wedge \mu$ est consistant, $\bigwedge \Psi_2 \wedge \mu$ n'est pas consistant et $\bigwedge \Psi_1 \wedge \mu \not\leq \mu$. Dans ce cas on a $\Delta_{\mu}(\Psi_1) = \bigwedge \Psi_1 \wedge \mu$ et $\Delta_{\mu}(\Psi_2) = \mu$, et $\Delta_{\mu}(\Psi_1 \sqcup \Psi_2) = \mu$. Ce qui contredit (IC6). □

Une autre méthode généralement acceptée pour la fusion est de prendre la conjonction des bases de connaissance si celle-ci est consistante et de prendre la disjonction sinon. La généralisation de cette méthode en présence de contraintes d'intégrité est la suivante :

Définition 88 Soit $\Psi = \{\varphi_1, \dots, \varphi_n\}$, l'opérateur de fusion basique est défini par :

$$\Delta_\mu^b(\Psi) = \begin{cases} \bigwedge \varphi_i \wedge \mu & \text{si consistant, sinon} \\ \bigvee \varphi_i \wedge \mu & \text{si consistant} \\ \mu & \text{sinon} \end{cases}$$

Cet opérateur n'est pas un opérateur de fusion contrainte car il ne satisfait pas (IC6).

Théorème 88 L'opérateur de fusion basique satisfait les postulats (IC0)-(IC5), (IC7) et (IC8). Mais il ne satisfait pas (IC6) et (IC6').

Preuve : Les preuves pour les postulats (IC0)-(IC3), (IC7) et (IC8) sont directes par définition.

Pour (IC4) on doit considérer deux cas. Soit $\varphi \wedge \varphi' \not\vdash \perp$, dans ce cas $\Delta_\mu(\varphi \sqcup \varphi') = \varphi \wedge \varphi'$ et (IC4) est satisfait. Soit $\varphi \wedge \varphi' \vdash \perp$, alors $\Delta_\mu(\varphi \sqcup \varphi') = \varphi \vee \varphi'$ et (IC4) est satisfait.

Pour montrer (IC5) il suffit d'énumérer les différents cas possibles pour Ψ_1 et Ψ_2 et constater que (IC5) est vérifié dans tous les cas.

(IC6) et (IC6') ne sont pas vérifiés. En effet il suffit de choisir Ψ_1 , Ψ_2 et μ tels que $\bigwedge \Psi_1 \wedge \mu$ est consistant, $\bigvee \Psi_2 \wedge \mu$ inconsistant et $\bigwedge \Psi_1 \wedge \mu \not\vdash \bigvee \Psi_1 \wedge \mu$. On a alors $\Delta_\mu(\Psi_1) = \bigwedge \Psi_1 \wedge \mu$, $\Delta_\mu(\Psi_2) = \mu$ et $\Delta_\mu(\Psi_1 \sqcup \Psi_2) = \bigvee (\Psi_1 \sqcup \Psi_2) \wedge \mu$. Ce qui contredit (IC6) et (IC6').

□

7.2 Distances

Les opérateurs présentés dans les prochaines sections sont définis à partir d'une distance sur les interprétations.

Définition 89 Une distance sur les interprétations est une fonction $d : \mathcal{W} \times \mathcal{W} \mapsto \mathbb{N}$ telle que :

$$\begin{aligned} d(I, J) &= d(J, I) \\ d(I, J) &= 0 \text{ ssi } I = J \end{aligned}$$

On dit que l'on dispose d'une distance topologique si d vérifie l'inégalité triangulaire :

$$d(I, J) \leq d(I, J') + d(J', J)$$

Cette distance sur les interprétations induit de manière naturelle une distance entre une interprétation et une base de connaissance, et une distance entre bases de connaissance.

Cette dernière distance n'est pas nécessaire pour la définition des opérateurs mais sera utile pour les preuves.

Définition 90 Soit d une distance sur les interprétations. On définit la distance (induite) entre une interprétation I et une base de connaissance φ de la façon suivante :

$$d(I, \varphi) = \min_{J \models \varphi} d(I, J)$$

La distance (induite) entre deux bases de connaissance φ et φ' est la suivante :

$$d(\varphi, \varphi') = \min_{I \models \varphi, J \models \varphi'} d(I, J)$$

La distance entre une interprétation et une base de connaissance est donc la distance minimale entre cette interprétation et les modèles de la base de connaissance.

Les familles d'opérateurs que nous allons définir dans la suite de ce chapitre sont des constructions à partir d'une distance donnée. Pour définir un opérateur particulier, il suffit donc de choisir sa famille et une distance entre interprétations.

Une distance entre interprétations bien connue est la *distance de Dalal* [Dal88b, Dal88a]. Elle consiste à calculer la distance de Hamming entre les interprétations, c'est-à-dire que la distance entre deux interprétations est le nombre de variables propositionnelles sur lesquelles ces interprétations diffèrent. Par exemple, la distance de Dalal entre (0,0,1,1) et (0,0,0,0) est 2, car ces deux interprétations diffèrent sur les deux dernières variables.

Une remarque importante à propos de la définition de ces opérateurs est que l'on utilise les théorèmes de représentation de la section 6.3. Cela montre le caractère constructif de cette représentation.

Une autre remarque est que tous ces opérateurs classent l'ensemble des interprétations suivant leur vraisemblance pour le groupe (l'ensemble de connaissance). Le résultat de la fusion est alors l'ensemble des interprétations les plus vraisemblables.

Cette vraisemblance est donnée par une notion de proximité, les interprétations les plus vraisemblables sont celles les plus proches de l'ensemble de connaissance. Et la différence entre les opérateurs réside dans la définition de cette notion de proximité.

7.3 Opérateurs Max

Nous définissons dans cette section les opérateurs $\Delta^{d,Max}$. Comme nous le verrons ces opérateurs ne sont pas assez fins pour être des opérateurs de fusion contrainte. Ce sont néanmoins des opérateurs de quasi-fusion.

Et le comportement de cette famille d'opérateurs peut être vu comme une approximation du comportement que nous attendons des opérateurs d'arbitrage.

Définition 91 Soient un ensemble de connaissance Ψ , une interprétation I et une distance entre interprétations d . On définit la distance *Max* entre une interprétation et un ensemble de connaissance comme :

$$d_{d,Max}(I, \Psi) = \max_{\varphi \in \Psi} d(I, \varphi)$$

Cela induit le pré-ordre suivant :

$$I \leq_{\Psi}^{d,Max} J \text{ ssi } d_{d,Max}(I, \Psi) \leq d_{d,Max}(J, \Psi)$$

Et l'opérateur $\Delta^{d,Max}$ est défini par :

$$Mod(\Delta_{\mu}^{d,Max}(\Psi)) = \min(Mod(\mu), \leq_{\Psi}^{d,Max})$$

L'idée de ces opérateurs est de trouver les interprétations les plus proches de l'ensemble de connaissance dans son ensemble. Cela semble être une bonne définition pour un opérateur d'arbitrage mais ces opérateurs ne vérifient pas toutes les propriétés.

Théorème 89 $\Delta^{d,Max}$ est un opérateur de quasi-fusion. Il vérifie (Arb), (IC_{it}) and (MI), et ne vérifie pas (IC6) et (Maj).

Preuve : Notons tout d'abord que, par définition, nous avons la propriété suivante :

$$d_{d,Max}(I, \Psi_1 \sqcup \Psi_2) = \max\{d_{d,Max}(I, \Psi_1), d_{d,Max}(I, \Psi_2)\} \quad (7.1)$$

Nous montrons que l'assignement $\Psi \mapsto \leq_{\Psi}^{d,Max}$ est un assignement quasi-synchrétique. Donc par le théorème 84 $\Delta^{d,Max}$ est un opérateur de quasi-fusion.

1. Si $I \models \Psi$ et $J \models \Psi$, alors $d_{d,Max}(I, \Psi) = 0$ et $d_{d,Max}(J, \Psi) = 0$, donc $I \simeq_{\Psi} J$.
2. Si $I \models \Psi$ et $J \not\models \Psi$, alors $d_{d,Max}(I, \Psi) = 0$ et $d_{d,Max}(J, \Psi) > 0$, donc $I <_{\Psi} J$.
3. Direct.
4. On veut montrer $\forall I \models \varphi \exists J \models \varphi' J \leq_{\varphi \sqcup \varphi'} I$. On sait que $d(I, \varphi) = 0$ et $d(I, \varphi') = \min_{J \models \varphi'} d(I, J)$, choisissons donc $J \models \varphi'$ tel que $d(I, J) = d(I, \varphi')$. Alors $d(J, \varphi) = \min_{I' \models \varphi} d(J, I') \leq d(J, I)$, et $d(J, \varphi') = 0$. D'où $d_{d,Max}(J, \varphi \sqcup \varphi') = d_{d,Max}(J, \varphi) \leq d_{d,Max}(I, \varphi') = d_{d,Max}(I, \varphi \sqcup \varphi')$. Donc par définition $J \leq_{\varphi \sqcup \varphi'} I$.
5. Si $I \leq_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors on a $d_{d,Max}(I, \Psi_1) \leq d_{d,Max}(J, \Psi_1)$ et $d_{d,Max}(I, \Psi_2) \leq d_{d,Max}(J, \Psi_2)$, donc par la propriété 7.1 $d_{d,Max}(I, \Psi_1 \sqcup \Psi_2) \leq d_{d,Max}(J, \Psi_1 \sqcup \Psi_2)$.
- 6'. Si $I <_{\Psi_1} J$ et $I <_{\Psi_2} J$, alors on a $d_{d,Max}(I, \Psi_1) < d_{d,Max}(J, \Psi_1)$ et $d_{d,Max}(I, \Psi_2) < d_{d,Max}(J, \Psi_2)$, donc par la propriété 7.1 $d_{d,Max}(I, \Psi_1 \sqcup \Psi_2) < d_{d,Max}(J, \Psi_1 \sqcup \Psi_2)$.

Prouvons à présent que (Arb) est vérifié : Supposons que $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$, $\Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2), \mu_1 \not\models \mu_2, \mu_2 \not\models \mu_1$. On désire montrer $\Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow \Delta_{\mu_1}(\varphi_1)$. D'abord $\Delta_{\mu_1}(\varphi_1) \vdash \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Comme $\mu_1 \not\models \mu_2$ on sait que $\Delta_{\mu_1}(\varphi_1) \not\models \mu_1$, donc $\exists J \models \mu_1 \wedge \neg \mu_2$ tel que $\forall I \models \Delta_{\mu_1}(\varphi_1) I <_{\varphi_1} J$, donc par hypothèse $\forall J \models \mu_1 \leftrightarrow \neg \mu_2 \forall I \models \Delta_{\mu_1}(\varphi_1) I <_{\varphi_1} J$. Comme $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$ on a également $\forall J \models \mu_1 \leftrightarrow \neg \mu_2 \forall I \models \Delta_{\mu_1}(\varphi_1) I <_{\varphi_2} J$. On a donc $\forall J \models \mu_1 \leftrightarrow \neg \mu_2 \forall I \models \Delta_{\mu_1}(\varphi_1) I <_{\varphi_1 \sqcup \varphi_2} J$. De plus on a $\forall J \models \mu_1 \wedge \mu_2 \forall I \models \Delta_{\mu_1}(\varphi_1) I \leq_{\varphi_1 \sqcup \varphi_2} J$. D'où $\Delta_{\mu_1}(\varphi_1) \vdash \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Il reste à montrer $\Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \vdash \Delta_{\mu_1}(\varphi_1)$ Soit $I \models \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$ et par l'absurde, supposons que $I \not\models \Delta_{\mu_1}(\varphi_1)$. Alors $\exists J \models \mu_1$ tel que $J \models \Delta_{\mu_1}(\varphi_1)$, donc $J <_{\varphi_1} I$. Par hypothèse $\Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2)$, donc $J <_{\varphi_2} I$. Alors $J <_{\varphi_1 \sqcup \varphi_2} I$, donc $I \not\models \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2)$. Contradiction.

Nous montrons à présent que (IC_{it}) est vérifié. Supposons que $\varphi \vdash \mu$ et $\varphi' \vdash \mu$. Il suffit de montrer $\exists n \Delta_{\mu}^{d,Max^n}(\varphi', \varphi) \vdash \varphi$ car on en déduit alors (IC_{it}) avec $\varphi' = \Delta_{\mu}^{d,Max}(\Psi \sqcup \varphi)$. Soit a la distance entre φ et φ' . On raisonne par induction sur a :

Si $a = 0$, alors $\Delta_{\mu}^{d,Max}(\varphi' \sqcup \varphi) = \varphi \wedge \varphi'$ et $\Delta_{\mu}^{d,Max}(\Delta_{\mu}^{d,Max}(\varphi' \sqcup \varphi) \sqcup \varphi) \vdash \varphi$. Supposons que $a \geq 1$ et pour tout i si $i < a$ on a que $d(\varphi, \varphi') = i$ implique $\exists n \Delta_{\mu}^{d,Max^n}(\varphi', \varphi) \vdash \varphi$. On veut montrer que si $d(\varphi, \varphi') = a$ alors $\exists n \Delta_{\mu}^{d,Max^n}(\varphi', \varphi) \vdash \varphi$. Puisque $d(\varphi, \varphi') = a$ on peut choisir I, I' tels que $I \models \varphi$, $I' \models \varphi'$ et $d(I, I') = a$. Si $\min\{d_{d,Max}(I', \varphi \sqcup \varphi') : I' \models \mu\} = a$, alors $I \models \Delta_{\mu}^{d,Max}(\varphi' \sqcup \varphi)$ et $\Delta_{\mu}^{d,Max}(\Delta_{\mu}^{d,Max}(\varphi' \sqcup \varphi) \sqcup \varphi) \vdash \varphi$. Sinon, $\min\{d_{d,Max}(I', \varphi \sqcup \varphi') : I' \models \mu\} = i_0 < a$, alors $Mod(\Delta_{\mu}^{d,Max}(\varphi' \sqcup \varphi)) = \{J \models \mu : d_{d,Max}(J, \varphi \sqcup \varphi') = i_0\}$. Soit φ'' cette base de

connaissance. Notons que $\varphi'' \vdash \mu$, donc soit a' la distance entre φ et φ'' . Notons que $a' \leq i_0 < a$ donc par l'hypothèse d'induction $\exists n \Delta_\mu^{d,Max^n}(\varphi'', \varphi) \vdash \varphi$. Alors $\exists n \Delta_\mu^{d,Max^n}(\varphi', \varphi) \vdash \varphi$.

La preuve que $\Delta^{d,Max}$ satisfait (MI) suit directement de l'équation 7.1. Donc d'après le théorème 72, $\Delta^{d,Max}$ ne satisfait pas (IC6).

Finalement on peut noter que $\Delta_\mu^{d,Max}(\Psi_1 \cup \Psi_2^n) = \Delta_\mu^{d,Max}(\Psi_1 \cup \Psi_2)$. De ce fait et de (IC4), il est facile de voir que (Maj) n'est pas satisfait.

□

Revesz donna l'opérateur $\Delta^{d_H,Max}$ (où d_H est la distance de Dalal) comme exemple à ses opérateurs de model fitting [Rev93].

Il se trouve également que cet opérateur est connu en théorie de la décision sous le nom de règle du *minimax* et a été largement étudié dans ce domaine, notamment par Savage [Sav71].

Cette méthode est également bien connue en économie, en théorie du vote et en théorie du choix social où elle illustre ce que l'on nomme l'approche égalitariste (voir [Mou88]).

7.4 Opérateurs Σ et Σ^n

On définit dans cette section une famille d'opérateurs de fusion majoritaire.

Définition 92 Soient un ensemble de connaissance Ψ , une interprétation I et une distance entre interprétation d . On définit la distance Σ entre une interprétation et un ensemble de connaissance comme :

$$d_{d,\Sigma}(I, \Psi) = \sum_{\varphi \in \Psi} d(I, \varphi)$$

Cela induit le pré-ordre suivant :

$$I \leq_{\Psi}^{d,\Sigma} J \text{ ssi } d_{d,\Sigma}(I, \Psi) \leq d_{d,\Sigma}(J, \Psi)$$

Et l'opérateur $\Delta^{d,\Sigma}$ est défini par :

$$Mod(\Delta^{d,\Sigma}_\mu(\Psi)) = \min(Mod(\mu), \leq_{\Psi}^{d,\Sigma})$$

La définition des opérateurs $\Delta^{d,\Sigma}$ repose sur la distance entre une interprétation et un ensemble de connaissance, définie comme la somme des distances entre cette interprétation et les bases de connaissance de l'ensemble.

Le résultat donné par un tel opérateur $\Delta^{d,\Sigma}$ est en quelque sorte l' "élection" des interprétations les plus populaires parmi celles permises par les contraintes d'intégrité.

Ces opérateurs sont des opérateurs de fusion majoritaire comme montré dans le théorème suivant.

Théorème 90 $\Delta^{d,\Sigma}$ est un opérateur de fusion majoritaire.

Preuve : Tout d'abord on peut remarquer le fait suivant, conséquence directe de la définition :

$$d_{d,\Sigma}(I, \Psi_1 \sqcup \Psi_2) = d_{d,\Sigma}(I, \Psi_1) + d_{d,\Sigma}(I, \Psi_2) \quad (7.2)$$

Nous montrons que l'assignement $\Psi \mapsto \leq_{\Psi}^{d,\Sigma}$ est un assignement syncrétique majoritaire. Puis par le théorème 85 on en conclut que $\Delta_{d,\Sigma}$ satisfait (IC0 – IC8) et (Maj).

Vérifions les conditions de l'assignement syncrétique majoritaire :

1. Si $I \models \Psi$ et $J \models \Psi$, alors $d_{d,\Sigma}(I, \Psi) = 0$ et $d_{d,\Sigma}(J, \Psi) = 0$, donc $I \simeq_{\Psi} J$.
2. Si $I \models \Psi$ et $J \not\models \Psi$, alors $d_{d,\Sigma}(I, \Psi) = 0$ et $d_{d,\Sigma}(J, \Psi) > 0$, donc $I <_{\Psi} J$.
3. Direct.
4. On veut montrer que $\forall I \models \varphi \exists J \models \varphi' J \leq_{\varphi \sqcup \varphi'} I$. On sait que $d(I, \varphi) = 0$ et $d(I, \varphi') = \min_{J \models \varphi'} d(I, J)$, on choisit donc $J \models \varphi'$ tel que $d(I, J) = d(I, \varphi')$. Alors $d(J, \varphi) = \min_{I' \models \varphi} d(J, I') \leq d(J, I)$, et $d(J, \varphi') = 0$. D'où $d_{d,\Sigma}(J, \varphi \sqcup \varphi') = d_{d,\Sigma}(J, \varphi) \leq d_{d,\Sigma}(I, \varphi') = d_{d,\Sigma}(I, \varphi \sqcup \varphi')$. Et par définition $J \leq_{\varphi \sqcup \varphi'} I$.
5. Si $I \leq_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors $d_{d,\Sigma}(I, \Psi_1) \leq d_{d,\Sigma}(J, \Psi_1)$ et $d_{d,\Sigma}(I, \Psi_2) \leq d_{d,\Sigma}(J, \Psi_2)$, donc d'après le fait 7.2 $d_{d,\Sigma}(I, \Psi_1 \sqcup \Psi_2) \leq d_{d,\Sigma}(J, \Psi_1 \sqcup \Psi_2)$.
6. Se déduit du fait 7.2 comme la condition précédente.
7. Si $I <_{\Psi_2} J$, alors $d_{d,\Sigma}(I, \Psi_2) < d_{d,\Sigma}(J, \Psi_2)$. On veut montrer $\exists n I <_{\Psi_1 \sqcup \Psi_2^n} J$, c'est-à-dire

$$\exists n \ d_{d,\Sigma}(I, \Psi_1) + n * d_{d,\Sigma}(I, \Psi_2) < d_{d,\Sigma}(J, \Psi_1) + n * d_{d,\Sigma}(J, \Psi_2)$$

il suffit donc de choisir $n > \frac{d_{d,\Sigma}(I, \Psi_1) - d_{d,\Sigma}(J, \Psi_1)}{d_{d,\Sigma}(J, \Psi_2) - d_{d,\Sigma}(I, \Psi_2)}$.

□

Théorème 91 Si la distance $d : \mathcal{W} \times \mathcal{W} \mapsto \mathbb{N}$ est topologique alors $\Delta^{d,\Sigma}$ satisfait (IC_{it}).

Preuve : Tout d'abord supposons que $\varphi \vdash \mu$ et $\varphi' \vdash \mu$. Nous montrons que $\exists n \ \Delta^{d,\Sigma}_{\mu}(\varphi', \varphi) \vdash \varphi$. Soit a la distance entre φ et φ' , i.e. $d(\varphi, \varphi') = a$. Prenons $I \models \varphi$ et $J \models \varphi'$ tels que $dist(I, J) = a$. En utilisant l'inégalité triangulaire, il est facile de voir que $a = \min\{d_{d,\Sigma}(I', \varphi \sqcup \varphi') : I' \models \mu\}$ d'où $I \models \Delta^{d,\Sigma}_{\mu}(\varphi \sqcup \varphi')$ et alors par (IC2) $\Delta^{d,\Sigma}_{\mu}(\Delta^{d,\Sigma}_{\mu}(\varphi' \sqcup \varphi) \sqcup \varphi) \vdash \varphi$. De ce fait $\exists n \ \Delta^{d,\Sigma}_{\mu}(\varphi', \varphi) \vdash \varphi$. Et donc, en prenant $\varphi' = \Delta^{d,\Sigma}_{\mu}(\Psi \sqcup \varphi)$, on obtient (IC_{it}).

□

Lin et Mendelzon ont donné l'opérateur $\Delta^{d_H, \Sigma}$ (où d_H la distance de Dalal) comme exemple de leurs opérateurs de *theory merging by majority* dans [LM98], et en ont proposé une méthode de calcul syntaxique lorsque les bases de connaissance sont exprimées en forme normale disjonctive (DNF).

Et indépendamment, Revesz donna le même opérateur comme exemple de model fitting pondéré dans [Rev93].

Cette méthode est très proche de la règle de Borda dans le domaine de la théorie du vote et illustre ce que l'on appelle l'approche utilitariste en théorie du choix social [Kel78, Arr63].

Mais la famille $\Delta^{d,\Sigma}$ n'est qu'un cas particulier d'une famille plus générale :

Définition 93 Soient un ensemble de connaissance Ψ , une interprétation I et une distance entre interprétations s . On définit la distance Σ^n entre une interprétation et un ensemble de connaissance comme :

$$d_{d,\Sigma^n}(I, \Psi) = \sqrt[n]{\sum_{\varphi \in \Psi} (d(I, \varphi))^n}$$

Cela induit le pré-ordre suivant :

$$I \leq_{\Psi}^{d,\Sigma^n} J \text{ ssi } d_{d,\Sigma^n}(I, \Psi) \leq d_{d,\Sigma^n}(J, \Psi)$$

Et l'opérateur Δ^{d,Σ^n} est défini par :

$$Mod(\Delta_{\mu}^{d,\Sigma^n}(\Psi)) = \min(Mod(\mu), \leq_{\Psi}^{d,\Sigma^n})$$

Les opérateurs Δ^{d,Σ^n} sont également des opérateurs de fusion majoritaire mais la puissance ajoutée permet de modifier notablement le comportement de ces opérateurs (voir section 7.7).

7.5 Opérateurs GMax

On définit dans cette section une nouvelle méthode de fusion que l'on a nommée $\Delta^{d,GMax}$ (pour *Max Généralisé*).

Le but est de capturer le comportement d' "arbitrage" de la famille $\Delta^{d,Max}$ mais sans perdre les propriétés des opérateurs de fusion contrainte.

Définition 94 Soit un ensemble de connaissance $\Psi = \{\varphi_1 \dots \varphi_n\}$. Pour chaque interprétation I on construit la liste $(d_1^I \dots d_n^I)$ des distances entre cette interprétation et les n bases de connaissance de Ψ , i.e. $d_j^I = d(I, \varphi_j)$. Soit L_I^{Ψ} la liste obtenue à partir de $(d_1^I \dots d_n^I)$ en la triant dans l'ordre décroissant. On définit $d_{d,GMax}(I, \Psi) = L_I^{\Psi}$. Soit $\leq_{lex}$ l'ordre lexicographique entre les séquences d'entiers.

On définit le pré-ordre suivant :

$$I \leq_{\Psi}^{d,GMax} J \text{ ssi } d_{d,GMax}(I, \Psi) \leq_{lex} d_{d,GMax}(J, \Psi)$$

Et l'opérateur $\Delta^{d,GMax}$ est défini par :

$$Mod(\Delta_{\mu}^{d,GMax}(\Psi)) = \min(Mod(\mu), \leq_{\Psi}^{d,GMax}).$$

Par définition il est facile de voir que les opérateurs $\Delta^{d,GMax}$ sont un raffinement des opérateurs $\Delta^{d,Max}$.

Remarque 6 $\Delta_{\mu}^{d,GMax}(\Psi) \vdash \Delta_{\mu}^{d,Max}(\Psi)$

On donne quelques résultats préliminaires pour faciliter les preuves du théorème 94.

Définition 95 Soient deux listes d'entiers L_1 and L_2 . On définit $L_1 \odot L_2$ comme la liste résultat du tri par ordre décroissant de la concaténation des listes L_1 et L_2 .

Lemme 92 Soient L_1, L'_1, L_2, L'_2 , quatre listes d'entiers triés dans l'ordre décroissant et telles que $\text{card}(L_1) = \text{card}(L'_1)$ et $\text{card}(L_2) = \text{card}(L'_2)$. Si $L_1 \leq_{\text{lex}} L'_1$ et $L_2 \leq_{\text{lex}} L'_2$, alors $L_1 \odot L_2 \leq_{\text{lex}} L'_1 \odot L'_2$.

Preuve : Supposons que $L_1 \leq_{\text{lex}} L'_1$ et $L_2 \leq_{\text{lex}} L'_2$. Il est facile de montrer les inégalités suivantes : $L_1 \odot L_2 \leq_{\text{lex}} L'_1 \odot L_2$ et $L'_1 \odot L_2 \leq_{\text{lex}} L'_1 \odot L'_2$. Et par transitivité $L_1 \odot L_2 \leq_{\text{lex}} L'_1 \odot L'_2$.

□

Lemme 93 Soient L_1, L'_1, L_2, L'_2 , quatre listes d'entiers triés dans l'ordre décroissant et telles que $\text{card}(L_1) = \text{card}(L'_1)$ et $\text{card}(L_2) = \text{card}(L'_2)$. Si $L_1 \leq_{\text{lex}} L'_1$ et $L_2 <_{\text{lex}} L'_2$, alors $L_1 \odot L_2 <_{\text{lex}} L'_1 \odot L'_2$.

Preuve : Avec les hypothèses il est facile de voir que $L_1 \odot L_2 \leq_{\text{lex}} L'_1 \odot L_2$ et $L'_1 \odot L_2 <_{\text{lex}} L'_1 \odot L'_2$. On conclut par transitivité de $\leq_{\text{lex}}$.

□

Théorème 94 $\Delta^{d, \text{GMax}}$ est un opérateur d'arbitrage et satisfait (IC_{it}).

Preuve :

Pour montrer que $\Delta^{d, \text{GMax}}$ satisfait les postulats (IC0 - IC8) et (Arb), on utilise le théorème de représentation et on montre que l'assignement $\Psi \mapsto \leq_{\Psi}^{d, \text{GMax}}$ est un assignement synchrétique juste.

1. Si $I \models \Psi$ et $J \models \Psi$, alors $\forall \varphi_i \in \Psi \ I \models \varphi_i$ et $J \models \varphi_i$, alors $L_I = (0, \dots, 0)$ et $L_J = (0, \dots, 0)$, donc $I \simeq_{\Psi} J$.
2. Si $I \models \Psi$ et $J \not\models \Psi$, alors $L_I = (0, \dots, 0)$ et $L_J \neq (0, \dots, 0)$, donc $I <_{\Psi} J$.
3. Si $\Psi_1 \leftrightarrow \Psi_2$, alors il est clair que $\leq_{\Psi_1} = \leq_{\Psi_2}$.
4. On veut montrer $\forall I \models \varphi \ \exists J \models \varphi' \ J \leq_{\varphi \sqcup \varphi'} I$. On sait que $d(I, \varphi) = 0$ et $d(I, \varphi') = \min_{J \models \varphi'} d(I, J)$, on choisit donc $J \models \varphi'$ tel que $d(I, J) = d(I, \varphi')$. Alors $d(J, \varphi) = \min_{I' \models \varphi} d(J, I') \leq d(J, I)$, et $d(J, \varphi') = 0$. Donc $d_{d, \text{GMax}}(J, \varphi \sqcup \varphi') \leq_{\text{lex}} d_{d, \text{GMax}}(I, \varphi \sqcup \varphi')$. D'où par définition $J \leq_{\varphi \sqcup \varphi'} I$.
5. Si $L_I^{\Psi_1} \leq_{\text{lex}} L_J^{\Psi_1}$ et $L_I^{\Psi_2} \leq_{\text{lex}} L_J^{\Psi_2}$. D'après le lemme 92, on a $L_I^{\Psi_1 \sqcup \Psi_2} \leq_{\text{lex}} L_J^{\Psi_1 \sqcup \Psi_2}$.
6. Si $L_J^{\Psi_1} <_{\text{lex}} L_I^{\Psi_1}$ et $L_J^{\Psi_2} \leq_{\text{lex}} L_I^{\Psi_2}$, du lemme 93 s'ensuit $L_J^{\Psi_1 \sqcup \Psi_2} <_{\text{lex}} L_I^{\Psi_1 \sqcup \Psi_2}$.
8. Si $L_I^{\varphi_1} <_{\text{lex}} L_J^{\varphi_1}$, alors $d(I, \varphi_1) < d(J, \varphi_1)$ et de la même façon $d(I, \varphi_2) < d(J, \varphi_2)$. De plus comme $L_J^{\varphi_1 \sqcup \varphi_2} \simeq_{\text{lex}} L_{J'}^{\varphi_1 \sqcup \varphi_2}$ c'est-à-dire soit $d(J, \varphi_1) = d(J', \varphi_1)$ et $d(J, \varphi_2) = d(J', \varphi_2)$, soit $d(J, \varphi_1) = d(J', \varphi_2)$ et $d(J, \varphi_2) = d(J', \varphi_1)$. Dans le premier cas puisque $L_I^{\varphi_1} <_{\text{lex}} L_J^{\varphi_1}$ et $L_I^{\varphi_2} <_{\text{lex}} L_J^{\varphi_2}$ on a par le lemme 93 que $L_I^{\varphi_1 \sqcup \varphi_2} <_{\text{lex}} L_J^{\varphi_1 \sqcup \varphi_2}$. Dans le second cas $d(I, \varphi_1) < d(J, \varphi_1)$ et $d(I, \varphi_2) < d(J, \varphi_1)$, donc $L_I^{\varphi_1 \sqcup \varphi_2} <_{\text{lex}} L_J^{\varphi_1 \sqcup \varphi_2}$.

La preuve de (IC_{it}) est similaire à celle pour $\Delta^{d,Max}$.

□

Il s'avère que cette idée d'améliorer l'opérateur $\Delta^{d,Max}$ a également été proposée en théorie du choix social [Mou88] où cette méthode est appelée fonction *leximin*. Cela a été également proposé en théorie de la décision, voir par exemple [DFP96].

7.6 Exemples

De façon à illustrer le comportement des trois familles d'opérateurs définis précédemment, nous donnons dans cette section deux exemples concrets de fusion.

Nous utiliserons comme distance entre interprétations pour les trois opérateurs la distance de Dalal [Dal88a], notée d_H .

Le premier exemple a été donné dans [Rev97] : un cours de bases de données est suivi par trois étudiants : $E = \{\varphi_1, \varphi_2, \varphi_3\}$. Le professeur peut enseigner SQL, Datalog et O_2 . Il demande l'avis de ses étudiants pour pouvoir satisfaire la classe du mieux possible. Le premier étudiant veut étudier SQL ou O_2 : $\varphi_1 = (S \vee O) \wedge \neg D$. Le second veut apprendre Datalog ou O_2 mais pas les deux : $\varphi_2 = (\neg S \wedge D \wedge \neg O) \vee (\neg S \wedge \neg D \wedge O)$. Le troisième veut étudier les trois langages : $\varphi_3 = (S \wedge D \wedge O)$. En considérant les variables propositionnelles S , D et O dans cet ordre on a : $Mod(\varphi_1) = \{(1,0,0), (0,0,1), (1,0,1)\}$, $Mod(\varphi_2) = \{(0,1,0), (0,0,1)\}$, $Mod(\varphi_3) = \{(1,1,1)\}$. Le professeur n'impose aucune contrainte *a priori* donc $\mu = \top$.

Que doit enseigner le professeur pour satisfaire au mieux sa classe ?

Le tableau 7.6 contient les distances nécessaires aux calculs pour les trois opérateurs $\Delta^{d,Max}(\Psi)$, $\Delta^{d,\Sigma}(\Psi)$ et $\Delta^{d,GMax}(\Psi)$.

	φ_1	φ_2	φ_3	$d_{d_H,Max}$	$d_{d_H,\Sigma}$	$d_{d_H,GMax}$
(0,0,0)	1	1	3	3	5	(3,1,1)
(0,0,1)	0	0	2	2	2	(2,0,0)
(0,1,0)	2	0	2	2	4	(2,2,2)
(0,1,1)	1	1	1	1	3	(1,1,1)
(1,0,0)	0	2	2	2	4	(2,2,0)
(1,0,1)	0	1	1	1	2	(1,1,0)
(1,1,0)	1	1	1	1	3	(1,1,1)
(1,1,1)	1	2	0	2	3	(2,1,0)

TAB. 7.1 – Distances

Le minimum de la colonne de $dist_{d_H, \Delta^{d,Max}}$ est 1, nous avons donc $Mod(\Delta^{d_H,Max}(E)) = \{(0,1,1), (1,0,1), (1,1,0)\}$, le professeur doit donc enseigner deux des trois langages pour satisfaire au mieux sa classe lorsque l'on choisit $\Delta^{d_H,Max}$ comme opérateur de concertation.

On peut illustrer sur cet exemple la raison pour laquelle cet opérateur n'est pas un opérateur de fusion contrainte. L'interprétation (0,1,1) est aussi bonne que (1,0,1) pour l'opérateur $\Delta^{d_H,Max}$. Mais on peut voir aisément que les deux interprétations sont aussi

bonnes pour $\Psi_1 = \varphi_2 \sqcup \varphi_3$ alors que (0,1,1) est meilleure que (1,0,1) pour $\Psi_2 = \varphi_1$. On peut donc satisfaire davantage φ_1 en choisissant (0,1,1) plutôt que (1,0,1), en laissant φ_2 et φ_3 aussi satisfaits. C'est ce que demande (IC6) et on voit bien ici que $\Delta^{d,Max}$ ne distingue pas ces nuances.

Le minimum de la colonne de $dist_{d_H, \Sigma}$ est 2, nous avons donc $Mod(\Delta^{d_H, \Sigma}(E)) = \{(0,0,1), (1,0,1)\}$, le professeur doit donc enseigner SQL et O_2 ou O_2 seul pour satisfaire au mieux sa classe lorsque l'on choisit $\Delta^{d_H, \Sigma}$ comme opérateur.

Le minimum de la colonne $dist_{d_H, GMax}$ est (1,1,0) et nous avons $Mod(\Delta^{d_H, GMax}(E)) = \{(1,0,1)\}$, donc le professeur doit enseigner SQL et O_2 pour satisfaire sa classe lorsque l'opérateur de concertation est $\Delta^{d_H, GMax}$. On voit bien sur cet exemple comment les opérateurs $\Delta^{d, GMax}$ précisent les choix des opérateurs $\Delta^{d, Max}$.

Sur ce premier exemple on a supposé pour simplifier qu'il n'y avait pas de contraintes d'intégrité pour le résultat de la fusion. L'exemple suivant illustre le cas où la fusion est contrainte.

A une réunion de co-propriétaires d'une résidence, le président propose pour l'année à venir la construction d'une piscine, d'un court de tennis et d'un parking privé. On notera respectivement S, T, P la construction de la piscine, du court de tennis et du parking. I dénotera l'augmentation du loyer. Le président souligne le fait que si deux des trois items sont construits le loyer augmentera significativement : $\mu = ((S \wedge T) \vee (S \wedge P) \vee (T \wedge P)) \rightarrow I$.

Il y a quatre co-propriétaires $\Psi = \varphi_1 \sqcup \varphi_2 \sqcup \varphi_3 \sqcup \varphi_4$. Deux d'entre eux veulent construire les trois items et ne se soucient pas de l'augmentation de loyer : $\varphi_1 = \varphi_2 = S \wedge T \wedge P$. Le troisième pense que construire la moindre chose se répercutera inexorablement un jour sur les loyers et ne tient absolument pas à voir son loyer augmenter, il est donc opposé à toute construction : $\varphi_3 = \neg S \wedge \neg T \wedge \neg P \wedge \neg I$. Le dernier trouve que la résidence a réellement besoin d'un court de tennis et d'un parking privé mais ne voudrait pas subir une forte augmentation de loyer : $\varphi_4 = T \wedge P \wedge \neg I$.

On considérera les quatre variables propositionnelles S, T, P, I dans cet ordre pour les interprétations :

$$Mod(\mu) = \mathcal{W} \setminus \{(0,1,1,0), (1,0,1,0), (1,1,0,0), (1,1,1,0)\}$$

$$Mod(\varphi_1) = \{(1,1,1,1), (1,1,1,0)\}$$

$$Mod(\varphi_2) = \{(1,1,1,1), (1,1,1,0)\}$$

$$Mod(\varphi_3) = \{(0,0,0,0)\}$$

$$Mod(\varphi_4) = \{(1,1,1,0), (0,1,1,0)\}$$

Les calculs sont résumés dans le tableau 7.6, pour chaque interprétation on donne la distance entre celle-ci et les quatre bases de connaissance et la distance entre cette interprétation et l'ensemble de connaissance selon les 3 opérateurs que l'on a défini $\Delta^{d_H, Max}$, $\Delta^{d_H, \Sigma}$ et $\Delta^{d_H, GMax}$. Les lignes grisées correspondent aux interprétations rejetées par les contraintes d'intégrité. Le résultat de la fusion doit donc être cherché parmi les interprétations non grisées.

Avec $\Delta^{d_H, Max}$ comme opérateur de fusion, la distance minimum entre une interprétation

	φ_1	φ_2	φ_3	φ_4	$\text{dist}_{d_H, \text{Max}}$	$\text{dist}_{d_H, \Sigma}$	$\text{dist}_{d_H, \text{GMax}}$
(0,0,0,0)	3	3	0	2	3	8	(3,3,2,0)
(0,0,0,1)	3	3	1	3	3	10	(3,3,3,1)
(0,0,1,0)	2	2	1	1	2	6	(2,2,1,1)
(0,0,1,1)	2	2	2	2	2	8	(2,2,2,2)
(0,1,0,0)	2	2	1	1	2	6	(2,2,1,1)
(0,1,0,1)	2	2	2	2	2	8	(2,2,2,2)
(0,1,1,0)	1	1	2	0	2	4	(2,1,1,0)
(0,1,1,1)	1	1	3	1	3	6	(3,1,1,1)
(1,0,0,0)	2	2	1	2	2	7	(2,2,2,1)
(1,0,0,1)	2	2	2	3	3	9	(3,2,2,2)
(1,0,1,0)	1	1	2	1	2	5	(2,1,1,1)
(1,0,1,1)	1	1	3	2	3	7	(3,2,1,1)
(1,1,0,0)	1	1	2	1	2	5	(2,1,1,1)
(1,1,0,1)	1	1	3	2	3	7	(3,2,1,1)
(1,1,1,0)	0	0	3	0	3	3	(3,0,0,0)
(1,1,1,1)	0	0	4	1	4	5	(4,1,0,0)

TAB. 7.2 – Distances

et l'ensemble de connaissance est 2, et les interprétations suivantes sont donc retenues : $\text{Mod}(\Delta_{\mu}^{d_H, \text{Max}}(\Psi)) = \{(0,0,1,0), (0,0,1,1), (0,1,0,0), (0,1,0,1), (1,0,0,0)\}$. La décision qui est conforme aux vœux du groupe est alors de ne pas augmenter le loyer et de construire l'un des trois items, ou d'augmenter le loyer et de construire soit le court de tennis, soit le parking privé.

Comme nous l'avons vu précédemment, l'opérateur $\Delta^{d_H, \text{Max}}$ n'est pas un opérateur de fusion contrainte parce qu'il n'est pas assez précis. L'opérateur $\Delta^{d_H, \text{GMax}}$ précise les choix de $\Delta^{d_H, \text{Max}}$. Avec cet opérateur on a comme résultat $\text{Mod}(\Delta_{\mu}^{d_H, \text{GMax}}(\Psi)) = \{(0,0,1,0), (0,1,0,0)\}$, la décision prise dans ce cas sera donc de construire soit le court de tennis, soit le parking et de ne pas augmenter le loyer.

En revanche, si on choisit $\Delta^{d_H, \Sigma}$ pour résoudre le conflit en se rangeant au vœux de la majorité, le résultat est alors $\text{Mod}(\Delta_{\mu}^{d_H, \Sigma}(\Psi)) = \{(1,1,1,1)\}$, et la solution adoptée est de construire les trois items et d'augmenter le loyer.

Le vote majoritaire semble plus démocratique que les autres méthodes mais par exemple, dans ce cas, cela ne marche que si φ_3 accepte de se conformer à cette décision qui va complètement à l'encontre de ses vœux. Il se pourrait très bien que, fâché de cette décision, il décide de ne pas payer son augmentation de loyer et aucun des trois items ne pourrait être construit.

Donc si une décision, comme ici ou comme lors d'un accord commercial ou d'un accord de paix, nécessite l'adhésion de tous les participants, un opérateur d'arbitrage semble plus adéquat qu'un opérateur majoritaire.

7.7 Opérateurs d'arbitrage et opérateurs majoritaires

Nous montrons dans cette section qu'il est possible d'être à la fois un opérateur d'arbitrage et un opérateur majoritaire. Nous définissons tout d'abord une distance drastique entre interprétations. Les opérateurs $\Delta^{d, GMax}$ et $\Delta^{d, \Sigma}$ définis à partir de cette distance coïncident. Ensuite, nous montrons que, quelle que soit la distance choisie, certains opérateurs Δ^{d, Σ^n} sont à la fois des opérateurs d'arbitrage et des opérateurs majoritaires.

7.7.1 Distance drastique

La distance la plus simple que l'on peut définir entre deux interprétations est celle donnant 0 si les deux interprétations sont égales et 1 sinon.

$$d_D(I, J) = \begin{cases} 0 & \text{si } I = J \\ 1 & \text{sinon} \end{cases}$$

La distance entre une interprétation et une base de connaissance est alors également 0 ou 1 suivant que l'interprétation satisfait ou non la base de connaissance.

Il est facile de voir alors que les opérateurs obtenus à partir des familles $\Delta^{d, GMax}$ et $\Delta^{d, \Sigma}$ et de cette distance coïncident. On a donc le résultat suivant, conséquence des théorèmes 90 et 94 :

Théorème 95 *L'opérateur $\Delta^{d_D, \Sigma} = \Delta^{d_D, GMax}$ satisfait les postulats (IC0)-(IC8), (Maj), (Arb) et (IC_{it}).*

Un aspect intéressant de cet opérateur est qu'il est quasiment syntaxique, puisque l'on calcule la distance d'une interprétation à une base de connaissance par un test de satisfiabilité. La distance obtenue étant la plus simple possible, il n'a qu'un comportement très simple mais cette définition peut sembler moins arbitraire que celles qui utilisent une distance plus évoluée comme la distance de Dalal par exemple.

7.7.2 Etude graphique

On a illustré sur des exemples les points communs et les différences entre les trois méthodes de fusion $\Delta^{d, Max}$, $\Delta^{d, \Sigma}$ et $\Delta^{d, GMax}$. On peut pousser l'investigation un peu plus loin en examinant graphiquement le comportement de ces opérateurs explorant l'espace des solutions.

Pour que la représentation soit simple, on se limitera dans cette section à la fusion de deux bases de connaissance.

On place les interprétations dans le plan avec comme abscisse leur distance à la base φ_2 et comme ordonnée leur distance à la base φ_1 . Ainsi, le but de la fusion est de déterminer l'ensemble des interprétations les plus proches du point (0,0). La différence entre les différents opérateurs de fusion réside dans la définition de la "distance" utilisée.

Sur le graphique 7.1 on a représenté les courbes qui dénotent les interprétations à une distance 3 de l'ensemble de connaissance selon les opérateurs $\Delta^{d, Max}$, $\Delta^{d, \Sigma}$ et Δ^{d, Σ^2} . $\Delta^{d, Max}$ est représenté par un carré de côté a , $\Delta^{d, \Sigma}$ par une droite d'équation $x = a - y$ et Δ^{d, Σ^2} par un arc de cercle de rayon a , a étant la distance par rapport à l'ensemble de connaissance.

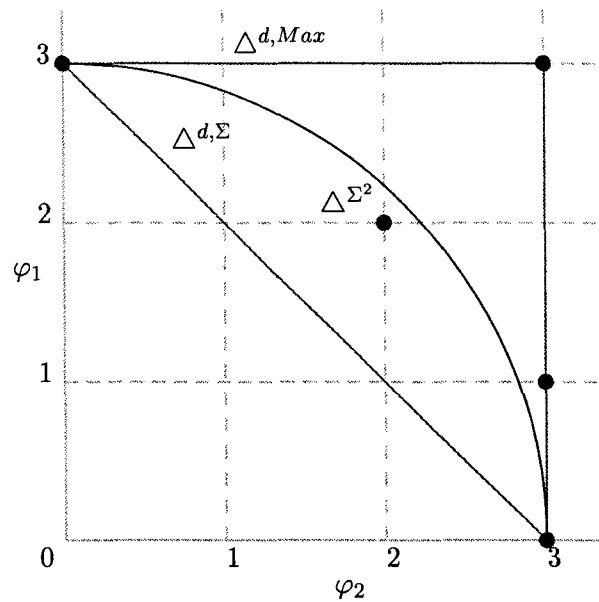


FIG. 7.1 – Fusion de deux bases de connaissance

L'opérateur $\Delta^{d,GMax}$ est difficilement représentable graphiquement mais il faut imaginer une courbe qui suit celle de $\Delta^{d,Max}$ mais préférant les interprétations proches des axes. Nous verrons ensuite comment approximer graphiquement l'opérateur $\Delta^{d,GMax}$.

Ainsi, le résultat de la fusion pour ces trois opérateurs est l'ensemble des interprétations rencontrées en premier par ces courbes lorsque l'on fait varier a de 0 à l'infini.

Sur cet exemple, le résultat pour $\Delta^{d,Max}$ et Δ^{d,Σ^2} sera l'interprétation placée en (2,2) et pour $\Delta^{d,\Sigma}$ le résultat sera les interprétations placées en (3,0) et (0,3).

De la même manière on peut reconstruire les pré-ordres $\leq_{\Psi}^{d,Max}$, $\leq_{\Psi}^{d,\Sigma}$ et $\leq_{\Psi}^{d,\Sigma^2}$ en considérant l'ordre dans lequel sont rencontrées les interprétations.

Sur le graphique, on peut se rendre compte de l'insuffisance de $\Delta^{d,Max}$ qui ne permet pas de faire la distinction entre les points (3,0) et (3,3). C'est pour cette raison que $\Delta^{d,Max}$ n'est pas un opérateur de fusion contrainte.

D'un autre côté, $\Delta^{d,\Sigma}$ ne fait aucune distinction sur l'origine du mécontentement. En effet, la distance de l'interprétation à l'ensemble de connaissance est une mesure du mécontentement qu'engendre cette interprétation sur l'ensemble de connaissance. Et l'opérateur $\Delta^{d,\Sigma}$ n'est absolument pas consensuel car il permet de choisir des interprétations satisfaisant totalement l'une des deux bases (par exemple celle située en (3,0)), alors que d'autres interprétations seraient plus "égalitaires" (comme par exemple celle située en (2,2)).

Cela peut sembler normal pour un opérateur majoritaire mais, contrairement à ce que l'on pourrait penser, ce n'est pas systématique. En effet, les opérateurs Δ^{d,Σ^n} avec $n > 1$ préféreront les choix consensuels, c'est-à-dire ceux situés à proximité de la droite $x = y$. Et

donc, sur la figure 7.1 l'interprétation située en (2,2) sera préférée à celle située en (3,0).

L'opérateur Δ^{d,Σ^2} est un représentant particulier de la classe des opérateurs Δ^{d,Σ^n} puisqu'il utilise la distance euclidienne pour calculer les distances entre les interprétations et l'ensemble de connaissance. Cette volonté d'être proche de l'ensemble de connaissance peut justifier l'utilisation de l'opérateur Δ^{d,Σ^2} puisqu'il donne une distance sphérique assez naturelle, qui est majoritaire sans avoir les excès de $\Delta^{d,\Sigma}$.

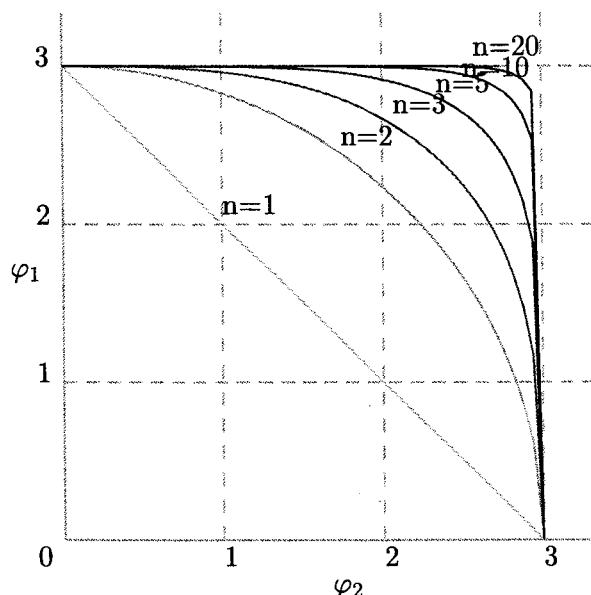


FIG. 7.2 – La famille Δ^{d,Σ^n}

De plus, on peut remarquer sur le graphique 7.2 que, lorsque l'on augmente la valeur de n , la courbe de $\Delta^{d,\Sigma}$ s'approche de celle de $\Delta^{d,Max}$. Donc, à partir d'un n suffisamment grand, on peut prendre $\Delta^{d,Max}$ comme approximation de la courbe de Δ^{d,Σ^n} . Mais, quelle que soit la valeur de n , une interprétation placée en (x,y) sera toujours préférée à une interprétation placée en $(x,y+1)$ ou en $(x+1,y)$. Mais ce parcours de la courbe $\Delta^{d,Max}$, en préférant les interprétations les plus proches des axes, est exactement celui de la courbe de $\Delta^{d,GMax}$. Donc, à partir d'un certain n , $\Delta^{d,\Sigma^n} = \Delta^{d,GMax}$. Plus formellement, on a le résultat suivant :

Théorème 96 *Lorsque le langage est fini, si l'on fusionne un ensemble de connaissance Ψ :*

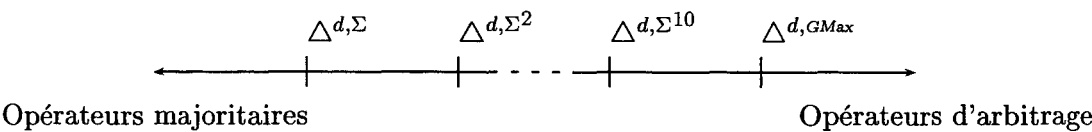
$$\lim_{n \rightarrow \infty} \Delta^{d,\Sigma^n}(\Psi) = \Delta^{d,GMax}(\Psi)$$

Preuve : Nous montrons ce résultat pour deux bases, mais il se généralise directement pour un nombre quelconque (fini) de bases. On montre qu'à partir d'un certain n_0 les pré-ordres $\leq_{\Psi}^{n_0}$ et $\leq_{\Psi}^{GMax}$ coïncident. Il est facile de montrer qu'une interprétation placée en (x,y) sera toujours préférée à une interprétation placée en $(x,y+1)$, puisque $\forall n \ x^n + y^n < x^n + (y+1)^n$. De même, symétriquement une interprétation placée en (x,y) sera toujours préférée à une interprétation placée en $(x+1,y)$. Pour montrer que les pré-ordres coïncident il ne reste donc plus qu'à montrer qu'à partir d'un certain n_0 une interprétation placée en (x,y) avec $x \geq y$ sera préférée à une interprétation placée en $(x+1,0)$. C'est-à-dire montrer que

$\forall n > n_0 \ 2x^n < (x + 1)^n$. Pour vérifier cette inégalité il suffit de choisir $n_0 \geq x$. Ce qui est possible lorsque la distance maximale est finie.

□

Ce résultat est une autre solution au problème de la partition opérateurs d'arbitrage - opérateurs majoritaires. Puisque les opérateurs Δ^{d,Σ^n} , pour tout n supérieur à un certain n_0 fixé par la distance maximale entre une interprétation et une base de connaissance, sont à la fois des opérateurs d'arbitrage et des opérateurs majoritaires. Ces deux ensembles ne sont donc pas disjoints. La frontière entre majorité et arbitrage est assez floue et il est possible, dans un sens, de passer continûment de l'un à l'autre.



Chapitre 8

Fusion contrainte et autres travaux

Dans ce chapitre nous comparons le cadre de la fusion contrainte avec d'autres travaux. Nous montrons que les opérateurs de fusion contrainte sont une généralisation des opérateurs de révision AGM. Nous étudions également comment définir un opérateur de fusion à partir d'un opérateur de révision. Un résultat intéressant est que l'opérateur de révision doit alors être défini à partir d'une distance.

Nous comparons ensuite le cadre de la fusion contrainte et celui de la fusion pure (i.e. sans contrainte d'intégrité) et soulignons le fait que, dans ce dernier cas, il est beaucoup moins facile de définir un théorème de représentation.

Nous montrons qu'il est possible de définir les opérateurs de révision commutative de Liberatore et Schaerf à partir des opérateurs de fusion contrainte. Cette définition en termes d'opérateur de fusion contrainte est intéressante car elle permet de généraliser, de manière naturelle, les opérateurs de Liberatore et Schaerf pour permettre la fusion d'un nombre quelconque de bases de connaissance.

Nous parlons pour finir des rapports entre fusion de bases de connaissance et théorie du choix social. La fusion de bases de connaissance a pour objet la fusion de connaissances et la théorie du choix social, la fusion de préférences. Les théorèmes de représentation pour la fusion construisants de tels ordres de préférences à partir des connaissances pour réaliser la fusion, il est intéressant de comparer ces deux approches. En particulier, il existe en théorie du choix social un théorème d'impossibilité exprimant le fait qu'il n'existe pas de "bonne" fonction d'agrégation de préférence. Nous voyons donc comment les opérateurs de fusion "échappent" à ce résultat.

8.1 Fusion contrainte et révision AGM

Nous montrons dans cette section que les opérateurs de fusion contrainte sont une généralisation des opérateurs de révision AGM à plusieurs bases de connaissance.

Nous montrons également comment construire des opérateurs de fusion contrainte à partir d'opérateurs de révision. Nous montrons que les opérateurs de révision donnant des opérateurs de fusion avec de bonnes propriétés sont ceux définis à partir d'une distance.

Intuitivement, en paraphrasant les théorèmes de représentation, les opérateurs de révision sélectionnent dans une formule (la nouvelle information), l'information la plus proche d'une information de base (l'ancienne base de connaissance). Identiquement, les opérateurs de fusion contrainte sélectionnent dans une formule (les contraintes d'intégrité), l'information la plus proche d'une information de base (un multi-ensemble de bases de connaissance). En suivant cette remarque, il est facile de donner la correspondance suivante entre opérateurs de fusion contrainte et opérateurs de révision (la preuve est triviale) :

Théorème 97 *Si Δ est un opérateur de fusion contrainte, alors l'opérateur $\circ$, défini par*

$$\varphi \circ \mu = \Delta_\mu(\varphi)$$

est un opérateur de révision AGM.

On dit alors que $\circ$ est l'opérateur de révision associé à Δ .

Inversement, on peut se demander s'il est possible de construire un opérateur de fusion contrainte à partir d'un opérateur de révision AGM donné.

Tout d'abord, il est important de noter qu'un opérateur de révision seul n'est pas suffisant pour définir un opérateur de fusion puisqu'il ne fournit pas d'information sur la manière de combiner les préférences (croyances) individuelles et qu'il n'y a pas de méthode canonique de le faire.

En fait, un opérateur de révision fournit un moyen de déterminer les relations de préférences individuelles associées à chaque base de connaissance, c'est-à-dire les pré-ordres donnés par l'assignement fidèle. Mais la définition d'un opérateur de fusion nécessite également une méthode d'agrégation de ces préférences individuelles en une relation de préférence globale.

Il est donc nécessaire d'examiner les propriétés de couples (opérateur de révision, méthode d'agrégation).

Cette méthode d'agrégation peut, par exemple, être une méthode à la $\Delta^{d,\Sigma}$ ou $\Delta^{d,GMax}$.

Nous proposons la construction suivante à partir d'un opérateur de révision $\circ$ et d'une méthode d'agrégation donnés :

Définition 96

- On définit $f_\varphi^\circ(I) = n$ où n est le niveau où l'interprétation I apparaît dans le pré-ordre $\leq_\varphi$. Plus formellement n est la taille de la plus longue chaîne d'inégalités strictes $I_0 < \dots < I_n$ avec $I_0 \models \varphi$ et $I_n = I$.
- On définit $f_\Psi^\circ(I)$ avec la méthode d'agrégation choisie (par exemple $f_\Psi^\circ(I) = \sum_{\varphi \in \Psi} (f_\varphi^\circ(I))$ si la méthode choisie est Σ).
- On définit $I \leq_\Psi J$ ssi $f_\Psi^\circ(I) \leq f_\Psi^\circ(J)$.
- Finalement $Mod(\Delta_\mu^\circ(\Psi)) = \min(Mod(\mu), \leq_\Psi)$.

La question est alors de trouver quelles sont les conditions sur l'opérateur de révision et sur la méthode d'agrégation pour assurer de bonnes propriétés à l'opérateur de fusion ainsi défini.

Par exemple, si l'on choisit $f_{\Psi}^{\circ}(I) = \sum_{\varphi \in \Psi} (f_{\varphi}^{\circ}(I))$, c'est-à-dire une méthode de fusion à la Σ et sans aucune propriété supplémentaire sur $\circ$, on a le résultat suivant :

Théorème 98 *Si un opérateur de fusion $\Delta^{\circ, \Sigma}$ est défini à partir d'un opérateur de révision $\circ$ et de la méthode d'agrégation $f_{\Psi}^{\circ}(I) = \sum_{\varphi \in \Psi} (f_{\varphi}^{\circ}(I))$, selon la définition 96, alors l'opérateur $\Delta^{\circ, \Sigma}$ satisfait (IC0-IC3), (IC5-IC8) et (Maj).*

Preuve : Nous vérifions les conditions sur l'assignement, et nous concluons par le corollaire 83.

1. Si $I \models \Psi$ et $J \models \Psi$, alors $f_{\Psi}^{\circ}(I) = 0$ et $f_{\Psi}^{\circ}(J) = 0$, donc $I \simeq_{\Psi} J$.
2. Si $I \models \Psi$ et $J \not\models \Psi$, alors $f_{\Psi}^{\circ}(I) = 0$ et $f_{\Psi}^{\circ}(J) > 0$, donc $I <_{\Psi} J$.
3. Direct.
5. Si $I \leq_{\Psi_1} J$ et $I \leq_{\Psi_2} J$, alors $f_{\Psi_1}^{\circ}(I) \leq f_{\Psi_1}^{\circ}(J)$ et $f_{\Psi_2}^{\circ}(I) \leq f_{\Psi_2}^{\circ}(J)$, d'où trivialement $f_{\Psi_1 \sqcup \Psi_2}^{\circ}(I) \leq f_{\Psi_1 \sqcup \Psi_2}^{\circ}(J)$.
6. Direct (similaire à 5).
7. Si $I <_{\Psi_2} J$, alors $f_{\Psi_2}^{\circ}(I) < f_{\Psi_2}^{\circ}(J)$. On désire montrer

$$\exists n \ f_{\Psi_1}^{\circ}(I) + n * f_{\Psi_2}^{\circ}(I) < f_{\Psi_1}^{\circ}(J) + n * f_{\Psi_2}^{\circ}(J)$$

on choisit donc simplement $n > \frac{f_{\Psi_1}^{\circ}(I) - f_{\Psi_1}^{\circ}(J)}{f_{\Psi_2}^{\circ}(J) - f_{\Psi_2}^{\circ}(I)}$.

□

Malheureusement, la condition 4 de l'assignement synchrétique n'est pas toujours satisfaite.

On étend la définition de $f_{\varphi}^{\circ}(\cdot)$ aux bases de connaissance naturellement :

$$f_{\varphi}^{\circ}(\varphi') = \min\{f_{\varphi}^{\circ}(I) : I \models \varphi'\}$$

Théorème 99 *Si un opérateur de fusion $\Delta^{\circ, \Sigma}$ est défini à partir d'un opérateur de révision $\circ$ selon la définition 96 en utilisant la méthode d'agrégation $f_{\Psi}^{\circ}(I) = \sum_{\varphi \in \Psi} f_{\varphi}^{\circ}(I)$, alors l'opérateur*

$\Delta^{\circ, \Sigma}$ est un opérateur de fusion majoritaire si et seulement si la condition suivante de symétrie est satisfaite :

$$(\text{Sym}) \ f_{\varphi}^{\circ}(\varphi') = f_{\varphi'}^{\circ}(\varphi) \quad (\text{symétrie})$$

Preuve : D'après les théorèmes 98 et 85 il reste simplement à montrer que la condition 4 de l'assignement correspond à la propriété (Sym).

On montre tout d'abord que la condition 4 implique $f_{\varphi}^{\circ}(\varphi') = f_{\varphi'}^{\circ}(\varphi)$. La condition 4 de l'assignement est équivalente pour $\Delta^{\circ, \Sigma}$ à $\forall I \models \varphi \ \exists J \models \varphi' \ f_{\varphi}^{\circ}(J) \leq f_{\varphi'}^{\circ}(I)$. De là on déduit aisément $f_{\varphi}^{\circ}(\varphi') \leq f_{\varphi'}^{\circ}(\varphi)$ et comme les rôles de φ et φ' sont symétriques dans la condition 4 on obtient également $f_{\varphi'}^{\circ}(\varphi) \leq f_{\varphi}^{\circ}(\varphi')$. D'où $f_{\varphi}^{\circ}(\varphi') = f_{\varphi'}^{\circ}(\varphi)$.

Inversement, supposons que $f_\varphi^\circ(\varphi') = f_{\varphi'}^\circ(\varphi)$. Et par l'absurde supposons que la condition 4 n'est pas satisfaite, c'est-à-dire $\exists I \models \varphi \forall J \models \varphi' J >_{\varphi \sqcup \varphi'} I$. Alors $\exists I \models \varphi \forall J \models \varphi' f_{\varphi \sqcup \varphi'}^\circ(J) > f_{\varphi \sqcup \varphi'}^\circ(I)$. De là, puisque $f_{\varphi'}^\circ(J) = f_\varphi^\circ(I) = 0$, on a $\exists I \models \varphi \forall J \models \varphi' f_\varphi^\circ(J) > f_{\varphi'}^\circ(I)$. D'où $f_\varphi^\circ(\varphi') > f_{\varphi'}^\circ(\varphi)$. Contradiction.

□

En fait, cette condition de symétrie est également valable pour d'autres méthodes de fusion. En particulier, si l'on définit une méthode à la *GMax*, c'est-à-dire si l'on modifie la définition 96 comme suit

- On définit $f_\Psi^\circ(I)$, la liste des $f_\varphi^\circ(I)$ triés par ordre décroissant, i.e.
 $f_\Psi^\circ(I) = (f_{\varphi_1}^\circ(I), \dots, f_{\varphi_n}^\circ(I))$ avec $\{\varphi_1, \dots, \varphi_n\} \leftrightarrow \Psi$ et $\forall i f_{\varphi_i}^\circ(I) \geq f_{\varphi_{i+1}}^\circ(I)$.
- On définit $I \leq_\Psi J$ ssi $f_\Psi^\circ(I) \leq_{lex} f_\Psi^\circ(J)$.

Pour cette famille d'opérateurs, on peut prouver de manière similaire au théorème 99 le résultat suivant :

Théorème 100 *Si un opérateur de fusion $\Delta^{\circ, GMax}$ est défini à partir d'un opérateur de révision $\circ$ selon la définition 96 en utilisant la méthode d'agrégation *GMax*, alors l'opérateur $\Delta^{\circ, GMax}$ satisfait (IC0-IC3), (IC5-IC8) and (Arb).*

De plus, l'opérateur $\Delta^{\circ, GMax}$ est un opérateur d'arbitrage si et seulement si la condition (Sym) est vérifiée.

En fait, comme nous le montrons ici, la condition (Sym) n'est satisfaite que si et seulement si l'opérateur de révision est défini à partir d'une distance.

Plus précisément on dit qu'un opérateur de révision est défini par une distance d ssi :

- d est une distance.
- soient une base de connaissance φ et une interprétation $I : d(I, \varphi) = \min\{d(I, J) : J \models \varphi\}$
- $I \leq_\varphi J$ ssi $d(I, \varphi) \leq d(J, \varphi)$
- $Mod(\varphi \circ \mu) = \min(Mod(\mu), \leq_\varphi)$

Notons que si $\circ$ est défini à partir d'une distance les méthodes d'agrégations Σ et *GMax* donneront respectivement les opérateurs $\Delta^{d, \Sigma}$ et $\Delta^{d, GMax}$ définis au chapitre 7.

Théorème 101 *Soit $\circ$ un opérateur de révision. Alors la condition (Sym) est satisfaite ssi $\circ$ est défini à partir d'une distance.*

Preuve : La partie **Si** est directe. Pour la partie **Seulement si** on définit la distance suivante : $d(I, J) = f_{\varphi_I}^\circ(J)$. Soit $\varphi \mapsto \leq_\varphi$ l'assignement représentant $\circ$. On veut montrer que $I \leq_\varphi J$ ssi $d(I, \varphi) \leq d(J, \varphi)$. Notons que par définition $d(I, \varphi) \leq d(J, \varphi)$ si et seulement si $\min\{f_{\varphi_{I'}}^\circ(I) : I' \models \varphi\} \leq \min\{f_{\varphi_{J'}}^\circ(J) : J' \models \varphi\}$ ce qui par (Sym) est équivalent à $\min\{f_{\varphi_I}^\circ(I') : I' \models \varphi\} \leq \min\{f_{\varphi_J}^\circ(I') : I' \models \varphi\}$. Ce qui est exactement $f_{\varphi_I}^\circ(\varphi) \leq f_{\varphi_J}^\circ(\varphi)$, ce qui est équivalent, en utilisant (Sym) une nouvelle fois, à $f_\varphi^\circ(I) \leq f_\varphi^\circ(J)$, c'est-à-dire $I \leq_\varphi J$.

□

Donc, comme corollaire des théorèmes 99, 100 et 101 on a :

Théorème 102 *Un opérateur de fusion Δ défini à partir d'un opérateur de révision $\circ$ et de la méthode d'agrégation Σ ou $Gmax$ est un opérateur de fusion contrainte si et seulement si $\circ$ est défini à partir d'une distance.*

En particulier, comme il est clair qu'il y a des opérateurs de révision qui ne sont pas définis à partir d'une distance, les opérateurs Δ° correspondant ne satisferont pas (IC4).

8.2 Fusion contrainte et fusion pure

Dans [KP98, KP97] nous avons étudié les opérateurs de fusion en l'absence de contrainte d'intégrité. Nous appellerons ces opérateurs opérateurs de fusion pure pour faire la distinction. Ce cadre est beaucoup moins expressif que le cas avec contraintes et le théorème de représentation n'est pas très constructif. Ce cadre pose également des problèmes en ce qui concerne la définition des opérateurs d'arbitrage.

8.2.1 Postulats

Dans cette section nous donnons la caractérisation des opérateurs de fusion pure proposée dans [KP98].

Soit un ensemble de connaissance Ψ et Δ un opérateur qui associe à chaque ensemble de connaissance Ψ une base de connaissance $\Delta(\Psi)$.

Définition 97 *Δ est un opérateur de fusion pure si il vérifie les propriétés suivantes :*

- (A1) $\Delta(\Psi)$ est consistant
- (A2) Si Ψ est consistant, alors $\Delta(\Psi) = \bigwedge \Psi$
- (A3) Si $\Psi_1 \leftrightarrow \Psi_2$, alors $\vdash \Delta(\Psi_1) \leftrightarrow \Delta(\Psi_2)$
- (A4) Si $\varphi \wedge \varphi'$ n'est pas consistant, alors $\Delta(\varphi \sqcup \varphi') \not\vdash \varphi$
- (A5) $\Delta(\Psi_1) \wedge \Delta(\Psi_2) \vdash \Delta(\Psi_1 \sqcup \Psi_2)$
- (A6) Si $\Delta(\Psi_1) \wedge \Delta(\Psi_2)$ est consistant, alors $\Delta(\Psi_1 \sqcup \Psi_2) \vdash \Delta(\Psi_1) \wedge \Delta(\Psi_2)$

On peut dès à présent noter que ces propriétés sont très proches de celles des opérateurs de fusion contrainte. Dans ce cadre, les opérateurs de fusion pure majoritaire et d'arbitrage pur ont été définis comme suit :

Un opérateur de *fusion pure majoritaire* est un opérateur de fusion pure qui satisfait le postulat suivant :

$$(M7) \quad \exists n \quad \Delta(\Psi \sqcup \varphi^n) \vdash \varphi$$

Un opérateur d'*arbitrage pur* est un opérateur de fusion pure qui satisfait le postulat suivant :

$$(A7) \quad \forall \varphi' \exists \varphi \quad \varphi' \not\vdash \varphi \quad \forall n \quad \Delta(\varphi' \sqcup \varphi^n) \leftrightarrow \Delta(\varphi' \sqcup \varphi)$$

8.2.2 Caractérisations sémantiques

Dans cette section nous donnons une caractérisation des opérateurs de fusion pure d'abord en termes de fonctions sur les ensembles d'interprétations, puis en termes de famille de pré-ordres sur les interprétations.

Cette caractérisation sémantique en termes de pré-ordres est très proche de la caractérisation logique. Cela est dû au fait que l'on ne peut pas définir le pré-ordre de manière aussi subtile que lorsque l'on autorise des contraintes d'intégrité.

On note $\mathcal{I}$ l'ensemble des ensembles d'interprétations non vides. Et on note $\mathcal{M}$ l'ensemble des multi-ensembles non vides d'éléments de $\mathcal{I}$.

On note $S/\leftrightarrow$ le quotient de S par la relation $\leftrightarrow$. Donc la fonction $\iota : S/\leftrightarrow \rightarrow \mathcal{M}$, définie par $\iota(\{\{\varphi_1, \dots, \varphi_n\}\}_{\leftrightarrow}) = \{Mod(\varphi_1), \dots, Mod(\varphi_n)\}$ est une bijection. Par abus, on écrira $\iota(\Psi)$ au lieu de $\iota(\{\Psi\}_{\leftrightarrow})$.

On définit d'abord ce qu'est une fonction de fusion :

Définition 98 Une fonction $\delta : \mathcal{M} \rightarrow \mathcal{I}$ est appelé fonction de fusion si les propriétés suivantes sont vérifiées : Soient $M, M_1, M_2 \in \mathcal{M}$ et $S, S' \in \mathcal{I}$:

1. Si $I \in \bigcap M$, alors $I \in \delta(M)$
2. Si $\bigcap M \neq \emptyset$ et $I \notin \bigcap M$, alors $I \notin \delta(M)$
3. Si $S \cap S' = \emptyset$, alors $\delta(S \sqcup S') \not\subseteq S$
4. Si $I \in \delta(M_1)$ et $I \in \delta(M_2)$, alors $I \in \delta(M_1 \sqcup M_2)$
5. Si $\delta(M_1) \cap \delta(M_2) \neq \emptyset$ et $I \notin \delta(M_1)$, alors $I \notin \delta(M_1 \sqcup M_2)$

Une fonction de fusion majoritaire est une fonction de fusion qui vérifie :

6. $\forall M \in \mathcal{M} \forall S \in \mathcal{I} \exists n \delta(M \sqcup S^n) \subseteq S$

Une fonction de fusion juste est une fonction de fusion qui vérifie :

7. $\forall S' \in \mathcal{I} \exists S \in \mathcal{I} S' \not\subseteq S \forall n \delta(S' \sqcup S^n) = \delta(S' \sqcup S)$

Il est facile de voir, d'après la bijection ι que les propriétés 1 – 5 sont la contrepartie syntaxique des postulats (A1 – A6) (en notant que le postulat (A₁) correspond au fait que $\emptyset \notin \mathcal{I}$), la propriété 6 correspond au postulat (M7) et la propriété 7 correspond au postulat (A7). Plus précisément, on obtient donc directement le théorème de représentation suivant :

Théorème 103 Un opérateur Δ est un opérateur de fusion pure si et seulement si il existe une fonction de fusion $\delta : \mathcal{M} \rightarrow \mathcal{I}$ telle que

$$Mod(\Delta(\Psi)) = \delta(\iota(\Psi)).$$

De plus, Δ est un opérateur de fusion pure majoritaire ssi δ est une fonction de fusion majoritaire; et Δ est un opérateur d'arbitrage pur ssi δ est une fonction de fusion juste.

On peut également supposer l'existence d'une relation qui représente intuitivement la crédibilité relative des interprétations pour un ensemble de connaissance donné.

On définit alors l'assignement suivant :

Définition 99 *Un assignement synchrétique pur est un assignement qui associe à chaque ensemble de connaissance Ψ un pré-ordre $\leq_\Psi$ sur les interprétations tel que pour tout $\Psi, \Psi_1, \Psi_2 \in \Psi$ et pour tout $\varphi, \varphi' \in \varphi$:*

1. Si $I \in \text{Mod}(\Psi)$ et $J \in \text{Mod}(\Psi)$, alors $I \simeq_\Psi J$
2. Si $I \in \text{Mod}(\Psi)$ et $J \notin \text{Mod}(\Psi)$, alors $I <_\Psi J$
3. Si $\Psi_1 \leftrightarrow \Psi_2$, et $\leq_{\Psi_1} = \leq_{\Psi_2}$
4. Si $\text{Mod}(\varphi) \cap \text{Mod}(\varphi') = \emptyset$, alors $\min(\mathcal{W}, \leq_{\varphi \sqcup \varphi'}) \not\subseteq \text{Mod}(\varphi)$
5. Si $I \in \min(\mathcal{W}, \leq_{\Psi_1})$ et $I \in \min(\mathcal{W}, \leq_{\Psi_2})$, alors $I \in \min(\mathcal{W}, \leq_{\Psi_1 \sqcup \Psi_2})$
6. Si $\min(\mathcal{W}, \leq_{\Psi_1}) \cap \min(\mathcal{W}, \leq_{\Psi_2}) \neq \emptyset$ et $I \notin \min(\mathcal{W}, \leq_{\Psi_1})$, alors $I \notin \min(\mathcal{W}, \leq_{\Psi_1 \sqcup \Psi_2})$

Un assignement synchrétique pur majoritaire est un assignement synchrétique pur qui vérifie :

7. $\forall \Psi \in \Psi \forall \varphi \in \varphi \exists n \min(\mathcal{W}, \leq_{\Psi \sqcup \varphi^n}) \subseteq \text{Mod}(\varphi)$

Un assignement synchrétique pur juste est un assignement synchrétique pur qui vérifie :

8. $\forall \varphi' \exists \varphi$ si $\text{Mod}(\varphi') \not\subseteq \text{Mod}(\varphi)$, alors $\forall n \min(\mathcal{W}, \leq_{\varphi' \sqcup \varphi^n}) = \min(\mathcal{W}, \leq_{\varphi' \sqcup \varphi})$

Si l'on dispose d'un assignement qui associe à chaque ensemble de connaissance Ψ un pré-ordre $\leq_\Psi$ sur $\mathcal{W}$, alors on peut définir une fonction $\delta : \mathcal{M} \rightarrow \mathcal{I}$ par : soit $M \in \mathcal{M}$ tel que $\iota(\Psi) = M$, alors

$$\delta(M) = \min(\mathcal{W}, \leq_\Psi) \quad (8.1)$$

Si l'assignement satisfait la condition 3 alors δ est bien défini.

De même, si on dispose d'une fonction $\delta : \mathcal{M} \rightarrow \mathcal{I}$ on peut définir une famille de relations sur les interprétations par $\forall \Psi \in \mathcal{S}$:

$$\leq_\Psi = [\delta(\iota(\Psi)) \times (\mathcal{W} \setminus \delta(\iota(\Psi)))] \cup \{\{I, I\} \mid I \in \mathcal{W}\} \quad (8.2)$$

Il est facile de montrer que, si on a un assignement synchrétique pur (juste, majoritaire), alors la fonction obtenue par l'équation 8.1 est une fonction de fusion (juste, majoritaire). De même, si on a une fonction de fusion (juste, majoritaire), alors la famille de relations obtenues par l'équation 8.2 est un assignement synchrétique pur (juste, majoritaire).

Cette observation et le théorème 103 donnent donc directement :

Théorème 104 *Un opérateur est un opérateur de fusion pure (respectivement un opérateur de fusion pure majoritaire ou un opérateur d'arbitrage pur) si et seulement si il existe un assignement synchrétique pur (respectivement un assignement synchrétique pur majoritaire ou un assignement synchrétique pur juste) qui associe à chaque ensemble de connaissance Ψ un pré-ordre $\leq_\Psi$ tel que*

$$\text{Mod}(\Delta(\Psi)) = \min(\mathcal{W}, \leq_\Psi).$$

Cette caractérisation est compatible avec celle des opérateurs de fusion contrainte. Plus exactement, on peut simuler l'absence de contraintes d'intégrité en prenant $\mu = \top$. Dans ce cas particulier, les postulats de la fusion contrainte sont notés $(\text{IC}_\top)$.

$(\text{IC}_{0_\top})$, $(\text{IC}_{7_\top})$ et $(\text{IC}_{8_\top})$ sont trivialement vrais. $(\text{Arb}_\top)$ est également vrai directement puisque la prémisse de l'implication est toujours fausse. En fait (Arb) n'est pas exprimable

lorsque l'on n'a pas de contraintes d'intégrité. C'est pour cela que, dans ce cadre, on a pris (WMI) comme approximation. Mais cette propriété est moins satisfaisante, car elle n'exprime qu'une sorte de non-majorité et n'est donc pas une caractérisation directe de l'arbitrage. Alors que (Arb) définit l'arbitrage de façon plus directe.

Les autres postulats sont les mêmes dans les deux cadres, les principales différences sont que le postulat (IC4_⊤) est plus fort que (A4) et que le postulat (Maj_⊤) est plus fort que (M7).

Le résultat suivant résume cette compatibilité entre la caractérisation dans le cas avec et sans contrainte d'intégrité.

Théorème 105 *Si Δ est un opérateur de fusion contrainte, alors $\Delta_{\top}$ est un opérateur de fusion pure. De plus, si Δ est un opérateur de fusion contrainte majoritaire, alors $\Delta_{\top}$ est un opérateur de fusion pure majoritaire.*

8.3 Fusion contrainte et révision commutative

Les opérateurs de révision commutative définis par Liberatore et Schaerf effectuent une fusion de deux bases de connaissance. Le résultat de cette fusion doit impliquer la disjonction des deux bases. Cette contrainte peut être exprimée dans le cadre de la fusion contrainte en fixant comme contraintes d'intégrité de la fusion la disjonction des deux bases de connaissance. On peut alors se demander s'il est possible d'utiliser les opérateurs de fusion contrainte pour définir des opérateurs de révision commutative.

Supposons que Δ soit un opérateur de fusion contrainte. La question est comment définir un opérateur de révision commutatif $\diamond_{\Delta}$ à partir de Δ ?

Cela peut être fait de deux manières assez naturelles. On peut poser d'une part

$$\varphi_1 \diamond_{\Delta} \varphi_2 = \Delta_{\varphi_1 \vee \varphi_2}(\varphi_1 \sqcup \varphi_2)$$

C'est-à-dire forcer le résultat de la fusion à choisir dans la disjonction des deux bases comme l'impose (LS7). L'autre possibilité est d'utiliser l'idée de la révision commutative :

$$\varphi_1 \diamond_{\Delta} \varphi_2 = \varphi_1 \circ \varphi_2 \vee \varphi_2 \circ \varphi_1$$

où $\circ$ est l'opérateur de révision associé à Δ (cf section 8.1). En fait ces deux définitions coïncident :

Théorème 106 *Si Δ est un opérateur de fusion contrainte, alors il satisfait*

$$\Delta_{\varphi \vee \mu}(\varphi \sqcup \mu) \leftrightarrow \Delta_{\varphi}(\mu) \vee \Delta_{\mu}(\varphi)$$

Ce qui peut se formuler également :

$$\Delta_{\varphi \vee \mu}(\varphi \sqcup \mu) \leftrightarrow \mu \circ \varphi \vee \varphi \circ \mu$$

où $\varphi \circ \mu = \Delta_{\mu}(\varphi)$ est l'opérateur de révision AGM associé à Δ .

Preuve : Soit $\varphi \mapsto \leq_\varphi$ l'assignement synchrétique représentant Δ . Si $\varphi \wedge \mu$ est consistant le résultat est trivial. Supposons donc $\varphi \wedge \mu$ inconsistent. On a

$$\min(\varphi \vee \mu, \leq_{\varphi \sqcup \mu}) = \begin{cases} \min(\varphi, \leq_{\varphi \sqcup \mu}) & \text{ou} \\ \min(\mu, \leq_{\varphi \sqcup \mu}) & \text{ou} \\ \min(\varphi, \leq_{\varphi \sqcup \mu}) \cup \min(\mu, \leq_{\varphi \sqcup \mu}) \end{cases}$$

Mais les deux premiers cas sont impossibles d'après la condition 4 (équité). Donc seul le dernier cas est possible i.e. $\min(\varphi \vee \mu, \leq_{\varphi \sqcup \mu}) = \min(\varphi, \leq_{\varphi \sqcup \mu}) \cup \min(\mu, \leq_{\varphi \sqcup \mu})$, ce qui peut être réécrit $\min(\varphi \vee \mu, \leq_{\varphi \sqcup \mu}) = \min(\varphi, \leq_\mu) \cup \min(\mu, \leq_\varphi)$ ce qui, par définition, donne $\Delta_{\varphi \vee \mu}(\varphi \sqcup \mu) \leftrightarrow \Delta_\varphi(\mu) \vee \Delta_\mu(\varphi)$. □

Bien que les deux définitions coïncident, comme nous venons de le voir, nous préférons la première, car elle peut se généraliser simplement à plus de deux bases de connaissance.

Définition 100 Soit Δ , un opérateur de fusion contrainte. On définit un opérateur de révision commutative $\diamond_\Delta$ par

$$\varphi \diamond_\Delta \mu = \Delta_{\varphi \vee \mu}(\varphi \sqcup \mu)$$

On dit que $\diamond_\Delta$ est l'opérateur de révision commutative associé à Δ .

Théorème 107 Si Δ est un opérateur de fusion contrainte, alors l'opérateur $\diamond_\Delta$ associé satisfait (LS1-LS5), (LS7) et (LS8).

Preuve : (LS1) direct par définition. (LS2) directement de (IC2). (LS3) directement de (IC2). (LS4) directement de (IC0) et (IC1). (LS5) directement de (IC3). (LS7) direct par définition. (LS8) directement de (IC4) et (IC2). □

Pour que les opérateurs $\diamond_\Delta$ définis d'une telle manière soient des opérateurs de révision commutative, nous avons besoin d'une propriété supplémentaire sur Δ . Nous verrons cela après avoir donné un théorème de représentation pour les opérateurs de révision commutative en termes de familles de pré-ordres sur les interprétations. Ce théorème est plus naturel que celui donné dans [LS98] par Liberatore et Schaerf, et permet de mettre en valeur la propriété nécessaire aux opérateurs de révision pour donner des opérateurs de révision commutative.

Soit un assignement fidèle $\varphi \mapsto \leq_\varphi$. On définit le *relevage* de cet assignement comme étant la fonction $F \mapsto \leq_F$ où F est un ensemble d'interprétations et $\leq_F$ est un pré-ordre total sur les ensembles d'interprétations tel que

$$A \leq_F B \text{ ssi } \exists I \in A \forall J \in B \ I \leq_\varphi J$$

où φ est une formule telle que $\text{Mod}(\varphi) = F$.

Le résultat suivant est implicite dans le travail de Liberatore et Schaerf. C'est en fait le noyau de leur théorème de représentation.

Observation 7 $\diamond$ est un opérateur de révision commutative ssi il existe un unique opérateur de révision $\circ$ tel que :

(i) le relevage de l'assignement fidèle représentant $\circ$ est un assignement commutatif, et

$$(ii) \mu \diamond \theta = (\mu \circ \theta) \vee (\mu \circ \theta)$$

Nous voyons à présent quelle condition est nécessaire sur les opérateurs de révision pour qu'ils donnent des opérateurs de révision commutative.

Théorème 108 *Un opérateur $\diamond$ défini à partir d'un opérateur de révision $\circ$ par $\varphi \diamond \mu = \varphi \circ \mu \vee \mu \circ \varphi$ satisfait (LS1-LS8) si et seulement si $\circ$ satisfait la propriété suivante :*

$$(\mu \vee \theta) \circ \varphi = \begin{cases} \mu \circ \varphi & \text{si } \varphi \circ (\mu \vee \theta) \vdash \neg \theta \\ \theta \circ \varphi & \text{si } \varphi \circ (\mu \vee \theta) \vdash \neg \mu \\ \mu \circ \varphi \vee \theta \circ \varphi & \text{sinon} \end{cases} \quad (8.3)$$

Preuve : Notons d'abord qu'en utilisant le théorème de représentation de la révision [KM91b] on obtient directement que la condition (8.3) correspond à la condition suivante sur l'assignement fidèle représentant $\circ$:

$$\min(\text{Mod}(\varphi), \leq_{\mu \vee \theta}) = \begin{cases} \min(\text{Mod}(\varphi), \leq_{\mu}) & \text{si } \mu <_{\varphi} \theta \\ \min(\text{Mod}(\varphi), \leq_{\theta}) & \text{si } \theta <_{\varphi} \mu \\ \min(\text{Mod}(\varphi), \leq_{\mu}) \cup \min(\text{Mod}(\varphi), \leq_{\theta}) & \text{sinon} \end{cases} \quad (8.4)$$

Où $\mu \leq_{\varphi} \theta$ est défini de manière naturelle :

$$\mu \leq_{\varphi} \theta \text{ ssi } \exists I \models \mu \forall J \models \theta \ I \leq_{\varphi} J$$

(Si) Supposons que $\circ$ est un opérateur de révision vérifiant la condition (8.3) et définissons $\diamond$ en posant $\varphi \diamond \mu = \varphi \circ \mu \vee \mu \circ \varphi$. On veut montrer que $\diamond$ est un opérateur de révision commutative. Considérons l'assignement fidèle $\varphi \mapsto \leq_{\varphi}$ représentant $\circ$ et soit $F \mapsto \leq_F$ le relevage associé. Par le théorème 71 il suffit de montrer que $F \mapsto \leq_F$ est un assignement commutatif.

On vérifie facilement que l'assignement $F \mapsto \leq_F$ satisfait L1-L4. Il reste à vérifier L5. Nous allons montrer la propriété suivante (qui est équivalente à L5) :

$$A \leq_{C \cup D} B \Leftrightarrow \begin{cases} (i) & C <_{A \cup B} D \text{ et } A \leq_C B & \text{ou} \\ (ii) & D <_{A \cup B} C \text{ et } A \leq_D B & \text{ou} \\ (iii) & D \simeq_{A \cup B} C \text{ et } (A \leq_C B \text{ ou } A \leq_D B) \end{cases}$$

Soient les formules φ, μ, θ telles que $\text{mod}(\varphi) = A \cup B$, $\text{mod}(\mu) = C$, $\text{mod}(\theta) = D$. De droite à gauche on considère trois cas :

- (i) est vérifié. Par l'absurde, supposons que $A \leq_{C \cup D} B$ n'est pas vrai, alors $B <_{C \cup D} A$. On a donc $\min(A \cup B, \leq_{\mu \vee \theta}) \cap A = \emptyset$. De $C <_{A \cup B} D$ i.e. $\mu <_{\varphi} \theta$ et de la condition (8.4) on a que $\min(A \cup B, \leq_{\mu \vee \theta}) = \min(A \cup B, \leq_{\mu})$. De $A \leq_C B$ on a $\min(A \cup B, \leq_{\mu}) \cap A \neq \emptyset$. D'où $\min(A \cup B, \leq_{\mu \vee \theta}) \cap A \neq \emptyset$. Contradiction.
- (ii) est vérifié. La preuve est similaire à celle de (i).
- (iii) est vérifié. Supposons que $A \leq_{C \cup D} B$ n'est pas vrai, c'est-à-dire $A >_{C \cup D} B$. De $D \simeq_{A \cup B} C$, i.e. $\mu \simeq_{\varphi} \theta$ et de la propriété (8.4), on a

$$\min(A \cup B, \leq_{\mu \vee \theta}) = \min(A \cup B, \leq_{\mu}) \cup \min(A \cup B, \leq_{\theta}) \quad (*)$$

Par hypothèse on sait que $A \leq_C B$ or $A \leq_D B$, supposons sans perte de généralité que $A \leq_C B$, alors $\min(A \cup B, \leq_{\mu}) \cap A \neq \emptyset$. D'où, d'après l'équation (*) on obtient $\min(A \cup B, \leq_{\mu \vee \theta}) \cap A \neq \emptyset$. Mais $A >_{C \cup D} B$ implique que $\min(A \cup B, \leq_{\mu \vee \theta}) \cap A = \emptyset$. Contradiction.

De gauche à droite on considère également trois cas : soit $C <_{A \cup B} D$ ou $D <_{A \cup B} C$ ou $C \simeq_{A \cup B} D$.

Cas $C <_{A \cup B} D$: on veut montrer que $A \leq_C B$. Supposons, vers une contradiction, que ce n'est pas le cas, i.e. $B <_C A$. De $A \leq_{C \cup D} B$ on a que $\min(A \cup B, \leq_{\mu \vee \theta}) \cap A \neq \emptyset$. Mais de $C <_{A \cup B} D$ on a $\min(A \cup B, \leq_{\mu \vee \theta}) = \min(A \cup B, \leq_C)$. A présent de $B <_C A$ on obtient $\min(A \cup B, \leq_C) \cap A = \emptyset$. Contradiction.

Cas $D <_{A \cup B} C$: Ce cas est similaire au précédent.

Cas $C \simeq_{A \cup B} D$: Supposons que $A \leq_C B$ or $A \leq_D B$ n'est pas vérifié, c'est-à-dire $A >_C B$ and $A >_D B$, alors $\min(A \cup B, \leq_{\mu}) \cap A = \emptyset$ et $\min(A \cup B, \leq_{\theta}) \cap A = \emptyset$. A partir de $C \simeq_{A \cup B} D$ on déduit $\min(A \cup B, \leq_{\mu \vee \theta}) = \min(A \cup B, \leq_{\mu}) \cup \min(A \cup B, \leq_{\theta})$. De $A \leq_{C \cup D} B$ on a que $\min(A \cup B, \leq_{\mu \vee \theta}) \cap A \neq \emptyset$. Mais par hypothèse on a que $(\min(A \cup B, \leq_{\mu}) \cup \min(A \cup B, \leq_{\theta})) \cap A = \emptyset$. Contradiction.

(Seulement Si) Supposons que l'on dispose d'un opérateur de révision commutative $\diamond$ satisfaisant les postulats (LS1-LS8), qui est défini à partir d'un opérateur de révision $\circ$ par $A \diamond B = A \circ B \vee B \circ A$. On veut montrer que $\circ$ satisfait la condition (8.3). De (LS6) on obtient par définition

$$\varphi \circ (\mu \vee \theta) \vee (\mu \vee \theta) \circ \varphi = \begin{cases} \varphi \circ \mu \vee \mu \circ \varphi & \text{ou} \\ \varphi \circ \theta \vee \theta \circ \varphi & \text{ou} \\ \varphi \circ \mu \vee \mu \circ \varphi \vee \varphi \circ \theta \vee \theta \circ \varphi & \end{cases} \quad (8.5)$$

Notons que l'on a également la propriété classique des opérateurs de révision [Gär88] :

$$\varphi \circ (\mu \vee \theta) = \begin{cases} \varphi \circ \mu & \text{ou} \\ \varphi \circ \theta & \text{ou} \\ \varphi \circ \mu \vee \varphi \circ \theta & \end{cases} \quad (8.6)$$

On considère les trois cas de la condition (8.3) :

cas 1 : Supposons que $\varphi \circ (\mu \vee \theta) \vdash \neg \theta$. Dans ce cas les deux dernières alternatives de (8.6) sont clairement impossibles, on a donc nécessairement $\varphi \circ (\mu \vee \theta) = \varphi \circ \mu$. L'hypothèse implique également que $\varphi \wedge \theta \vdash \perp$ et il est alors facile de voir que les deux derniers cas de (8.5) sont impossibles, on a donc nécessairement

$$\varphi \circ (\mu \vee \theta) \vee (\mu \vee \theta) \circ \varphi = \varphi \circ \mu \vee \mu \circ \varphi \quad (8.7)$$

A présent supposons que $\varphi \wedge \mu \vdash \perp$. De cela, de l'hypothèse et de l'équation (8.7) on obtient $(\mu \vee \theta) \circ \varphi = \mu \circ \varphi$. Sinon, supposons que $\varphi \wedge \mu \not\vdash \perp$. Alors, par les propriétés des opérateurs de révision et le fait que par hypothèse $\varphi \wedge \theta \vdash \perp$, on a $(\mu \vee \theta) \circ \varphi = \mu \wedge \varphi = \mu \circ \varphi$. C'est-à-dire que le premier cas de la condition (8.3) est vérifié.

cas 2 : Supposons $\varphi \circ (\mu \vee \theta) \vdash \neg \mu$. Dans ce cas on suit le même raisonnement que dans le cas 1 en intervertissant les rôles de μ et θ . On obtient alors $(\mu \vee \theta) \circ \varphi = \theta \circ \varphi$. C'est-à-dire que le second cas de la condition (8.3) est vérifié.

cas 3 : Dans ce cas il est facile de voir que

$$\varphi \circ (\mu \vee \theta) = (\varphi \circ \mu) \vee (\varphi \circ \theta) \quad (8.8)$$

On considère trois possibilités. Supposons d'abord que $\varphi \wedge (\mu \vee \theta) \not\models \perp$. Alors $\varphi \circ (\mu \vee \theta) = (\varphi \wedge \mu) \vee (\varphi \wedge \theta)$ avec $\varphi \wedge \mu \not\models \perp$ et $\varphi \wedge \theta \not\models \perp$. Il s'ensuit alors $(\mu \vee \theta) \circ \varphi = (\mu \wedge \varphi) \vee (\theta \wedge \varphi) = (\mu \circ \varphi) \vee (\theta \circ \varphi)$.

Supposons à présent que $\varphi \wedge (\mu \vee \theta) \vdash \perp$. De là et de la formulation de (LS6) qui suit

$$\varphi \circ (\mu \vee \theta) \vee (\mu \vee \theta) \circ \varphi = \begin{cases} (\varphi \circ \mu) \vee (\mu \circ \varphi) & \text{ou} \\ (\varphi \circ \theta) \vee (\theta \circ \varphi) & \text{ou} \\ (\varphi \circ \mu) \vee (\varphi \circ \theta) \vee (\mu \circ \varphi) \vee (\theta \circ \varphi) & \end{cases} \quad (8.9)$$

on a facilement que $Mod((\mu \vee \theta) \circ \varphi) \subseteq Mod(\mu \circ \varphi) \cup Mod(\theta \circ \varphi)$. Pour établir l'inclusion inverse prenons $I \models \mu \circ \varphi$ (l'autre cas est analogue). Alors $\varphi_{\{I\}} \leq_{\mu} \varphi$. D'où, en notant que $\mu \simeq_{\varphi} \theta$ (ce qui est une conséquence de (8.8)), on a $\varphi_{\{I\}} \leq_{\mu \vee \theta} \varphi$, c'est-à-dire que I est minimal dans $Mod(\varphi)$ pour $\leq_{\mu \vee \theta}$. □

Remarque 8 Il est facile de voir que si l'opérateur $\circ$ est défini à partir d'une distance, la propriété (8.9) est vérifiée. Dans ce cas, l'opérateur $\diamond$ associé par $\varphi \diamond \mu = (\varphi \circ \mu) \vee (\mu \circ \varphi)$ est un opérateur de révision commutative.

A partir de ce résultat, il est facile d'établir un théorème de représentation pour les opérateurs de révision commutative.

Définition 101 Un assignement fidèle commutatif est une fonction qui associe à chaque base de connaissance φ un pré-ordre $\leq_{\varphi}$ sur les interprétations tel que :

1. Si $I \models \varphi$ et $J \models \varphi$, alors $I \simeq_{\varphi} J$
2. Si $I \models \varphi$ et $J \not\models \varphi$, alors $I <_{\varphi} J$
3. Si $\varphi \leftrightarrow \varphi'$, alors $\leq_{\varphi} = \leq_{\varphi'}$
4. $\min(Mod(\varphi), \leq_{\mu \vee \theta}) = \begin{cases} \min(Mod(\varphi), \leq_{\mu}) & \text{si } \mu <_{\varphi} \theta \\ \min(Mod(\varphi), \leq_{\theta}) & \text{si } \theta <_{\varphi} \mu \\ \min(Mod(\varphi), \leq_{\mu}) \cup \min(Mod(\varphi), \leq_{\theta}) & \text{sinon} \end{cases}$

Où $\mu \leq_{\varphi} \theta$ est défini de façon naturelle :

$$\mu \leq_{\varphi} \theta \text{ ssi } \exists I \models \mu \forall J \models \theta \ I \leq_{\varphi} J$$

Théorème 109 Un opérateur $\diamond$ satisfait (LS1-LS8) si et seulement si il existe un assignement fidèle commutatif qui associe à chaque base de connaissance φ un pré-ordre $\leq_{\varphi}$ tel que :

$$\text{mod}(\varphi \diamond \varphi') = \min(Mod(\varphi), \leq_{\varphi'}) \cup \min(Mod(\varphi'), \leq_{\varphi})$$

Preuve : La partie **Si** est une conséquence du théorème 108. Pour la partie **Seulement si** notons qu'à partir d'un opérateur de révision commutative $\diamond$ on peut définir un opérateur de révision $\circ$ et par le théorème 108 conclure que cet opérateur vérifie la condition 4 de l'assignement fidèle commutatif. □

Nous avons vu, théorème 108, comment caractériser les opérateurs de révision donnant des opérateurs de révision commutative. Il est alors évident que tous les opérateurs de révision commutative définis à partir d'opérateurs de fusion contrainte ne satisfont pas automatiquement (LS6). En fait, à partir du théorème 106, on peut facilement voir que la propriété suivante correspond à la propriété (8.3) :

$$\Delta_{\varphi}(\mu \vee \theta) = \begin{cases} \Delta_{\varphi}(\mu) & \text{si } \Delta_{\mu \vee \theta}(\varphi) \vdash \neg \theta \\ \Delta_{\varphi}(\theta) & \text{si } \Delta_{\mu \vee \theta}(\varphi) \vdash \neg \mu \\ \Delta_{\varphi}(\mu) \vee \Delta_{\varphi}(\theta) & \text{sinon} \end{cases} \quad (8.10)$$

Le théorème suivant peut donc, d'après les théorèmes 97 et 106, être simplement considéré comme une reformulation du théorème 108 :

Théorème 110 *Si Δ est un opérateur de fusion contrainte, alors l'opérateur $\diamond$ défini par $\varphi \diamond \mu = \Delta_{\varphi \vee \mu}(\varphi \sqcup \mu)$ satisfait (LS1-LS8) si et seulement si Δ satisfait la propriété (8.10).*

Comme corollaire du dernier théorème, on peut noter que les opérateurs $\diamond$ définis à partir des familles $\Delta^{d, \Sigma}$ et $\Delta^{d, GMax}$ satisfont tous les postulats de la révision commutative. Cela se montre facilement puisque ces opérateurs sont définis à partir d'une distance. Les opérateurs de révision correspondants sont également définis à partir d'une distance et on conclut par la remarque 8.

Il est également intéressant de noter que la propriété (8.10) implique la propriété suivante :

Lemme 111 *La propriété (8.10) implique $\Delta_{\mu}(\varphi) = \Delta_{\mu}(\Delta_{\varphi}(\mu))$*

Preuve : Soit $\varphi' = \min(\text{Mod}(\varphi), \leq_{\mu})$, c'est-à-dire $\varphi' = \Delta_{\varphi}(\mu)$. On définit φ'' comme $\text{Mod}(\varphi'') = \{I \models \varphi \text{ and } I \not\models \varphi'\}$. De la propriété (8.10) on a que $\min(\text{Mod}(\mu), \leq_{\varphi' \vee \varphi''}) = \min(\text{Mod}(\mu), \leq_{\varphi'})$ puisque par hypothèse $\varphi' <_{\mu} \varphi''$. Donc $\min(\text{Mod}(\mu), \leq_{\varphi}) = \min(\text{Mod}(\mu), \leq_{\varphi'})$. On a alors $\Delta_{\mu}(\varphi) = \Delta_{\mu}(\varphi')$, et à partir de la définition de φ' cela est $\Delta_{\mu}(\varphi) = \Delta_{\mu}(\Delta_{\varphi}(\mu))$

□

Cette propriété ne semble peut-être pas très naturelle, mais elle est très expressive quand elle est représentée en termes d'opérateur de révision puisque l'opérateur de révision correspondant à Δ vérifie :

$$\varphi \circ \mu = (\mu \circ \varphi) \circ \mu \quad (8.11)$$

C'est-à-dire que le résultat de la révision de φ par μ dépend seulement des modèles de φ les plus proches de μ . Les opérateurs de révision définis à partir d'une distance vérifient cette propriété.

Un sérieux défaut des opérateurs de révision commutative est qu'ils ne permettent pas de fusionner plus de deux bases de connaissance puisqu'ils ne sont pas associatifs (voir [LS95, LS98]), mais exiger que le résultat de la fusion doit impliquer la disjonction des bases de connaissance peut être très utile pour certaines applications.

Donc, de manière à généraliser les opérateurs de Liberatore et Schaerf à n bases de connaissance, on peut alors définir la révision commutative d'un ensemble de connaissance $\{\varphi_1 \sqcup \dots \sqcup \varphi_n\}$ comme :

$$\Delta_{\varphi_1 \vee \dots \vee \varphi_n}(\varphi_1 \sqcup \dots \sqcup \varphi_n)$$

8.4 Fusion contrainte et théorie du choix social

En économie, le problème de l'agrégation de préférences individuelles en une préférence globale est de première importance. On s'est intéressé très tôt à l'étude comparative des différents systèmes de vote [dC85, Par97]. Il est facile de montrer que le résultat d'une élection dépend autant du système électoral que des préférences individuelles, c'est-à-dire qu'avec des mêmes préférences individuelles on peut choisir (ou tout au moins modifier) le vainqueur de l'élection en changeant le mode de scrutin. Il est alors important de pouvoir répondre à un certain nombre de questions. Comment caractériser les différentes méthodes de votes? Qu'est-ce qu'une bonne méthode de vote? Comment dire qu'une méthode est meilleure qu'une autre?

Il est donc clair que la théorie de la fusion de bases de connaissance à beaucoup à voir avec ces méthodes d'agrégation de préférences.

Or, le résultat le plus célèbre de la théorie du choix social, le théorème d'impossibilité d'Arrow [Arr63, Kel78], énonce qu'il n'est pas possible de définir une "bonne" méthode d'agrégation. C'est-à-dire qu'Arrow donne un ensemble de cinq propriétés qui semblent toutes nécessaires à une méthode d'agrégation et montre ensuite qu'aucune méthode ne satisfait toutes ces propriétés simultanément.¹

Il serait donc intéressant de voir en quoi les opérateurs de fusion de bases de connaissance diffèrent des méthodes d'agrégation de préférences de façon à comprendre comment ils échappent à ce théorème d'impossibilité.

8.4.1 Le théorème d'impossibilité d'Arrow

Nous avons besoin de définir tout d'abord quelques notions :

On note X un ensemble non vide, dont les éléments seront appelés *alternatives*. Ces alternatives doivent être exclusives et on supposera que ce sont des descriptions complètes du monde.

Généralement, lorsque l'on doit opérer un choix, toutes les alternatives possibles ne sont pas disponibles. Certaines contraintes peuvent limiter le nombre d'alternatives à un sous-ensemble de X , noté v , que l'on nommera *agenda*.

On appelle *relation de préférence* (individuelle), un pré-ordre total $\leq_i$ sur X dénotant les préférences de l'individu i sur l'ensemble des alternatives de X .

1. Une conséquence de ce résultat est qu'il n'existe pas de bon système électoral !

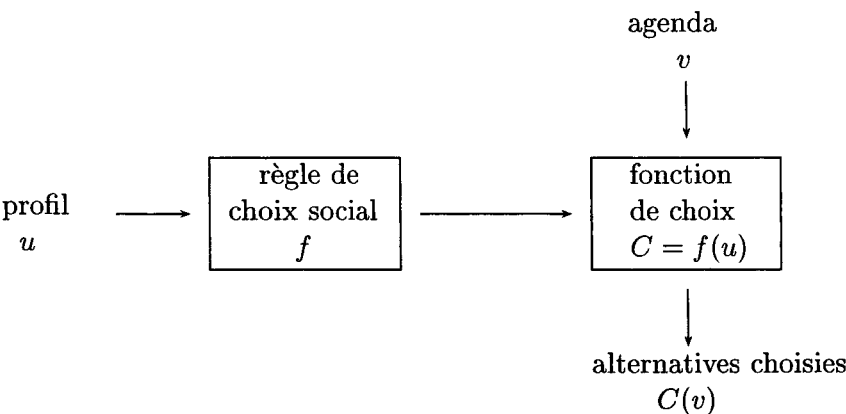


FIG. 8.1 – Règle de choix social

On appelle *profil* un (multi-)ensemble de relations de préférences individuelles.

Une *fonction de choix* C est une fonction qui choisit parmi les alternatives d'un agenda v un ensemble d'alternatives admissibles $C(v)$ tel que $C(v) \neq \emptyset$ et $C(v) \subseteq v$.

Le but d'une *règle de choix social* est, étant donné un ensemble de préférences individuelles (un profil), de déterminer une fonction de choix correspondante. C'est-à-dire qu'une règle de choix social f est une fonction f telle que pour tout profil u , $f(u) = C$ où C est une fonction de choix (voir figure 8.1).

L'objet de la théorie de choix social est l'étude des règles de choix social. Un simple calcul combinatoire montre qu'il n'est pas possible d'étudier les règles de choix social individuellement comme illustré dans l'exemple 17 [Kel88]. C'est pour cela que l'on s'intéresse à des classes de règles qui satisfont certaines propriétés. C'est la même idée qui justifie l'approche axiomatique des relations d'inférence non monotones, de la révision et de la fusion de bases de connaissance.

Exemple 17 *Considérons un exemple simple où l'on a cinq individus et quatre alternatives. Il y a 75 relations de préférences possibles sur ces quatres alternatives. Ce qui donne donc $75^5 = 2\,373\,000\,000$ profils possibles. D'un autre coté, quatre alternatives donnent $15 \cdot 7^4 \cdot 3^6 = 26\,254\,935$ fonctions de choix possibles. Le nombre de règles de choix social possibles est donc*

$$(15 \cdot 7^4 \cdot 3^6)^{75^5} \simeq 10^{235\,000\,000\,000}$$

On peut noter que cet exemple n'est pas très compliqué car il n'est pas rare de trouver des cas où le nombre d'individus est supérieur à cinq et le nombre d'alternatives supérieur à 4 ! Le nombre de règles de choix social possibles est donc quelque chose de difficilement imaginable.

On pourrait argumenter que ce nombre ne représente que les règles de choix social mathématiquement convenables, mais en aucun cas les règles "raisonnables". En effet la règle de choix social choisissant, par exemple, toujours la première alternative de l'agenda, sans tenir compte du profil, est compté parmi ces règles. Mais il n'est pas concevable de l'utiliser pour une quelconque prise de décision.

Il est donc possible de réduire considérablement le nombre de règles de choix social en ne considérant que celles qui satisfont un ensemble de critères de rationalité.

Nous allons à présent énumérer quelques uns de ces critères.

La contrainte du domaine standard :

- Il y a au moins trois alternatives dans X
- Il y a au moins trois individus
- Une règle de choix social a comme domaine l'ensemble des profils possibles à partir des relations de préférences sur X
- Toute fonction de choix, résultat de la règle de choix social, a comme domaine tous les agendas possibles.

Cette contrainte est très naturelle, les deux premiers points enlèvent simplement les cas les plus simples, et la contrainte demande simplement que l'application de la règle de choix social donne un résultat quel que soit le profil et l'agenda donnés en entrée.

La condition forte de Pareto :

Soit un profil u , et soit C_u la fonction de choix associée à u par la règle de choix social f^2 . Si tous les individus de u trouvent une alternative x au moins aussi bonne qu'une alternative y , et si au moins un des individus préfère strictement x à y , alors si x est une alternative possible ($x \in v$), alors y ne sera pas choisi ($y \notin C_u(v)$).

Cette condition permet d'exclure une alternative y du choix final s'il existe une alternative que personne ne trouve moins bonne que y et qu'au moins un individu trouve strictement préférable à y .

L'indépendance des alternatives non disponibles :

Si les restrictions de deux profils u et u' à un agenda v coïncident, alors les choix faits à partir de cet agenda seront les mêmes : $C_u(v) = C_{u'}(v)$

Cette propriété assure que lorsque l'on désire faire un choix dans un ensemble d'alternatives disponibles v , ce choix ne sera fait qu'en fonction des préférences entre ces alternatives. Celui-ci ne dépendra donc pas d'alternatives non-disponibles. Par exemple, si l'on doit choisir à une élection entre trois candidats : Garry, Anatoly et Bobby. Si, juste avant l'élection Bobby décide de retirer sa candidature, il semble raisonnable de ne pas faire intervenir les préférences relatives entre Bobby, Garry et Anatoly dans notre décision mais simplement les préférences entre Garry et Anatoly.

La propriété d'avoir des explications transitives :

On dit qu'une fonction de choix C a des explications transitives s'il existe un pré-ordre $\prec_C$ tel que $C(v) = \{x \in v : x \prec_C y \text{ pour tout } y \in v\}$.

On dit qu'une règle de choix social a des explications transitives si toutes les fonctions de choix qu'elle produit ont des explications transitives.

Cette propriété assure une certaine rationalité au niveau des préférences globales. C'est-à-dire, en quelque sorte, que l'on dispose d'une explication au choix effectué. Cela assure que lorsque l'on garde un même profil, mais que l'on fait varier l'agenda, on a un comportement

2. i.e. $f(u) = C_u$.

raisonnable. Par exemple si les deux alternatives possibles sont $v = \{x, y\}$ et que $C(v) = \{x\}$ et si en prenant comme alternatives $v' = \{y, z\}$ alors $C(v') = \{y\}$, alors si on a le choix entre l'ensemble des alternatives $v'' = \{x, y, z\}$, ce sera x qui sera choisi $C(v'') = \{x\}$. Cela semble raisonnable puisque le profil (i.e. les préférences individuelles) ne changent pas.

L'absence de dictateur:

Une règle de choix social n'a pas de dictateur s'il n'existe pas d'individu i tel que $\forall x, y \in v$ si $x <_i y$ alors $y \notin C_u(v)$.

Cette propriété assure que l'on n'a pas un individu qui dispose de tous les pouvoirs. Cela est une propriété difficilement critiquable pour une règle de choix social.

Nous venons d'énumérer des propriétés qui caractérisent des règles de choix social "raisonnables". Nous espérons donc ainsi faire chuter le nombre de règles suffisamment pour entreprendre une étude un peu plus poussée. Et bien il s'avère que nous avons dépasser notre objectif puisque nous avons réduit ce nombre à zéro. Il n'existe aucune règle de choix social satisfaisant ces cinq règles. Ce résultat est connu sous le nom de *théorème d'impossibilité d'Arrow* [Arr63]:

Théorème 112 *Il n'existe pas de règle de choix social satisfaisant toutes les propriétés suivantes :*

- la contrainte du domaine standard
- la condition forte de Pareto
- l'indépendance des alternatives non disponibles
- le fait d'avoir des explications transitives
- l'absence de dictateur

Ce résultat semble d'autant plus surprenant que, prises individuellement, ces conditions semblent tout à fait acceptables, voire nécessaires, pour une règle de choix social.

Ce théorème est l'un des résultats les plus importants en théorie du choix social. Depuis d'autres théorèmes d'impossibilité ont été montrés (cf e.g. [Sen79]).

8.4.2 Impossibilité de la fusion contrainte?

Nous allons voir dans cette section quelles sont les propriétés énumérées précédemment vérifiées par les opérateurs de fusion contrainte. Et nous expliquerons pourquoi certaines ne le sont pas.

Nous allons tout d'abord expliciter la correspondance entre fusion contrainte et théorie du choix social. Celle-ci sera principalement basée sur la représentation des opérateurs de fusion contrainte sous forme de familles de pré-ordres sur les interprétations.

Une alternative est une interprétation.

Les individus pour la fusion contrainte sont représentés par leur base de connaissance φ . Etant donné un opérateur de fusion (i.e. une règle de choix social), cette base de connaissance correspond au pré-ordre $\leq_\varphi$. C'est donc ce pré-ordre qui tiendra lieu de relation de préférence

individuelle.

Un profil est un ensemble de préférences individuelles, c'est-à-dire un ensemble de connaissance dans le cadre de la fusion.

Un agenda est un sous-ensemble des alternatives, c'est-à-dire la base μ représentant les contraintes d'intégrité pour la fusion.

Les règles de choix social que nous allons considérer sont des opérateurs de fusion contrainte, ils effectuent donc l'agrégation des relations de préférences individuelles en une relation de préférence globale. Les alternatives choisies par la fonction de choix sont alors les alternatives minimales pour cette relation de préférence globale. Plus exactement les fonctions de choix correspondantes s'écrivent $f(u) = \min(v, \leq_u)$ où le pré-ordre $\leq_u$ est déterminé par la règle de choix social.

On peut résumer cette correspondance par la figure 8.2.

	Fusion contrainte	Théorie du choix social
individu	φ_i	i
préférence individuelle	$\leq_{\varphi_i}$	$\leq_i$
profil	$\Psi = \{\leq_{\varphi_1}, \dots, \leq_{\varphi_n}\}$	$u = \{\leq_1, \dots, \leq_n\}$
agenda	μ	v
fonction de choix	$\leq_{\Psi}$	C_u
alternatives choisies	$\Delta_{\mu}(\Psi) = \min(\text{Mod}(\mu), \leq_{\Psi})$	$C_u(v)$

FIG. 8.2 – Correspondance fusion contrainte / théorie du choix social

On peut alors examiner les différentes propriétés.

La contrainte du domaine standard est vérifiée. Les deux premiers points de cette condition, comme nous l'avons dit dans la section précédente, ne servent qu'à enlever les cas simples. On peut donc ne considérer que les opérateurs de fusion contrainte opérant sur au moins trois bases de connaissances et avec des contraintes d'intégrité ayant au moins trois modèles.

La condition forte de Pareto est vérifiée. Cela est assuré principalement par les postulats (IC5) et (IC6) des opérateurs de fusion contrainte.

Le fait d'avoir des explications transitives est satisfait trivialement puisque la "règle de choix social" définie par un opérateur de fusion contrainte produit un pré-ordre représentant la préférence globale. Ce pré-ordre est donc une relation d'explication transitive pour la fonction de choix.

L'absence de dictateur est assurée principalement par le postulat d'équité (IC4).

L'unique propriété non vérifiée est l'indépendance des alternatives non disponibles. Il est facile de voir, par exemple, que les opérateurs de la famille $\Delta^{d,\Sigma}$ ne vérifient pas ces propriétés car le "score" de chaque interprétation est calculé en tenant compte de toutes les interprétations (y compris les non disponibles). Nous allons illustrer cela sur l'exemple 18.

Exemple 18 *Supposons que l'on dispose de quatre interprétations $I_0 = \{00\}$, $I_1 = \{01\}$, $I_2 = \{10\}$ et $I_3 = \{11\}$ et de trois bases de connaissance φ_1 , φ_2 et φ_3 telles que $Mod(\varphi_1) = \{I_0\}$, $Mod(\varphi_2) = \{I_1, I_3\}$, et $Mod(\varphi_3) = \{I_0, I_2\}$. Avec l'opérateur $\Delta^{d_H, \Sigma}$ on a les trois pré-ordres de la figure 8.3.*

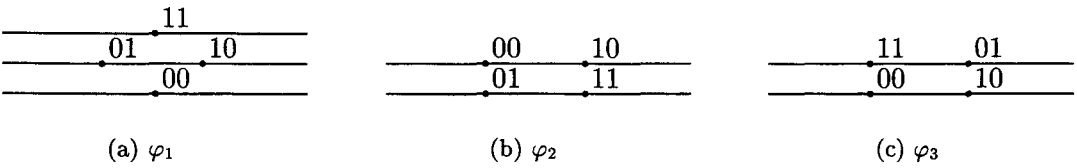


FIG. 8.3 – Pré-ordres

Si on effectue les deux fusions $\Delta^{d_H, \Sigma}_{\{I_0, I_3\}}(\varphi_1 \sqcup \varphi_2) = \varphi_{\{I_0\}}$ et $\Delta^{d_H, \Sigma}_{\{I_0, I_3\}}(\varphi_3 \sqcup \varphi_1) = \varphi_{\{I_0, I_3\}}$ on a bien des résultats différents alors que la restriction des deux profils $\varphi_1 \sqcup \varphi_2$ et $\varphi_3 \sqcup \varphi_2$ à l'agenda $\{I_0, I_3\}$ sont identiques puisque $I_0 <_{\varphi_1} I_1$, $I_0 <_{\varphi_3} I_1$, et $I_1 <_{\varphi_2} I_0$.

Ce qui modifie le résultat est le fait que dans φ_1 la préférence pour I_0 par rapport à I_3 est plus marquée que dans φ_3 .

Cette indépendance des alternatives indépendantes est surtout motivée par l'interdiction de comparer des utilités individuelles que l'on impose en économie. L'utilité peut-être, très sommairement, vue comme le "score" que l'on attribue à chaque alternative, dénotant l'intérêt qu'on lui porte. Mais il est très difficile de travailler avec une notion qui est aussi subjective. En effet comment choisir entre deux alternatives A et B lorsqu'un individu A déclare préférer trois fois plus une alternative x à une alternative y alors qu'un individu B déclare préférer y dix fois plus que x ? L'individu B est peut-être quatre fois plus sensible que l'individu A . Le problème est qu'il n'y a aucune échelle objective derrière cela, et il semble sage dans ce cas de ne pas autoriser cette comparaison et de se limiter à une échelle ordinale (un pré-ordre) n'indiquant pas d'information quantitative sur l'utilité des alternatives. Arrow cite Bentham dans [Arr63] page 11, pour illustrer cette idée :

"Tis in vain to talk of adding quantities which after the addition will continue distinct as they were before, one man's happiness will never be another man's happiness: a gain to one man is no gain to another: you might as well pretend to add 20 apples to 20 pears..."

Mais dans le cas de la fusion contrainte cette utilité individuelle n'est pas posée *a priori* mais calculée uniformément (objectivement) par l'opérateur. Puisque l'on dispose alors de la même "échelle de valeur", il n'y a donc aucune raison de s'interdire cette comparaison. Cela permet donc de différencier l'intensité avec laquelle une interprétation est considérée par une

base de connaissance.

En conclusion, on peut donc dire que l'on échappe au théorème d'impossibilité dans le cadre de la fusion contrainte car on peut se permettre de comparer les utilités, puisque celles-ci sont calculées de manière uniforme. On n'obéit donc pas de ce fait à la condition d'indépendance des alternatives non disponibles.

Chapitre 9

Fusion syntaxique

On étudie dans ce chapitre les propriétés logiques des opérateurs de combinaison de bases de connaissance proposés dans la littérature [BKM91, BKMS92]. On appelle opérateurs de combinaison les opérateurs basés sur une union de toutes les bases de connaissance et sur la sélection de sous-ensembles maximaux consistants pour un ordre donné (pas nécessairement l'inclusion ensembliste).

Nous examinons ces opérateurs à la lumière des propriétés de rationalité proposées section 6.1. Nous illustrons ainsi l'utilisation de ces propriétés pour comparer différents opérateurs de fusion et souligner les qualités et défauts de chacun d'eux. En particulier, un important défaut des opérateurs de combinaison est que la source de chaque connaissance est perdue dans le processus de fusion.

Nous proposons alors l'utilisation de fonctions de sélection pour les opérateurs de combinaison. Ces fonctions sont inspirées de la définition des opérateurs de révision par intersection partielle (section 1.2.1) et permettent de prendre en compte la source de chaque information pour la combinaison. Cela permet donc de définir des opérateurs de combinaison avec de meilleures propriétés.

De manière à motiver le besoin de prendre en compte la source des informations, considérons le scénario suivant : On désire combiner les bases de connaissance suivantes : $\varphi_1 = \varphi_2 = \{a, b\}$, $\varphi_3 = \{a, b \rightarrow c\}$, $\varphi_4 = \{\neg a, d\}$. Alors l'union des bases de connaissance est $\{a, \neg a, b, b \rightarrow c, d\}$. Et, avec un opérateur de combinaison, les ensembles maximaux consistants (pour l'inclusion) seront $\{a, b, b \rightarrow c, d\}$ et $\{\neg a, b, b \rightarrow c, d\}$. Avec un tel résultat, on ne peut pas décider lequel de a ou $\neg a$ est vrai. Pourtant, a est supporté par trois des quatre bases de connaissance, alors qu'une seule supporte $\neg a$. Il peut donc être sensé de souhaiter que a soit vrai dans la base de connaissance résultat de la fusion. Les opérateurs de combinaison usuels ne permettent pas de prendre de tels arguments en compte. Nous montrons comment l'utilisation de fonctions de sélection permet de construire des opérateurs qui sont sensibles à ce genre d'arguments.

Dans ce chapitre, une base de connaissance K est un ensemble fini de formules. Cette base n'est pas nécessairement close par déduction. Un ensemble de connaissance Ψ est toujours un multi-ensemble de bases de connaissance.

Pour les preuves, nous devons définir un opérateur de partition $\oplus$. $K = K' \oplus K''$ signifie

que $K' \cup K'' = K$ et $K' \cap K'' = \emptyset$.

Soit K et K' deux bases de connaissance, $K \wedge K'$ dénotera de manière naturelle la base de connaissance $K \cup K'$. Soit un ensemble de connaissance $\Psi = \{K_1, \dots, K_n\}$, on notera $\Psi \wedge K$ l'ensemble de connaissance $\{K_1 \wedge K, \dots, K_n \wedge K\}$.

Les opérateurs de combinaison étudiés dans ce chapitre auront comme résultat un ensemble de bases de connaissance. Ces ensembles ont été appelés *flocks* par Fagin et al. [FKUV86] (voir également chapitre 4). Notons que *flocks* et ensembles de connaissance sont tous deux des ensembles de bases de connaissance (en fait les ensembles de connaissance sont des multi-ensembles). Mais la différence est que dans le cas des ensembles de connaissance, l'ensemble dénote différentes sources d'information, alors que dans un flock l'ensemble dénote, en quelque sorte, l'incertitude (les alternatives) sur le résultat de la combinaison. Les flocks peuvent être comparés aux extensions dans le cadre de la logique des défauts [Rei80, Eth88]. Pour souligner cette différence, nous noterons K les éléments d'un ensemble de connaissance et M, P, Q, R les éléments d'un flock.

Pour étudier les propriétés logiques des opérateurs de combinaison, nous devons définir ce qu'est le contenu logique (*i.e.* les conséquences) d'un flock. Nous adopterons ce que l'on appelle l'approche prudente en logique des défauts qui considère, dans un sens, un flock comme une disjonction de bases de connaissance.

Définition 102 Soit un flock $\mathcal{F} = \langle M_1, \dots, M_n \rangle$. Si $\mathcal{F} = \emptyset$, alors $\mathcal{F}$ est inconsistant et $Cn(\mathcal{F})$ est l'ensemble de toutes les formules. Sinon on définit

$$Cn(\mathcal{F}) = \bigcap_{i=1}^n Cn(M_i)$$

On peut alors définir l'équivalence entre flocks :

Définition 103 Soient deux flocks $\mathcal{F}$ et $\mathcal{F}'$. On dit que $\mathcal{F}$ implique $\mathcal{F}'$, noté $\mathcal{F} \vdash \mathcal{F}'$ si $Cn(\mathcal{F}) \supseteq Cn(\mathcal{F}')$. $\mathcal{F}$ et $\mathcal{F}'$ sont équivalents, noté $\mathcal{F} \leftrightarrow \mathcal{F}'$, si $\mathcal{F} \vdash \mathcal{F}'$ et $\mathcal{F}' \vdash \mathcal{F}$. Similairement, si K est une base de connaissance, on définit $\mathcal{F} \vdash K$ comme $Cn(\mathcal{F}) \supseteq Cn(K)$.

Définition 104 Soient deux flocks $\mathcal{F} = \langle M_1, \dots, M_n \rangle$ et $\mathcal{F}' = \langle M'_1, \dots, M'_m \rangle$. $\mathcal{F} \vee \mathcal{F}'$ dénote le flock $\langle M_1, \dots, M_n, M'_1, \dots, M'_m \rangle$ et $\mathcal{F} \wedge \mathcal{F}'$ dénote le flock $\langle M''_{1,1}, \dots, M''_{n,m} \rangle$ où $M''_{j,k} = M_j \cup M'_k$.

Notons que si $\mathcal{F} = \langle M_1, \dots, M_n \rangle$ et $\mathcal{F}' = \mathcal{F} \vee \langle P_1, \dots, P_m \rangle$ avec $\forall i P_i$ inconsistant, alors $\mathcal{F} \leftrightarrow \mathcal{F}'$. Donc, lorsque l'on considère un flock, il suffit de considérer ses bases de connaissance consistantes.

9.1 Opérateurs de combinaison

Baral, Kraus, Minker et Subrahmanian ont proposé dans [BKM91, BKMS92] plusieurs opérateurs de *combinaison*, ces opérateurs sont basés sur une sélection de sous-ensembles maximaux consistants de l'union des bases de connaissance de l'ensemble de connaissance.

Une fois l'union des bases de connaissance réalisée, le problème devient alors d'extraire une information cohérente d'une base de connaissance inconsistante. De ce fait, une telle définition

est très proche des sous-théories préférées de Brewka [Bre89] et des travaux de Benferhat et al. sur l'inférence dans les bases de connaissance inconsistantes [BCD⁺93, BDP97, BDL⁺98].

Définition 105 Soient K et IC deux bases de connaissance et Ψ un ensemble de connaissance. $\text{MAXCONS}(K, IC)$ dénote l'ensemble des sous-ensembles maximaux (pour l'inclusion ensembliste) consistants de $K \wedge IC$ qui contiennent IC . On pose alors $\text{MAXCONS}(\Psi, IC) = \text{MAXCONS}(\bigwedge \Psi, IC)$. On utilisera l'indice $\text{MAXCONS}_{\text{card}}(\Psi, IC)$ lorsque la maximalité des ensembles se calculera d'après leur cardinalité.

Cette définition donne les maximaux consistants de Ψ qui impliquent les contraintes IC . Cette définition est à rapprocher de celle de $K \perp A$ pour les révisions par intersection partielle.

On définit alors les opérateurs suivants : soient un ensemble de connaissance Ψ et une base de connaissance IC :

Définition 106

$$\begin{aligned} \Delta^{C^1}_{IC}(\Psi) &= \text{MAXCONS}(\Psi, IC) \\ \Delta^{C^3}_{IC}(\Psi) &= \{M : M \in \text{MAXCONS}(\Psi, \top) \text{ et } M \wedge IC \text{ consistant}\} \\ \Delta^{C^4}_{IC}(\Psi) &= \{M : M \in \text{MAXCONS}_{\text{card}}(\Psi, IC)\} \\ \Delta^{C^5}_{IC}(\Psi) &= \begin{cases} \{M \wedge IC : M \in \text{MAXCONS}(\Psi, \top) \text{ et } M \wedge IC \text{ consistant}\} & \text{si cet ensemble est non vide} \\ IC & \text{sinon} \end{cases} \end{aligned}$$

L'opérateur $\Delta^{C^1}_{IC}(\Psi)$ considère comme résultat de la combinaison l'ensemble des sous-ensembles maximaux consistants de Ψ contenant les contraintes. L'opérateur $\Delta^{C^3}_{IC}(\Psi)$ calcule d'abord l'ensemble des sous-ensembles maximaux de Ψ , puis ne garde ensuite que ceux consistants avec les contraintes. L'opérateur $\Delta^{C^4}_{IC}(\Psi)$ considère l'ensemble des sous-ensembles consistants de cardinalité maximale contenant les contraintes. Les trois opérateurs $\Delta^{C^1}_{IC}(\Psi)$, $\Delta^{C^3}_{IC}(\Psi)$ et $\Delta^{C^4}_{IC}(\Psi)$ correspondent respectivement aux opérateurs $\text{Comb1}(\Psi, IC)$, $\text{Comb3}(\Psi, IC)$ et $\text{Comb4}(\Psi, IC)$ de [BKMS92]. L'opérateur Δ^{C^5} est une légère modification de Δ^{C^3} de manière à capturer plus de propriétés logiques.

Dans [BKMS92], Baral et al. "montrent" les résultats suivants :

Théorème 113 $\Delta^{C^1}_{IC}(\Psi) \supseteq \Delta^{C^3}_{IC}(\Psi)$

Théorème 114 $\Delta^{C^1}_{IC}(\Psi) \supseteq \Delta^{C^4}_{IC}(\Psi)$

Le théorème 114 est aisé à vérifier. Mais il est facile de voir que le théorème 113 est faux puisqu'un élément du flock résultat de $\Delta^{C^3}_{IC}(\Psi)$ ne contient pas forcément IC et n'est donc pas forcément un élément de $\Delta^{C^1}_{IC}(\Psi)$.

Nous étudions à présent les propriétés logiques des opérateurs définis ci-dessus. On peut souligner immédiatement que tous ces opérateurs satisfont (MI) puisque, pour les opérateurs de combinaison, les ensembles de connaissance sont considérés comme des ensembles. Aucun de ces opérateurs ne vérifie (IC3) puisqu'ils sont tous sensibles à la syntaxe des bases de connaissance. Pour illustrer ceci, considérons l'exemple suivant :

Exemple 19 Soient les bases de connaissance $K_1 = \{a, b\}$, $K_2 = \{a \wedge b\}$, et $K_3 = \{\neg b\}$. Considérons les ensembles de connaissance $\Psi_1 = K_1 \sqcup K_3$ et $\Psi_2 = K_2 \sqcup \varphi_3$. Les sous-ensembles maximaux consistants de Ψ_1 sont $\{a, b\}$ et $\{a, \neg b\}$, ceux de Ψ_2 sont $\{a \wedge b\}$ et $\{\neg b\}$.

On constate que tous les sous-ensembles maximaux de Ψ_1 contiennent a , alors que ce n'est pas le cas pour Ψ_2 . Donc, bien que les bases de connaissance K_1 et K_2 soient logiquement équivalentes, avec les opérateurs suivants le résultat de la fusion de Ψ_1 impliquera a , alors que le résultat de la fusion de Ψ_2 non.

Théorème 115 *L'opérateur Δ^{C1} satisfait (IC0), (IC1), (IC2), (IC4), (IC5), (IC7), et (MI). Il ne satisfait pas (IC6), (IC8) et (Maj).*

Preuve :

(IC0), (IC1) et (IC2) sont satisfaits directement par définition de Δ^{C1} .

(IC4) est satisfait puisque Δ^{C1} vérifie une propriété plus forte :

$$\text{Si } K \vdash IC \text{ alors } \Delta^{C1}_{IC}(K \sqcup K') \wedge K \not\vdash \perp$$

Puisque K est un sous-ensemble consistant de $K \wedge K' \wedge IC$ et par hypothèse $K \vdash IC$, alors il existe un sous-ensemble maximal consistant de $K \wedge K' \wedge IC$ qui contient $K \wedge IC$. Et donc $\Delta^{C1}_{IC}(K \sqcup K') \wedge K \not\vdash \perp$.

(IC5) est vérifié. C'est trivialement vrai lorsque $\Delta_{IC}(\Psi_1) \wedge \Delta_{IC}(\Psi_2)$ n'est pas consistant. Et si $\Delta_{IC}(\Psi_1) \wedge \Delta_{IC}(\Psi_2)$ est consistant, notons P_i les éléments de $\text{MAXCONS}(\Psi_1 \sqcup \Psi_2, IC)$, Q_i les éléments de $\text{MAXCONS}(\Psi_1, IC)$, et R_i les éléments de $\text{MAXCONS}(\Psi_2, IC)$. Pour montrer que (IC5) est vérifié il suffit de montrer que si $Q_j \wedge R_k$ est consistant alors $\exists i P_i = Q_j \wedge R_k$. Posons d'abord $L = IC \cup \bigwedge \Psi_1 \cup \bigwedge \Psi_2$. $Q_j \wedge R_k$ est un sous-ensemble consistant de L (contenant IC), et comme les P_i sont des sous-ensembles maximaux consistants de L (contenant IC), alors $\exists i P_i \supseteq Q_j \wedge R_k$. On a de même $P_i \subseteq Q_j \wedge R_k$. Car sinon $P_i \supset Q_j \wedge R_k$. Alors on peut décomposer $P_i = P_1 \wedge P_2$ avec $P_1 = P_i \cap (\bigwedge \Psi_1 \wedge IC)$ et $P_2 = P_i \cap (\bigwedge \Psi_2 \wedge IC)$. Et alors on a que $P_1 \supset Q_j$ ou $P_2 \supset R_k$. Donc Q_j ou R_k n'est pas un sous-ensemble maximal consistant de Ψ_1 ou Ψ_2 respectivement. Contradiction.

(IC6) n'est pas vérifié. Considérons l'exemple suivant : $K_1 = \{a \rightarrow c, e \rightarrow \neg c, b \rightarrow \neg c\}$, $K_2 = \{a, e\}$, $K_3 = \{b\}$ et $IC = \top$. Alors $\Delta^{C1}(K_1 \sqcup K_2) \wedge \Delta^{C1}(K_3)$ est consistant mais $\Delta^{C1}(K_1 \sqcup K_2 \sqcup K_3) \not\vdash \Delta^{C1}(K_1 \sqcup K_2) \wedge \Delta^{C1}(K_3)$.

(IC7) est satisfait. Lorsque $\Delta_{IC_1}(\Psi) \wedge IC_2$ n'est pas consistant (IC7) est trivial. Sinon nommons P_i les éléments de $\text{MAXCONS}(\Psi, IC_1)$ et Q_i les éléments de $\text{MAXCONS}(\Psi, IC_1 \wedge IC_2)$. Il suffit de montrer que si $P_i \wedge IC_2$ est consistant alors $\exists Q_j$ tel que $P_i \wedge IC_2 = Q_j$. Soit $P_i \wedge IC_2$ un sous-ensemble consistant de $\bigwedge \Psi \wedge IC_1 \wedge IC_2$, alors il existe un sous-ensemble maximal consistant Q_j qui contient $P_i \wedge IC_2$, donc $\exists j P_i \subseteq Q_j$. De plus, on a que $P_i \wedge IC_2 \supseteq Q_j$, sinon $P_i \wedge IC_2 \subset Q_j$, et de là il est facile de voir que $P_i \subset Q_j \cap (\bigwedge \Psi \wedge IC_1)$. Donc P_i n'est pas un maximal. Contradiction.

(IC8) n'est pas vérifié. On utilise le contre-exemple à (IC6) légèrement modifié :

$K = \{a \rightarrow c, e \rightarrow \neg c, b \rightarrow \neg c\}$, $IC_1 = \{a, e\}$ et $IC_2 = \{b\}$. Alors $\Delta^{C1}_{IC_1}(K) \wedge IC_2$ est consistant mais $\Delta^{C1}_{IC_1 \wedge IC_2}(K) \not\vdash \Delta^{C1}_{IC_1}(K) \wedge IC_2$.

(Maj) n'est pas satisfait. C'est une conséquence du théorème 73.

□

Théorème 116 *L'opérateur Δ^{C3} satisfait (IC4), (IC5), (IC7), (IC8) et (MI). Il ne satisfait pas (IC0), (IC1), (IC2), (IC6) et (Maj).*

Preuve :

- (IC0) n'est pas satisfait, puisque le résultat est seulement consistant avec IC .
- (IC1) n'est pas satisfait puisque lorsqu'il n'y a pas de maximal consistant avec IC le flock est vide et le résultat est $\perp$.
- (IC2) n'est pas satisfait, on a le plus faible : si $\bigwedge \Psi$ est consistant avec IC , alors $\Delta^{C^3}_{IC}(\Psi) = \bigwedge \Psi$.
- (IC4) et (IC5) sont satisfaits. Les preuves sont similaires à celles de Δ^{C^1} .
- (IC6) n'est pas satisfait. On a le même contre-exemple que pour Δ^{C^1} .
- (IC7) et (IC8) sont satisfaits directement. Si $\Delta^{C^3}_{IC_1}(\Psi) \wedge IC_2$ n'est pas consistant (IC7) est trivial. Sinon puisque par définition $\Delta^{C^3}_{IC_1}(\Psi) = \{Q : Q \in \text{MAXCONS}(\Psi, \top) \text{ and } Q \wedge IC_1 \text{ consistant}\}$, alors $\Delta^{C^3}_{IC_1}(\Psi) \wedge IC_2 \leftrightarrow \{Q : Q \in \text{MAXCONS}(\Psi, \top) \text{ and } Q \wedge IC_1 \wedge IC_2 \text{ consistant}\}$ (cet ensemble est non vide par hypothèse) ce qui est, par définition $\Delta^{C^3}_{IC_1 \wedge IC_2}(\Psi)$.
- (Maj) n'est pas vérifié. Il suffit de considérer $IC = \top$ et $\Psi = K_1 \sqcup K_2$ avec $K_1 = \{a\}$ et $K_2 = \{\neg a\}$. On a alors que $\forall n \Delta^{C^3}_{\top}(K_1 \sqcup K_2^n) = \Delta^{C^3}_{\top}(K_1 \sqcup K_2) = \{\{a\}, \{\neg a\}\}$ ce qui n'implique pas $\neg a$.

□

Théorème 117 *L'opérateur Δ^{C^4} satisfait (IC0), (IC1), (IC2), (IC7), (IC8), (MI). Il ne satisfait pas (IC4), (IC5), (IC6) et (Maj).*

Preuve :

- (IC0) et (IC1) sont satisfaits par définition.
- (IC2) est satisfait car lorsque Ψ est consistant avec IC , il n'y a qu'un seul maximal consistant $\bigwedge \Psi \wedge IC$.
- (IC4) n'est pas vérifié : Prenons $IC = \{a, b\}$, $K = \{a, b, c, d\}$, $K' = \{a, b, c \rightarrow \neg d, d \rightarrow \neg c\}$. Alors $\Delta^{C^4}_{IC}(K \sqcup K') = \{\{a, b, c \rightarrow \neg d, d \rightarrow \neg c\}, \{a, b, d, c \rightarrow \neg d, d \rightarrow \neg c\}\}$, d'où $\Delta^{C^4}_{IC}(K \sqcup K') \wedge K' \not\vdash \perp$ et $\Delta^{C^4}_{IC}(K \sqcup K') \wedge K \vdash \perp$.
- (IC5) n'est pas vérifié. Soit l'exemple suivant : Soient $K_1 = \{a\}$, $K_2 = \{\neg a \wedge b\}$, $K_3 = \{\neg a \wedge c\}$. Alors avec $\Psi_1 = K_1 \sqcup K_2$, $\Psi_2 = K_1 \sqcup K_3$ et $IC = \top$, on a $\Delta^{C^4}_{IC}(\Psi_1) \wedge \Delta^{C^4}_{IC}(\Psi_2) \not\vdash \Delta^{C^4}_{IC}(\Psi_1 \sqcup \Psi_2)$.
- (IC6) n'est pas vérifié. Considérons l'exemple suivant : Soient $K_1 = \{a\}$, $K_2 = \{\neg a \wedge (b \vee c)\}$, $K_3 = \{a, a \wedge z\}$, $K_4 = \{\neg a \wedge \neg b, \neg a \wedge \neg c\}$. Alors avec $\Psi_1 = K_1 \sqcup K_2$, $\Psi_2 = K_3 \sqcup K_4$ et $IC = \top$, on a $\Delta^{C^4}_{IC}(\Psi_1) \wedge \Delta^{C^4}_{IC}(\Psi_2)$ consistant mais $\Delta^{C^4}_{IC}(\Psi_1 \sqcup \Psi_2) \not\vdash \Delta^{C^4}_{IC}(\Psi_1) \wedge \Delta^{C^4}_{IC}(\Psi_2)$.
- (IC7) et (IC8) sont vérifiés. Lorsque $\Delta^{C^4}_{IC_1}(\Psi) \wedge IC_2$ n'est pas consistant (IC7) est vérifié trivialement. Supposons donc que $\Delta^{C^4}_{IC_1}(\Psi) \wedge IC_2$ est consistant. Notons P un élément de $\text{MAXCONS}_{card}(\Psi, IC_1 \wedge IC_2)$. Et soit Q un élément de $\text{MAXCONS}_{card}(\Psi, IC_1)$ consistant avec IC_2 . Nous voulons montrer que $Q \cup IC_2 \in \text{MAXCONS}_{card}(\Psi, IC_1 \wedge IC_2)$ et que P peut s'écrire $P_1 \oplus P_2$ avec $P_1 \in \text{MAXCONS}_{card}(\Psi, IC_1)$ et $P_2 \subseteq IC_2$. Ce qui est suffisant pour montrer que $\Delta^{C^4}_{IC_1}(\Psi) \wedge IC_2 \leftrightarrow \Delta^{C^4}_{IC_1 \wedge IC_2}(\Psi)$.
On définit $A_1 = \bigwedge \Psi \cup IC_1$ et $A_2 = IC_2 \setminus A_1$. On peut alors écrire $P = P_1 \oplus P_2$ avec $P_1 = P \cap A_1$ et $P_2 = P \cap A_2$. Similairement on définit $Q \cup IC_2 = Q_1 \oplus Q_2$ tel que

$Q_1 = (Q \cup IC_2) \cap A_1$ et $Q_2 = (Q \cup IC_2) \cap A_2$. Comme $IC_2 \subseteq P$ et $IC_2 \subseteq Q \cup IC_2$, par construction il est facile de voir que $P_2 = A_2 = Q_2$. En termes de cardinalités $|P| = |P_1| + |P_2|$ et $|Q \cup IC_2| = |Q_1| + |Q_2|$. Mais on a que $|P_2| = |Q_2|$. Comme P est dans $\text{MAXCONS}_{card}(\Psi, IC_1 \wedge IC_2)$ et $Q \cup IC_2 \subseteq \Psi \cup IC_1 \cup IC_2$, alors $|P| \geq |Q \cup IC_2|$, d'où $|P_1| \geq |Q_1|$. De la même manière comme $Q_1 = Q$ est dans $\text{MAXCONS}_{card}(\Psi, IC_1)$ et $P_1 \subseteq \Psi \cup IC_1$, alors $|P_1| \leq |Q_1|$. Il est alors facile de voir que l'on a $|P| = |Q \cup IC_2|$ et $|Q| = |P_1|$. Donc $Q \cup IC_2$ est dans $\text{MAXCONS}_{card}(\Psi, IC_1 \wedge IC_2)$, et P_1 est dans $\text{MAXCONS}_{card}(\Psi, IC_1)$.

(Maj) n'est pas satisfait. C'est une conséquence du théorème 73.

□

Théorème 118 *L'opérateur Δ^{C5} satisfait (IC0), (IC1), (IC2), (IC4), (IC5), (IC7), (IC8) et (MI). Il ne satisfait pas (IC6) et (Maj).*

La preuve de ce théorème est la même que pour Δ^{C3} , mis à part pour (IC0), (IC1) et (IC2) qui sont vérifiés à présent par définition.

Les résultats des théorèmes précédents sont résumés dans le tableau 9.1. Les opérateurs de combinaison ne satisfont pas (IC3), ce qui est normal vu que la syntaxe des bases de connaissance intervient sur le résultat de la combinaison. Mais on pourrait espérer qu'ils satisfassent les autres postulats. Or ce n'est pas le cas, et on peut noter en particulier qu'aucun des opérateurs ne satisfait (IC6). De plus le fait qu'ils satisfassent (MI) n'est pas un point positif pour la fusion puisque cela implique que ces opérateurs "oublient" l'origine des informations.

	IC0	IC1	IC2	IC3	IC4	IC5	IC6	IC7	IC8	MI	Maj
Δ^{C1}	✓	✓	✓		✓	✓		✓		✓	
Δ^{C3}					✓	✓		✓	✓	✓	
Δ^{C4}	✓	✓	✓					✓	✓	✓	
Δ^{C5}	✓	✓	✓		✓	✓		✓	✓	✓	

TAB. 9.1 – Propriétés des opérateurs de combinaison

9.2 Fonctions de sélection

Les opérateurs de combinaison ne prennent pas en compte l'aspect "individuel" de la fusion, puisque les sources d'information n'interviennent pas dans le processus de combinaison. Il mettent simplement toutes les informations ensemble et sélectionnent alors certains sous-ensembles maximaux consistants. Donc, avec cette approche, il n'est pas possible de tenter de satisfaire au mieux les protagonistes. Il n'est pas possible, par exemple, de sélectionner seulement les maximaux consistants qui satisfont la majorité des bases. De la même manière on ne peut pas tenter de définir un arbitrage, c'est-à-dire de satisfaire chaque protagoniste du mieux possible. L'idée de cette section est d'utiliser une fonction de sélection pour choisir parmi les sous-ensembles maximaux consistants, ceux qui satisfont au mieux un critère de

fusion.

La motivation pour la définition de ces fonctions de sélection vient des fonctions de sélection dans le cadre de la révision AGM (voir section 1.2.1). En effet, l'opérateur Δ^{C1} correspond à l'opérateur de contraction par intersection totale, qui est considéré comme insatisfaisant pour la révision, puisqu'il abandonne trop d'informations. Les opérateurs de contraction par intersection partielle, définis à partir de fonctions de sélection, qui choisissent seulement certains maximaux consistants, ont un comportement beaucoup plus satisfaisant.

Dans cette section on étudie quelques opérateurs de fusion avec fonctions de sélection. L'idée de cette famille d'opérateurs est de donner la préférence aux maximaux consistants qui sont les plus proches des bases de connaissance. La différence entre les opérateurs réside d'abord dans la définition de la "distance" entre un maximal consistant et une base de connaissance et ensuite dans la méthode d'agrégation de ces résultats en une "distance" entre un maximal consistant et un ensemble de connaissance.

On se focalisera sur les opérateurs définis à partir de l'opérateur Δ^{C1} dont nous étudierons les propriétés logiques. Mais les méthodes données ci-dessous peuvent également être utilisées avec les autres opérateurs de combinaison.

9.2.1 Opérateur majoritaire drastique

La "distance" entre bases de connaissance que l'on considère ici est drastique. Cette distance vaut 0 si la conjonction des deux bases de connaissance est consistante et 1 sinon.

Définition 107 Soient un ensemble de connaissance Ψ et une base de connaissance M . La distance entre un maximal consistant M et une base de connaissance K est :

$$dist_D(M, K) = \begin{cases} 0 & \text{si } M \wedge K \text{ consistant} \\ 1 & \text{sinon} \end{cases}$$

La distance entre un maximal consistant M et un ensemble de connaissance Ψ est alors :

$$dist_D(M, \Psi) = \sum_{K \in \Psi} dist_D(M, K)$$

Et l'opérateur majoritaire drastique est défini par :

$$\Delta_{IC}^D(\Psi) = \{M \in \Delta_{IC}^{C1}(\Psi) : dist_D(M, \Psi) = \min_{M_i \in \Delta_{IC}^{C1}(\Psi)} (dist_D(M_i, \Psi))\}$$

Donc cette fonction de sélection choisit parmi les maximaux consistants ceux qui sont consistants avec le maximum de bases de connaissance de l'ensemble de connaissance.

Il est facile de voir que si un maximal consistant est consistant avec une base de connaissance de l'ensemble de connaissance, alors il contient cette base de connaissance. Cela revient donc alors à sélectionner les maximaux consistants qui contiennent le maximum de bases de l'ensemble de connaissance.

Théorème 119 L'opérateur Δ_{IC}^D vérifie (IC0), (IC1), (IC2), (IC4), (IC5), (IC7) et (Maj). Il ne vérifie pas (IC6), (IC8) et (MI).

Preuve :

(IC0), (IC1) et (IC2) sont satisfaits directement.

- (IC4) est satisfait. Puisque soit K est consistant avec K' et alors (IC4) est vérifié, ou K n'est pas consistant avec K' . Dans ce cas il n'existe pas de maximal consistant consistant avec les deux bases de connaissance, comme il existe un maximal consistant qui contient K on a $\Delta_{IC}^D(K \sqcup K') \wedge K \not\models \perp$.
- (IC5) est vérifié. Notons Q_i les éléments de $\Delta_{IC}^D(\Psi_1)$, et R_j les éléments de $\Delta_{IC}^D(\Psi_2)$. Remarquons que si Q_i est consistant avec $K \in \Psi_1$, alors $K \subseteq Q_i$ et donc si $Q_i \wedge R_j$ est consistant, alors $Q_i \wedge R_j \wedge K$ est consistant. Donc si $K \in \Psi_1$, alors $\text{dist}_D(Q_i \wedge R_j, K) = \text{dist}_D(Q_i, K)$. Et de même pour $K \in \Psi_2$, $\text{dist}_D(Q_i \wedge R_j, K) = \text{dist}_D(R_j, K)$. Donc si $Q_i \wedge R_j$ est consistant, alors $\text{dist}_D(Q_i, \Psi_1) + \text{dist}_D(R_j, \Psi_2) = \text{dist}_D(Q_i \wedge R_j, \Psi_1 \sqcup \Psi_2)$. Puisque nous savons que Δ^{C1} satisfait (IC5), on a que $Q_i \wedge R_j \in \text{MAXCONS}(\Psi_1 \sqcup \Psi_2, IC)$. Il reste à montrer que $\text{dist}_D(Q_i \wedge R_j, \Psi_1 \sqcup \Psi_2)$ est minimal. Si ce n'est pas le cas $\exists P \in \Delta_{IC}^D(\Psi_1 \sqcup \Psi_2)$ tel que $\text{dist}_D(P, \Psi_1 \sqcup \Psi_2) < \text{dist}_D(Q_i \wedge R_j, \Psi_1 \sqcup \Psi_2)$. Donc soit $\text{dist}_D(P, \Psi_1) < \text{dist}_D(Q_i \wedge R_j, \Psi_1)$, soit $\text{dist}_D(P, \Psi_2) < \text{dist}_D(Q_i \wedge R_j, \Psi_2)$. Supposons, sans perte de généralité, que $\text{dist}_D(P, \Psi_1) < \text{dist}_D(Q_i \wedge R_j, \Psi_1)$, alors $\text{dist}_D(P \cap (\bigwedge \Psi_1 \wedge IC), \Psi_1) < \text{dist}_D(Q_i, \Psi_1)$. Donc $Q_i \notin \Delta_{IC}^D(\Psi_1)$. Contradiction.
- (IC6) et (IC8) ne sont pas vérifiés. Les contre-exemples utilisés pour Δ^{C1} fonctionnent ici aussi.
- (IC7) est vérifié. La preuve est exactement la même que celle du théorème 115, puisque ajouter les contraintes d'intégrité IC_2 ne change pas le "score" de chaque maximal consistant.
- (MI) n'est pas satisfait. C'est une conséquence du théorème 73.

□

9.2.2 Les opérateurs cardinaux

L'opérateur précédent n'est pas très fin, car l'évaluation des maximaux consistant par les bases de connaissance est manichéenne.

On peut imaginer des moyens plus subtils d'évaluer les maximaux consistants. Un tel problème a déjà été étudié dans la littérature. Par exemple, dans le cas des mises à jour de bases de données, Fagin et al. [FUV83, FKUV86] ont proposé une notion de moindre changement (*fewer change*):

Définition 108 Soient trois bases de connaissance K_1, K_2 et K .

1. K_1 a moins d'insertions que K_2 par rapport à K si $K_1 \setminus K \subset K_2 \setminus K$.
2. K_1 a moins de suppressions que K_2 par rapport à K si $K \setminus K_1 \subset K \setminus K_2$.
3. K_1 a moins de changements que K_2 par rapport à K si K_1 a moins de suppressions que K_2 , ou K_1 et K_2 ont les mêmes suppressions et K_1 a moins d'insertions que K_2 .

Le problème avec cette définition de moindre changement est qu'elle ne donne qu'un ordre partiel. Ici nous avons besoin d'un pré-ordre total de façon à agréger ces préférences (croyances) individuelles en une préférence sociale.

Fagin et al. donnent plus d'importance aux suppressions qu'aux insertions. Même si cela peut être justifié pour la mise à jour de bases de données par la volonté de garder un maximum de l'ancienne connaissance, cela semble contredire le principe de "changement minimal", puisqu'un ensemble qui n'effectuerait aucune suppression, mais ajouterait des centaines de formules serait considéré meilleur qu'un ensemble qui accomplirait une suppression mais aucune insertion!

Pour la fusion, il semble naturel que nous ayons à porter la même importance aux insertions et aux suppressions.

Cela mène à l'opérateur suivant.

L'opérateur différence symétrique

L'opérateur suivant est défini à partir d'une distance qui représente la cardinalité de la différence symétrique entre les bases de connaissance.

Définition 109 Soient un ensemble de connaissance Ψ et une base de connaissance M . La distance différence symétrique entre un maximal consistant M et une base de connaissance K est

$$\text{dist}_S(M, K) = |K \setminus M| + |M \setminus K|$$

La distance entre un maximal consistant M et un ensemble de connaissance Ψ est alors :

$$\text{dist}_S(M, \Psi) = \sum_{K \in \Psi} \text{dist}_S(M, K)$$

Et l'opérateur différence symétrique est défini comme

$$\Delta_{IC}^{S, \Sigma}(\Psi) = \{M \in \Delta_{IC}^{C1}(\Psi) : \text{dist}_S(M, \Psi) = \min_{M_i \in \Delta_{IC}^{C1}(\Psi)} (\text{dist}_S(M_i, \Psi))\}$$

Théorème 120 L'opérateur $\Delta_{IC}^{S, \Sigma}$ satisfait (IC0), (IC1), (IC2), (IC4), (IC7), (IC8) et (Maj). Il ne satisfait pas (IC5), (IC6) et (MI).

Preuve :

(IC0), (IC1) et (IC2) sont satisfaits trivialement.

(IC4) est satisfait. Cela vient de la propriété suivante : $\forall M, K \cap K' \subseteq M \subseteq K \cup K' \text{ dist}_S(M, K \sqcup K') = |K| + |K'| - 2|K \cap K'|$. Alors si $\Delta_{IC}^{S, \Sigma}(K \sqcup K') \wedge K \not\models \perp$, c'est-à-dire K est consistant avec un maximal consistant M . Cela implique que $K \subseteq M$, donc par la propriété ci-dessus $\text{dist}_S(M, K \sqcup K') = |K| + |K'| - 2|K \cap K'|$. Mais K' est un sous-ensemble consistant de $K \cup K'$, donc il existe un élément $M' \in \Delta_{IC}^{C1}(K \sqcup K')$ tel que $K' \subseteq M'$. Par la même propriété on obtient $\text{dist}_S(M', K \sqcup K') = \text{dist}_S(M, K \sqcup K')$, donc $M' \in \Delta_{IC}^{S, \Sigma}(K \sqcup K')$. donc $\Delta_{IC}^{S, \Sigma}(K \sqcup K') \wedge K' \not\models \perp$.

(IC5) n'est pas vrai pour $\Delta_{IC}^{S, \Sigma}$. Considérons les bases de connaissance $K_1 = \{a, a \rightarrow b \wedge c\}$, $K_2 = \{\neg c\}$ et $K_3 = \{a \rightarrow \neg c\}$. Et posons $\Psi_1 = K_1 \sqcup K_2$, $\Psi_2 = K_3$ et $IC = \top$.

(IC6) n'est pas satisfait. Le contre-exemple pour Δ_{IC}^{C1} fonctionne ici aussi.

(IC7) et (IC8) sont satisfaits. Lorsque $\Delta_{IC_1}^{S, \Sigma}(\Psi) \wedge IC_2$ n'est pas consistant (IC7) est satisfait trivialement. Supposons donc $\Delta_{IC_1}^{S, \Sigma}(\Psi) \wedge IC_2$ consistant. Notons P un élément de $\Delta_{IC_1 \wedge IC_2}^{S, \Sigma}(\Psi)$, et soit Q un élément de $\Delta_{IC_1}^{S, \Sigma}(\Psi)$ consistant avec IC_2 . On définit $A_1 = \bigwedge \Psi \wedge IC_1$ et $A_2 = IC_2 \setminus A_1$. On peut alors partager $P = P_1 \oplus P_2$ avec $P_1 = P \cap A_1$ et $P_2 = P \cap A_2$. De même on définit $Q \cup IC_2 = Q_1 \oplus Q_2$ tel que $Q_1 = (Q \cup IC_2) \cap A_1$ et $Q_2 = (Q \cup IC_2) \cap A_2$. Comme $IC_2 \subseteq P$ et $IC_2 \subseteq Q \cup IC_2$, par construction il est aisé de voir que $P_2 = A_2 = Q_2$. Soit $K \in \Psi$ on a $\text{dist}_S(P, K) = \text{dist}_S(P_1, K) + |P_2 \setminus K| - |P_2 \cap K|$ et $\text{dist}_S(Q \cup IC_2, K) = \text{dist}_S(Q_1, K) + |Q_2 \setminus K| - |Q_2 \cap K|$. On a que $|P_2 \setminus K| - |P_2 \cap K| = |Q_2 \setminus K| - |Q_2 \cap K|$. Comme P est dans $\Delta_{IC_1 \wedge IC_2}^{S, \Sigma}(\Psi)$ et $Q \cup IC_2 \subseteq \Psi \cup IC_1 \cup IC_2$, alors $\text{dist}_S(P, K) \leq \text{dist}_S(Q \cup IC_2, K)$. De même comme

$Q_1 = Q$ est dans $\Delta_{IC_1}^{S,\Sigma}(\Psi)$ et $P_1 \subseteq \Psi \cup IC_1$, alors $dist_S(P_1, K) \geq dist_S(Q, K)$. On en conclut alors $dist_S(P, K) = dist_S(Q \cup IC_2, K)$ et $dist_S(Q, K) = dist_S(P_1, K)$. D'où $Q \cup IC_2$ est dans $\Delta_{IC_1 \wedge IC_2}^{S,\Sigma}(\Psi)$, et P_1 appartient à $\Delta_{IC_1}^{S,\Sigma}(\Psi)$.

(MI) n'est pas satisfait. C'est une conséquence du théorème 73.

□

Cet opérateur ne semble pas satisfaire beaucoup de propriétés. Mais plutôt que de se focaliser sur les différences entre les bases de connaissance, on peut s'intéresser à ce qu'il y a en commun entre ces bases. Ces deux approches sont très proches et semblent duales mais il se trouve que cette dernière est plus intéressante d'un point de vue logique.

L'opérateur Intersection

Cet opérateur est défini à partir d'une distance qui représente la cardinalité de l'intersection entre les bases de consistances.

Définition 110 Soient un ensemble de connaissance Ψ et une base de connaissance M . La distance entre un maximal consistant M et une base de connaissance K est

$$dist_{\cap}(M, K) = |K \cap M|$$

La distance entre un maximal consistant M et un ensemble de connaissance Ψ est alors :

$$dist_{\cap}(M, \Psi) = \sum_{K \in \Psi} dist_{\cap}(M, K)$$

L'opérateur $\Delta^{\cap, \Sigma}$ est défini par :

$$\Delta_{IC}^{\cap, \Sigma}(\Psi) = \{M \in \Delta_{IC}^C(\Psi) : dist_{\cap}(M, \Psi) = \max_{M_i \in \Delta_{IC}^{C1}(\Psi)} (dist_{\cap}(M_i, \Psi))\}$$

Théorème 121 L'opérateur $\Delta_{IC}^{\cap, \Sigma}$ vérifie (IC0), (IC1), (IC2), (IC5), (IC6), (IC7), (IC8) et (Maj). Il ne vérifie pas (IC4) et (MI).

Preuve :

(IC0), (IC1) et (IC2) sont satisfaits trivialement.

(IC4) n'est pas satisfait. Considérons l'exemple suivant : $K = \{a, b\}$ et $K' = \{\neg a \wedge \neg b\}$. Alors $\Delta_{\top}^{\cap, \Sigma}(K \sqcup K') = K$ et $\Delta_{\top}^{\cap, \Sigma}(K \sqcup K') \wedge K' \vdash \perp$.

(IC5) est vérifié. Le résultat est direct si $\Delta_{IC}^{\cap, \Sigma}(\Psi_1) \wedge \Delta_{IC}^{\cap, \Sigma}(\Psi_2)$ n'est pas consistant. Sinon il existe Q_i un élément de $\Delta_{IC}^{\cap, \Sigma}(\Psi_1)$ et R_j un élément de $\Delta_{IC}^{\cap, \Sigma}(\Psi_2)$ tel que $Q_i \wedge R_j$ est consistant. Notons que si $K_l \in \Psi_1$, alors $dist_{\cap}(Q_i \wedge R_j, K_l) = dist_{\cap}(Q_i, K_l)$. Car si ce n'est pas le cas, c'est-à-dire $dist_{\cap}(Q_i \wedge R_j, K_l) > dist_{\cap}(Q_i, K_l)$ alors il existe une formule $\alpha \in \Psi_1$ telle que $\alpha \notin Q_i$ et $Q_i \cup \alpha$ est consistant. Donc Q_i n'est pas un maximal consistant. Contradiction. De même si $K_l \in \Psi_2$, alors $dist_{\cap}(Q_i \wedge R_j, K_l) = dist_{\cap}(R_j, K_l)$. Donc si $Q_i \wedge R_j$ est consistant, alors $dist_{\cap}(Q_i \wedge R_j, \Psi_1 \sqcup \Psi_2) = dist_{\cap}(Q_i, \Psi_1) + dist_{\cap}(R_j, \Psi_2)$. D'après les propriétés de Δ^{C1} on sait que $Q_i \wedge R_j$ appartient à $MAXCONS(\Psi_1 \cup \Psi_2, IC)$. Il reste à montrer que $dist_{\cap}(Q_i \wedge R_j, \Psi)$ est maximum. Si ce n'est pas le cas $\exists P \in \Delta_{IC}^{\cap, \Sigma}(\Psi_1 \sqcup \Psi_2)$ tel que $dist_{\cap}(P, \Psi_1 \sqcup \Psi_2) > dist_{\cap}(Q_i \wedge R_j, \Psi_1 \sqcup \Psi_2)$. Donc soit $dist_{\cap}(P, \Psi_1) >$

$dist_{\cap}(Q_i \wedge R_j, \Psi_1)$, soit $dist_{\cap}(P, \Psi_2) > dist_{\cap}(Q_i \wedge R_j, \Psi_2)$. Supposons, sans perte de généralité, que $dist_{\cap}(P, \Psi_1) > dist_{\cap}(Q_i \wedge R_j, \Psi_1)$, alors $dist_{\cap}(P \cap (\Psi_1 \cup IC), \Psi_1) > dist_{\cap}(Q_i, \Psi_1)$. D'où $Q_i \notin \Delta_{IC}^{\cap, \Sigma}(\Psi_1)$. Contradiction.

(IC6) est vérifié. Notons P un élément de $\Delta_{IC}^{\cap, \Sigma}(\Psi_1 \sqcup \Psi_2)$. On décompose $P = P_1 \cup P_2$ avec $P_1 = P \cap (IC \cup \Psi_1)$ et $P_2 = P \cap (IC \cup \Psi_2)$. Si $P_1 \notin \Delta_{IC}^{\cap, \Sigma}(\Psi_1)$, alors il existe $Q \in \Delta_{IC}^{\cap, \Sigma}(\Psi_1)$ et $R \in \Delta_{IC}^{\cap, \Sigma}(\Psi_2)$ tels que $Q \wedge R$ est consistant et par définition $dist_{\cap}(Q, \Psi_1) > dist_{\cap}(P_1, \Psi_1)$ et $dist_{\cap}(R, \Psi_2) \geq dist_{\cap}(P_2, \Psi_2)$. Donc avec un argument similaire à celui utilisé pour (IC5) on a que $dist_{\cap}(Q, \Psi_1 \sqcup \Psi_2) > dist_{\cap}(P, \Psi_1 \sqcup \Psi_2)$. Donc $P \notin \Delta_{IC}^{\cap, \Sigma}(\Psi_1 \sqcup \Psi_2)$. Contradiction.

(IC7) et (IC8) sont satisfaits. La preuve est similaire à celle de $\Delta^{S, \Sigma}$. Lorsque $\Delta_{IC_1}^{\cap, \Sigma}(\Psi) \wedge IC_2$ n'est pas consistant (IC7) est satisfait trivialement. Supposons donc que $\Delta_{IC_1}^{\cap, \Sigma}(\Psi) \wedge IC_2$ est consistant. Soit P un élément de $\Delta_{IC_1 \wedge IC_2}^{\cap, \Sigma}(\Psi)$. Et soit Q un élément de $\Delta_{IC_1}^{\cap, \Sigma}(\Psi)$ consistant avec IC_2 . On pose $A_1 = \bigwedge \Psi \wedge IC_1$ et $A_2 = IC_2 \setminus A_1$. On peut alors décomposer $P = P_1 \oplus P_2$ avec $P_1 = P \cap A_1$ et $P_2 = P \cap A_2$. De même $Q \cup IC_2 = Q_1 \oplus Q_2$ tel que $Q_1 = (Q \cup IC_2) \cap A_1$ et $Q_2 = (Q \cup IC_2) \cap A_2$. Comme $IC_2 \subseteq P$ et $IC_2 \subseteq Q \cup IC_2$, par construction il est facile de voir $P_2 = A_2 = Q_2$. En terme de distances $dist_{\cap}(P, K) = dist_{\cap}(P_1, K) + dist_{\cap}(P_2, K)$ et $dist_{\cap}(Q, K) = dist_{\cap}(Q_1, K) + dist_{\cap}(Q_2, K)$. Mais on a que $dist_{\cap}(P_2, K) = dist_{\cap}(Q_2, K)$. Comme P appartient à $\Delta_{IC_1 \wedge IC_2}^{\cap, \Sigma}(\Psi)$ et $Q \cup IC_2 \subseteq \Psi \cup IC_1 \cup IC_2$, alors $dist_{\cap}(P, K) \geq dist_{\cap}(Q \cup IC_2, K)$. De même $Q_1 = Q$ est dans $\Delta_{IC_1}^{\cap, \Sigma}(\Psi)$ et $P_1 \subseteq \Psi \cup IC_1$, donc $dist_{\cap}(P_1, K) \leq dist_{\cap}(Q, K)$. On en déduit donc $dist_{\cap}(P, K) = dist_{\cap}(Q \cup IC_2, K)$ et $dist_{\cap}(Q, K) = dist_{\cap}(P_1, K)$. D'où $Q \cup IC_2$ appartient à $\Delta_{IC_1 \wedge IC_2}^{\cap, \Sigma}(\Psi)$, et P_1 appartient à $\Delta_{IC_1}^{\cap, \Sigma}(\Psi)$.

(MI) n'est pas satisfait. C'est une conséquence du théorème 73.

□

Les propriétés des opérateurs définis ci-dessus sont résumés dans le tableau 9.2.2.

Il est donc clair que l'on peut capturer davantage de propriétés logiques avec des fonctions de sélection adéquates. Les différents opérateurs que nous avons définis dans cette section satisfont tous (Maj) et sont donc plus à même d'être utilisés comme opérateurs de fusion que l'opérateur Δ^{C1} à partir duquel ils sont définis. On peut noter, en particulier, que l'opérateur $\Delta^{\cap, \Sigma}$ satisfait pratiquement toutes les propriétés de la fusion contrainte. (IC3) n'est forcément pas satisfait, donc le seul postulat "manquant" est (IC4). Remarquons que le seul opérateur satisfaisant autant de propriétés de base est Δ^{C5} , mais le postulat non vérifié par $\Delta^{\cap, \Sigma}$ est (IC4), alors que celui non vérifié par Δ^{C5} est (IC6). Or (IC5) et (IC6) sont les postulats qui traitent purement de l'agrégation des connaissances. Ce manquement semble donc plus "grave" que de ne pas satisfaire (IC4). Notons d'ailleurs que le comportement de Δ^{C5} est beaucoup plus simple (voire simpliste) que celui de $\Delta^{\cap, \Sigma}$. Signalons enfin que le fait de satisfaire (MI) est beaucoup plus discutable que de satisfaire (Maj) pour un opérateur de fusion puisque cela ne permet pas de prendre en compte la popularité des informations.

	IC0	IC1	IC2	IC3	IC4	IC5	IC6	IC7	IC8	MI	Maj
Δ^{C1}	✓	✓	✓		✓	✓		✓		✓	
Δ^D	✓	✓	✓		✓	✓		✓			✓
$\Delta^{S,\Sigma}$	✓	✓	✓		✓			✓	✓		✓
$\Delta^{\cap,\Sigma}$	✓	✓	✓			✓	✓	✓	✓		✓

TAB. 9.2 – *Propriété des opérateurs de fusion avec fonctions de sélection*

9.3 Remarques

Dans ce chapitre nous avons étudié les propriétés logiques des opérateurs de combinaison proposés par Baral et al. . Nous avons montré que des propriétés importantes n'étaient pas vérifiées.

Nous avons alors montré que l'utilisation de fonctions de sélection permet d'augmenter les propriétés logiques des opérateurs de combinaison et d'améliorer leur comportement. Nous avons seulement utilisé ici une méthode d'agrégation utilitariste en ajoutant les différentes distances lorsque l'on calcule la distance entre un maximal consistant et un ensemble de connaissance. Il pourrait être intéressant de voir si on obtient des résultats similaires avec une méthode égalitariste à la $\Delta^{d,GMax}$.

Une autre question intéressante est de trouver les propriétés générales qu'une fonction de sélection doit satisfaire pour obtenir des opérateurs de fusion vérifiant tous les postulats, comme pour les fonctions de contraction par intersection partielle relationnelles transitives pour la révision.

Conclusion

La motivation initiale de ce travail était d'utiliser et de définir des outils issus de la théorie de la révision (opérateurs de révision, de mise à jour, d'abduction et de fusion) pour modéliser les systèmes coopératifs [Kon96b, Kon96a]. La plupart de ces systèmes peuvent être représentés par la figure 9.1.

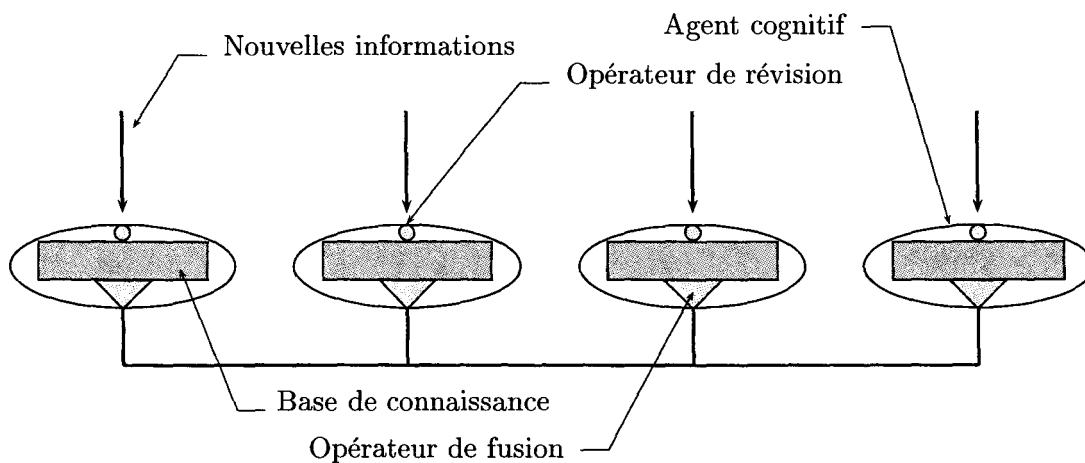


FIG. 9.1 – *Système coopératif*

C'est-à-dire que l'on considère un ensemble d'agents cognitifs qui sont principalement caractérisés par leur base de connaissance. Ces agents sont capables d'intégrer de nouvelles informations à leur base de connaissance grâce à un opérateur de révision. Et il est également possible d'effectuer une concertation entre agents à l'aide d'un opérateur de fusion. Cela permet alors de définir un état de connaissance du système dans son ensemble.

Beaucoup de systèmes coopératifs peuvent se réduire à un tel modèle. Dans [Kon96a] nous avons proposé une taxonomie des différents systèmes suivant différents critères, selon que l'opérateur de révision est le même pour tous les agents ou est différent, selon que les nouvelles informations sont communes à tous les agents ou non, etc.

L'étape suivante était d'étudier formellement les propriétés logiques des différents systèmes.

Mais il y avait de nombreux problèmes pour mener à bien cette étude.

Tout d'abord, de tels systèmes doivent pouvoir manipuler un grand nombre d'informations. Il est donc primordial que l'itération du processus de révision soit conduite de manière rationnelle. Or il n'y avait pas de réponse définitive au problème de la révision itérée.

Ensuite, pour pouvoir effectuer une concertation entre agents, il fallait disposer d'opérateurs de fusion. La seule proposition alors, celle de Revesz, n'était pas suffisamment contraignante et ne pouvait donc pas être utilisée comme l'outil de concertation que l'on souhaitait.

Enfin, de tels systèmes doivent pouvoir manipuler des quantités importantes d'information et il était donc important d'avoir des opérateurs "calculables" ou tout au moins des opérateurs syntaxiques.

Ce sont donc ces trois directions que nous avons suivies au cours de cette thèse.

Nous avons tout d'abord étudié le problème de la révision itérée. Bien que les propositions dans ce domaine soient nombreuses, aucune ne semble apporter de réponse définitive au problème. Nous avons comparé ces différentes méthodes et nous avons montré qu'après un certain nombre d'itérations la plupart d'entre elles coïncidaient. Toutes ces propositions utilisent plus ou moins clairement une notion d'état épistémique, c'est-à-dire qu'elles travaillent sur des objets plus complexes qu'une simple base de connaissance. Nous avons montré que ces propositions utilisaient le même type d'Etats Epistémiques sous différentes sémantiques. Nous avons défini les opérateurs de révision à mémoire qui permettent un bon comportement vis à vis de l'itération en gardant trace des révisions antérieures et en donnant une primauté forte à la nouvelle information. Le fait qu'aucune méthode de révision itérée ne semble s'imposer provient peut-être du fait que les Etats Epistémiques considérés ne sont toujours pas suffisamment complexes. Nous avons donc proposé la définition d'opérateurs de révision à mémoire dynamique qui utilisent des Etats Epistémiques plus complexes. Nous avons également des opérateurs de révision où la nouvelle information n'est pas qu'une simple formule mais également un état épistémique. Cette généralisation permet de réviser les connaissances de l'agent pas seulement par des formules, mais également par des informations conditionnelles.

En ce qui concerne la révision syntaxique, nous avons défini cinq opérateurs de changement basés sur le chaînage avant. Le premier, l'opérateur de mise à jour factuelle est un opérateur de mise à jour au sens de Katsuno Mendelson. Les autres sont basés sur une hiérarchie entre formules, induite par la base de connaissance elle-même. Cette hiérarchie est inspirée de la clôture rationnelle de bases de conditionnels proposée par Kraus, Lehmann et Magidor. Nous avons étudié les propriétés logiques de ces différents opérateurs et nous avons montré que les opérateurs permettant de garder plus de transitivité étaient ceux qui vérifiaient le moins de propriétés logiques.

La deuxième partie de notre étude concernait le problème de la fusion de bases de connaissance. Nous avons proposé une caractérisation logique des opérateurs de fusion contrainte et de certaines sous-classes particulières. Nous avons en particulier fait une distinction entre opérateurs majoritaires et opérateurs d'arbitrage. Nous avons proposé des théorèmes de représentation en terme de familles de pré-ordres sur les interprétations pour les différents

opérateurs définis. Nous avons ensuite montré le comportement de différents opérateurs sur des exemples et nous avons, en particulier, montré qu'il était possible pour un opérateur d'être à la fois majoritaire et d'arbitrage. Nous avons également comparé notre approche à divers travaux, montrant tout d'abord que la révision était un cas particulier de fusion contrainte. Nous avons ensuite montré que notre approche permettait de capturer les propositions existantes pour la fusion. Les opérateurs de Revesz sont caractérisés par un sous-ensemble strict de nos postulats et les opérateurs de Liberatore et Schaerf peuvent être obtenus à partir d'opérateurs de fusion contrainte. Nous avons d'ailleurs proposé un nouveau théorème de représentation pour les opérateurs de Liberatore et Schaerf. Un avantage d'utiliser les opérateurs de fusion contrainte pour définir les opérateurs de Liberatore et Schaerf est qu'il est facile dans ce cas de généraliser la définition pour fusionner un nombre quelconque de bases de connaissance, alors que la définition de Liberatore et Schaerf ne permet de fusionner que deux bases.

La dernière partie de notre contribution concerne la fusion syntaxique. Dans cette partie, nous illustrons en quoi la caractérisation logique que nous avons proposée peut être utilisée pour classer différents opérateurs de fusion suivant leurs propriétés. Cette caractérisation permet également de suggérer des "améliorations" de la définition de certains opérateurs. Nous avons étudié les propriétés logiques des opérateurs de combinaison proposés par Baral *et al.* Nous avons alors proposé l'utilisation de fonctions de sélection, inspirées de celles utilisées pour la révision, pour améliorer le comportement de ces opérateurs. Nous avons montré que l'utilisation d'une fonction de sélection adéquate permettait de définir des opérateurs avec plus de propriétés logiques.

En ce qui concerne les perspectives de ce travail, considérons d'abord les perspectives propres à chacun de ses aspects. Pour la révision itérée, nous avons proposé une famille d'opérateurs particuliers, mais il manque toujours une caractérisation générale des opérateurs de révision itérée. Nous continuerons donc nos recherches à ce sujet. Nous pensons que les opérateurs de révision à mémoire dynamique et les opérateurs de révision par états épistémiques sont deux pistes intéressantes.

En ce qui concerne la révision syntaxique et le problème de l'implémentation de la révision en général, il serait intéressant d'étudier plus finement la complexité des opérateurs "chaînage avant" que nous avons définis et d'évaluer pratiquement leurs performances.

Pour la fusion contrainte, les perspectives semblent plus importantes que le travail réalisé. En effet nous sommes convaincus qu'il est possible d'exporter bon nombre des résultats obtenus dans le cadre de la révision AGM à la fusion de bases de connaissances. Si on examine la figure 1.1, en tentant de faire le parallèle avec la fusion contrainte, on voit que nous avons donné le pendant des assignements fidèles avec les assignements syncrétiques. Nous avons également commencé l'investigation de l'utilisation de fonctions de sélection, ceci illustrant comment l'utilisation de résultats pour la révision (dans ce cas, révision par intersection partielle) pouvait s'avérer utile pour la fusion. Mais nous avons dans ce cadre simplement étudié quelques fonctions de sélection particulières et nous souhaitons montrer, comme pour la révision, que les opérateurs de fusion contrainte correspondent à une famille de fonctions de sélection particulière. Enfin, nous n'avons pas étudié les autres correspondances possibles. En particulier, il serait intéressant de tenter de caractériser la fusion contrainte par des enracinements épistémiques. Le but serait alors de définir un enracinement épistémique global à

partir d'enracinements épistémiques individuels. Cette définition permettrait d'identifier plus clairement les liens entre fusion contrainte, théorie du choix social et théorie de la décision.

Une dernière perspective de première importance pour la fusion est l'introduction de pondérations dans le cadre de la fusion contrainte. En effet, les individus d'un groupe ont rarement tous la même importance (ou la même fiabilité), et ceci doit pouvoir entrer en compte dans la détermination de la connaissance globale du groupe. Revesz a proposé des opérateurs qui permettent une pondération, mais la caractérisation logique ne contraint pas l'utilisation de cette pondération. La définition d'ensemble de connaissance comme des multi-ensembles permet un semblant de pondération par la répétition des bases, ce qui peut être manié de manière correcte par les opérateurs majoritaires. Ceci est plus discutable pour les opérateurs d'arbitrage. Il est donc nécessaire d'introduire des pondérations associées à chaque base de connaissance et de contraindre leur utilisation de manière satisfaisante dans la caractérisation logique. Nous souhaitons que cette caractérisation soit suffisamment souple pour pouvoir inclure à la fois des systèmes hiérarchiques, où une base de poids n est toujours préférée aux bases de poids $n - 1$, et des systèmes plus souples, que nous nommons démocratiques, où un certain nombre de bases de poids $n - 1$ seront préférées à une base de poids n . Dans cette optique, il serait alors intéressant de considérer des logiques pondérées, comme la logique des pénalités [dSCLS94, dSC96] par exemple.

En ce qui concerne les perspectives "croisées", l'ensemble des outils que nous avons définis peut être utilisé pour l'étude formelle de systèmes coopératifs à bases de connaissance. On peut alors envisager d'étudier les propriétés logiques (la rationalité) et le comportement de systèmes particuliers.

Table des figures

1.1	Théorèmes de représentation	30
1.2	Exemple de pré-ordre associé à φ	35
1.3	Révision de φ par μ	36
1.4	Système de sphères centré sur $[K]$	36
1.5	Système de sphères centré sur $[K]$: Révision par A	37
1.6	Représentation d'un modèle rangé	40
1.7	Information conditionnelle	50
1.8	Opérateur de révision AGM aléatoire	51
2.1	Révision Darwiche et Pearl	61
2.2	Révision naturelle	63
2.3	Modèle rangé	66
2.4	modèle rangé : nouvelle information compatible	66
2.5	modèle rangé : nouvelle information non compatible	67
2.6	révision rangée	67
2.7	Fonction Ordinale Conditionnelle: $(\varphi,1)$ -conditionnalisation	69
2.8	Fonction Ordinale Conditionnelle: $(\varphi,3)$ -ajustement	70
2.9	révision par φ	74
3.1	Exemples opérateur basique à mémoire	90
3.2	Opérateur basique à mémoire : révision par φ	93
3.3	Exemples opérateur Dalal à mémoire	93
3.4	clôture du pré-ordre	96
3.5	Distinction entre les opérateurs de révision à mémoire	97
4.1	Relations entre les opérateurs de révision chaînage avant	121
4.2	Propriétés des opérateurs syntaxiques	122
6.1	Les différentes expressions de la propriété d' équit��	140
6.2	Arbitrage	144
7.1	Fusion de deux bases de connaissance	163
7.2	La famille Δ^{d,Σ^n}	164
8.1	R��gle de choix social	181
8.2	Correspondance fusion contrainte / th��orie du choix social	184
8.3	Pr��-ordres	185
9.1	Syst��me coop��ratif	199

Index des propriétés logiques

(A1) $\Delta(\Psi)$ est consistant ,	171
(A2) Si Ψ est consistant, alors $\Delta(\Psi) = \bigwedge \Psi$,	171
(A3) Si $\Psi_1 \leftrightarrow \Psi_2$, alors $\vdash \Delta(\Psi_1) \leftrightarrow \Delta(\Psi_2)$,	171
(A4) Si $\varphi \wedge \varphi'$ n'est pas consistant, alors $\Delta(\varphi \sqcup \varphi') \not\vdash \varphi$,	171
(A5) $\Delta(\Psi_1) \wedge \Delta(\Psi_2) \vdash \Delta(\Psi_1 \sqcup \Psi_2)$,	171
(A6) Si $\Delta(\Psi_1) \wedge \Delta(\Psi_2)$ est consistant, alors $\Delta(\Psi_1 \sqcup \Psi_2) \vdash \Delta(\Psi_1) \wedge \Delta(\Psi_2)$,	171
(A7) $\forall \varphi' \exists \varphi \varphi' \not\vdash \varphi \forall n \Delta(\varphi' \sqcup \varphi^n) \leftrightarrow \Delta(\varphi' \sqcup \varphi)$,	171
(Ar3) $\left. \begin{array}{l} \forall \mu_i \exists \mu_i \forall \mu \Delta_{\mu_i \vee \mu}(\varphi_i) \vdash \mu \\ \forall i \mu' \wedge \mu_i \vdash \perp \end{array} \right\} \Rightarrow \Delta_{\mu' \vee \bigvee \mu_i}(\sqcup \varphi_i) \vdash \mu'$,	141
(Ar4) $\left. \begin{array}{l} \forall \mu_i \exists \mu_i \forall \mu \Delta_{\mu_i \vee \mu}(\varphi_i) \vdash \mu \\ \forall i \mu' \wedge \mu_i \vdash \perp \end{array} \right\} \Rightarrow \Delta_{\top}(\sqcup \varphi_i) \wedge \bigvee \mu_i \vdash \perp$,	141
(Arb) $\left. \begin{array}{l} \Delta_{\mu_1}(\varphi_1) \leftrightarrow \Delta_{\mu_2}(\varphi_2) \\ \Delta_{\mu_1 \leftrightarrow \neg \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow (\mu_1 \leftrightarrow \neg \mu_2) \\ \mu_1 \not\vdash \mu_2 \\ \mu_2 \not\vdash \mu_1 \end{array} \right\} \Rightarrow \Delta_{\mu_1 \vee \mu_2}(\varphi_1 \sqcup \varphi_2) \leftrightarrow \Delta_{\mu_1}(\varphi_1)$ (arbitrage),	135
(Ass) $\Delta_{\mu}(\Psi_1 \sqcup \Delta_{\mu}(\Psi_2)) = \Delta_{\mu}(\Psi_1 \sqcup \Psi_2)$ (associativité),	137
(C) Si $\varphi \wedge \mu$ est consistant, alors $\Phi \bullet \varphi \bullet \mu \leftrightarrow \Phi \bullet (\varphi \wedge \mu)$,	86
(C*) Si $\Theta \wedge \Gamma$ est satisfaisable, alors pour tout état épistémique Θ' tel que $\Theta' \leftrightarrow \Theta \wedge \Gamma$ on a $(\Phi \circ \Theta) \circ \Gamma \leftrightarrow \Phi \circ \Theta'$,	102
(C1*) Si $\Theta \vdash \Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \leftrightarrow \Phi \circ \Theta$,	103
(C1) Si $\alpha \vdash \mu$, alors $(\Phi \circ \mu) \circ \alpha \leftrightarrow \Phi \circ \alpha$,	60
(C2*) Si $\Theta \vdash \neg \Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \leftrightarrow \Phi \circ \Theta$,	103
(C2) Si $\alpha \vdash \neg \mu$, alors $(\Phi \circ \mu) \circ \alpha \leftrightarrow \Phi \circ \alpha$,	60
(C3*) Si $\Phi \circ \Theta \vdash \Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \vdash \Gamma$,	103
(C3) Si $\Phi \circ \alpha \vdash \mu$, alors $(\Phi \circ \mu) \circ \alpha \vdash \mu$,	60
(C4*) Si $\Phi \circ \Theta \not\vdash \neg \Gamma$, alors $(\Phi \circ \Gamma) \circ \Theta \not\vdash \neg \Gamma$,	103
(C4) Si $\Phi \circ \alpha \not\vdash \neg \mu$, alors $(\Phi \circ \mu) \circ \alpha \not\vdash \neg \mu$,	60
(CB) Si $\Phi \circ \psi \vdash \neg \mu$, alors $(\Phi \circ \psi) \circ \mu \leftrightarrow \Phi \circ \mu$,	63
(D1) $A \geq_c A$ (réflexivité),	41
(D2) $A \geq_c B$ ou $B \geq_c A$ (totalité),	41
(D3) Si $A \geq_c B$ et $B \geq_c C$, alors $A \geq_c C$ (transitivité),	41
(D4) $\top >_c \perp$ (non trivialité),	41
(D5) $\top \geq_c A$ (certitude de la tautologie),	41

(D6) Si $A \geq_c B$, alors $A \wedge C \geq_c B \wedge C$ (stabilité conjonctive),	41
(EE1) Si $A \leq B$ et $B \leq C$, alors $A \leq C$ (transitivité),	32
(EE2) Si $A \vdash B$, alors $A \leq B$ (domination),	32
(EE3) $A \leq A \wedge B$ ou $B \leq A \wedge B$ (conjonction),	33
(EE4) Si $K \neq K_\perp$, $A \notin K$ ssi $\forall B A \leq B$ (minimalité),	33
(EE5) Si $B \leq A \forall B$, alors $\vdash A$ (maximalité),	33
(Eq1) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \not\vdash \varphi \Rightarrow \Delta_\mu(\varphi' \sqcup \varphi) \neq \varphi'$,	139
(Eq2) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \vdash \varphi \Leftrightarrow \Delta_\mu(\varphi' \sqcup \varphi) = \varphi'$,	139
(Eq3) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \wedge \varphi \vdash \perp \Rightarrow \Delta_\mu(\varphi' \sqcup \varphi) \neq \varphi'$,	139
(Eq4) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \wedge \varphi \vdash \perp \Rightarrow \Delta_\mu(\varphi' \sqcup \varphi) \not\vdash \varphi'$,	139
(Eq5) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\varphi' \wedge \varphi \not\vdash \perp \Leftrightarrow \Delta_\mu(\varphi' \sqcup \varphi) \vdash \varphi'$,	139
(H'7) Si $\varphi \bullet \mu \leftrightarrow \mu$, alors $\Phi \bullet \varphi \bullet \mu \leftrightarrow \Phi \bullet \mu$,	86
(H'8) Si $\Phi' \vdash \mu$, alors $\Phi \bullet \Phi' \vdash \mu$,	86
(H*0) $K * A$ est consistant si A est consistant (consistance),	48
(H*1) $K * A \subseteq K \cup A$ (inclusion),	49
(H*2) Si $x \in K \setminus (K * A)$, alors $\exists K'$ tel que $K * A \subseteq K' \subseteq K \cup A$, K' est consistant et $K' \cup \{x\}$ est consistant (pertinence),	49
(H*3) $A \subseteq K * A$ (succès),	49
(H*4) Si, pour tout $K' \subseteq K$, $K' \cup A$ est inconsistant ssi $K' \cup B$ est inconsistant, alors $K \cap (K * A) = K \cap (K * B)$ (uniformité),	49
(H*4w) Si A et B sont inclus dans K et, pour tout $K' \subseteq K$, $K' \cup A$ est inconsistant ssi $K' \cup B$ est inconsistant, alors $K \cap (K * A) = K \cap (K * B)$ (uniformité faible),	49
(H*5) Si A est consistant, et $A \cup \{b\}$ est inconsistant pour chaque $b \in B$, alors $K * A = (K \cup B) * A$ (redondance),	49
(H*6) $K * A = (K \cup A) * A$ (pré-expansion),	49
(H0*) $Mod(\Phi \circ \Theta) = \min(Mod(\Theta), \leq_\Phi)$,	103
(H1*) $\Phi \circ \Theta \vdash \Theta$,	101
(H1) $\Phi \bullet \varphi \vdash \varphi$,	82
(H2*) Si $\Phi \wedge \Theta$ est consistant, alors $\Phi \circ \Theta \leftrightarrow \Phi \wedge \Theta$,	101
(H2) Si $\Phi \wedge \varphi$ est consistant, alors $\Phi \bullet \varphi \leftrightarrow \Phi \wedge \varphi$,	83
(H3*) Si Θ est consistant, alors $\Phi \circ \Theta$ est consistant,	101
(H3) Si φ est consistant, alors $\Phi \bullet \varphi$ est consistant,	83
(H4*) Si $\Phi_1 = \Phi_2$ et $\Theta_1 \leftrightarrow \Theta_2$, alors $\Phi_1 \circ \Theta_1 \leftrightarrow \Phi_2 \circ \Theta_2$,	101
(H4) Si $\Phi_1 = \Phi_2$ et $\varphi_1 \leftrightarrow \varphi_2$, alors $\Phi_1 \bullet \varphi_1 = \Phi_2 \bullet \varphi_2$,	83
(H5*) $(\Phi \circ \Theta) \wedge \Gamma \vdash \Phi \circ (\Theta \wedge \Gamma)$,	102
(H5) $(\Phi \bullet \varphi) \wedge \mu \vdash \Phi \bullet (\varphi \wedge \mu)$,	83
(H6*) Si $(\Phi \circ \Theta) \wedge \Gamma$ est consistant, alors $\Phi \circ (\Theta \wedge \alpha) \vdash \Phi \circ (\Theta \wedge \Gamma)$,	102
(H6) Si $(\Phi \bullet \varphi) \wedge \mu$ est consistant, alors $\Phi \bullet (\varphi \wedge \mu) \vdash (\Phi \bullet \varphi) \wedge \mu$,	83
(H7) $\Phi \bullet \Phi' \leftrightarrow \Phi \bullet \pi(\Phi')$,	83
(HI*) $(\Phi \circ \Theta) \circ \Gamma = \Phi \circ (\Theta \circ \Gamma)$,	101
(H÷1) $K \div A \subseteq K$ (inclusion),	48

(H÷2) Si $x \in K \setminus (K \div A)$, alors il existe un K' avec $K \div A \subseteq K' \subseteq K$, tel que $A \cap Cn(K') = \emptyset$ et $A \cap Cn(K' \cup \{x\}) \neq \emptyset$ (pertinence),	48
(H÷3) Si $A \cup Cn(\emptyset) = \emptyset$, alors $A \cup Cn(K \div A) = \emptyset$ (succès),	48
(H÷4) Si pour chaque sous-ensemble K' de K on a $A \cup Cn(K') = \emptyset$ si et seulement si $B \cup Cn(K') = \emptyset$, alors $K \div A = K \div B$ (uniformité),	48
(H÷5) Si $A \cup Cn(\emptyset) = \emptyset$, et si chaque élément de B implique un élément de A , alors $K \div A = (K \cap B) \div A$ (redondance),	48
(I1) $Bel(\Phi)$ est une théorie consistante ,	64
(I2) $\Phi \circ \varphi \vdash \varphi$,	64
(I3) Si $\Phi \circ \varphi \vdash \mu$, alors $\Phi \vdash \varphi \rightarrow \mu$,	64
(I4) Si $\Phi \vdash \mu$, alors $\Phi \circ \Phi' \leftrightarrow \Phi \circ \mu \circ \Phi'$,	64
(I5) Si $\varphi \vdash \mu$, alors $\Phi \circ \mu \circ \varphi \circ \Phi' \leftrightarrow \Phi \circ \varphi \circ \Phi'$,	64
(I6) Si $\Phi \circ \varphi \not\vdash \neg \mu$, alors $\Phi \circ \varphi \circ \mu \circ \Phi' \leftrightarrow \Phi \circ \varphi \circ (\varphi \wedge \mu) \circ \Phi'$,	64
(I7) $\Phi \circ \neg \varphi \circ \varphi \subseteq Cn(Bel(\Phi) \cup \{\varphi\})$,	64
(IC _{it}) Si $\varphi \vdash \mu$ alors $\exists n \Delta_\mu^n(\Psi, \varphi) \vdash \varphi$ (itération),	135
(IC0) $\Delta_\mu(\Psi) \vdash \mu$,	134
(IC1) Si μ est consistant, alors $\Delta_\mu(\Psi)$ est consistant ,	134
(IC2) Si Ψ est consistant avec μ , alors $\Delta_\mu(\Psi) = \bigwedge \Psi \wedge \mu$,	134
(IC3) Si $\Psi_1 \leftrightarrow \Psi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\Delta_{\mu_1}(\Psi_1) \leftrightarrow \Delta_{\mu_2}(\Psi_2)$,	134
(IC4) Si $\varphi \vdash \mu$ et $\varphi' \vdash \mu$, alors $\Delta_\mu(\varphi \sqcup \varphi') \wedge \varphi \not\vdash \perp \Rightarrow \Delta_\mu(\varphi \sqcup \varphi') \wedge \varphi' \not\vdash \perp$,	134
(IC5) $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2) \vdash \Delta_\mu(\Psi_1 \sqcup \Psi_2)$,	134
(IC6') Si $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant, alors $\Delta_\mu(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\mu(\Psi_1) \vee \Delta_\mu(\Psi_2)$, ...	135
(IC6) Si $\Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$ est consistant, alors $\Delta_\mu(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\mu(\Psi_1) \wedge \Delta_\mu(\Psi_2)$, ...	134
(IC7) $\Delta_{\mu_1}(\Psi) \wedge \mu_2 \vdash \Delta_{\mu_1 \wedge \mu_2}(\Psi)$,	134
(IC8) Si $\Delta_{\mu_1}(\Psi) \wedge \mu_2$ est consistant, alors $\Delta_{\mu_1 \wedge \mu_2}(\Psi) \vdash \Delta_{\mu_1}(\Psi) \wedge \mu_2$,	134
(K□1) $K \Box A$ est une théorie (clôture),	104
(K□2) Si $\neg A \notin K$, alors $A \in K \Box A$ (succès faible),	104
(K□3) $K \Box A \subseteq K + A$ (inclusion),	104
(K□4) Si $\neg A \notin K$, alors $K + A \subseteq K \Box A$ (vacuité),	104
(K□5) Si K est consistant, alors $K \Box A$ est consistant (maintien de la consistance),	104
(K□6) Si $A \leftrightarrow B$, alors $K \Box A = K \Box B$ (extensionnalité),	104
(K*0) $K_\perp * A = Cn(\{A\})$ (absurdité),	72
(K*1) $K * A$ est une théorie (clôture),	26
(K*2) $A \in K * A$ (succès),	26
(K*3) $K * A \subseteq K + A$ (inclusion),	26
(K*4) Si $\neg A \notin K$, alors $K + A \subseteq K * A$ (vacuité),	26
(K*5) $K * A = K_\perp$ ssi $\vdash \neg A$ (consistance),	26
(K*6) Si $A \leftrightarrow B$, alors $K * A = K * B$ (extensionnalité),	26
(K*7) $K * (A \wedge B) \subseteq (K * A) + B$ (inclusion conjonctive),	26
(K*7new) Si $A \wedge B \not\vdash \perp$, alors $(K * A) * B = K * (A \wedge B)$ (conjonction),	72
(K*8) Si $\neg B \notin K * A$, alors $(K * A) + B \subseteq K * (A \wedge B)$ (vacuité conjonctive),	26
(K*8new) Si $A \wedge B \vdash \perp$, alors $(K * A) * B = K * B$ (DP2'),	72

(K÷1) $K \div A$ est une théorie (clôture),	27
(K÷2) $K \div A \subseteq K$ (inclusion),	27
(K÷3) Si $A \notin K$, alors $K \div A = K$ (vacuité),	27
(K÷4) Si $\nVdash A$, alors $A \notin K \div A$ (succès),	27
(K÷5) Si $A \in K$, alors $K \subseteq (K \div A) + A$ (restauration),	28
(K÷6) Si $A \leftrightarrow B$, alors $K \div A = K \div B$ (préservation),	28
(K÷7) $(K \div A) \cap (K \div B) \subseteq K \div (A \wedge B)$ (intersection),	28
(K÷8) Si $A \notin K \div (A \wedge B)$, alors $K \div (A \wedge B) \subseteq K \div A$ (conjonction),	28
(K+1) $K + A$ est une théorie (clôture),	25
(K+2) $A \in K + A$ (succès),	25
(K+3) $K \subseteq K + A$ (inclusion),	25
(K+4) Si $A \in K$, alors $K + A = K$ (vacuité),	25
(K+5) Si $K \subseteq H$, alors $K + A \subseteq H + A$ (monotonie),	25
(K+6) $K + A$ est la plus petite base de connaissance satisfaisant (K+1)-(K+5) (minimalité),	25
(LM1) $\blacktriangle(\Psi)$ est consistant ,	131
(LM2) Si $\bigwedge \Psi$ est consistant, alors $\blacktriangle(\Psi) = \bigwedge \Psi$,	131
(LM3) Si $\Psi \leftrightarrow \Psi'$, alors $\blacktriangle(\Psi) \leftrightarrow \blacktriangle(\Psi')$,	131
(LM4) Soit un littéral l , si $ \{\varphi_i \in \Psi : \varphi_i \models l\} > \{\varphi_i \in \Psi : \varphi_i \models \neg l\} + \{\varphi_i \in \Psi : \varphi_i \triangleright \neg l\} $, alors $\blacktriangle(\Psi)$ implique l .	131
(LS1) $\varphi \diamond \mu \leftrightarrow \mu \diamond \varphi$,	129
(LS2) $\varphi \wedge \mu$ implique $\varphi \diamond \mu$,	129
(LS3) Si $\varphi \wedge \mu$ est consistant, alors $\varphi \diamond \mu$ implique $\varphi \wedge \mu$,	129
(LS4) $\varphi \diamond \mu$ est inconsistant ssi φ et μ sont inconsistants, ,	129
(LS5) Si $\varphi_1 \leftrightarrow \varphi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\varphi_1 \diamond \mu_1 \leftrightarrow \varphi_2 \diamond \mu_2$,	129
(LS6) $\varphi \diamond (\mu \vee \theta) = \begin{cases} \varphi \diamond \mu & \text{ou} \\ \varphi \diamond \theta & \text{ou} \\ (\varphi \diamond \mu) \vee (\varphi \diamond \theta) \end{cases}$,	129
(LS7) $\varphi \diamond \mu$ implique $\varphi \vee \mu$,	129
(LS78) Si φ et μ sont deux formules complètes, alors $\varphi \diamond \mu = \varphi \vee \mu$, ,	130
(LS8) If φ est consistant alors $\varphi \wedge (\varphi \diamond \mu)$ est consistant ,	129
(M1) $\Psi \triangleright \mu$ implique μ . ,	127
(M2) Si $\Psi \wedge \mu$ est consistant, alors $\Psi \triangleright \mu \leftrightarrow \Psi \wedge \mu$, ,	127
(M3) Si μ est consistant, alors $\Psi \triangleright \mu$ est consistant, ,	127
(M4) Si $\Psi_1 \Leftrightarrow \Psi_2$ et $\mu \leftrightarrow \phi$, alors $\Psi_1 \triangleright \mu \leftrightarrow \Psi_2 \triangleright \phi$, ,	127
(M5) $(\Psi \triangleright \mu) \wedge \phi$ implique $\Psi \triangleright (\mu \wedge \phi)$, ,	127
(M6) Si $(\Psi \triangleright \mu) \wedge \phi$ est consistant alors $\Psi \triangleright (\mu \wedge \phi)$ implique $(\Psi \triangleright \mu) \wedge \phi$, ,	127
(M7) $\exists n \Delta(\Psi \sqcup \varphi^n) \vdash \varphi$, ,	171
(M7) $(\Psi_1 \triangleright \mu) \wedge (\Psi_2 \triangleright \mu)$ implique $(\Psi_1 \sqcup \Psi_2) \triangleright \mu$, ,	127
(MI) $\forall n \Delta_\mu(\Psi_1 \sqcup \Psi_2^n) \leftrightarrow \Delta_\mu(\Psi_1 \sqcup \Psi_2)$ (indépendance de la majorité),	136
(Ma2) $\exists n \Delta_\mu(\Psi \sqcup \varphi^n) \vdash \Delta_\mu(\varphi)$, ,	140
(Maj) $\exists n \Delta_\mu(\Psi_1 \sqcup \Psi_2^n) \vdash \Delta_\mu(\Psi_2)$ (majorité),	135

(Mon) Si $\varphi_1 \vdash \varphi'_1, \dots, \varphi_n \vdash \varphi'_n$ alors $\Delta_\mu(\varphi_1 \sqcup \dots \sqcup \varphi_n) \vdash \Delta_\mu(\varphi'_1 \sqcup \dots \sqcup \varphi'_n)$ (monotonie),	137
(OEE1) Si $A \vdash B$, alors $E(A) \leq E(B)$,	71
(OEE2) $E(A) \leq E(A \wedge B)$ ou $E(B) \leq E(A \wedge B)$,	71
(OEE3) $\vdash A$ si et seulement si $E(A) = \mathcal{O}$,	71
(OEE4) Si A est inconsistant, alors $E(A) = 0$,	71
(R*1) $\Phi \circ \mu \vdash \mu$,	59
(R*1) $\Phi \bullet \mu \vdash \mu$,	82
(R*2) Si $\Phi \wedge \mu$ est consistant, alors $\Phi \circ \mu \leftrightarrow \Phi \wedge \mu$,	59
(R*2) Si $\Phi \wedge \mu$ est consistant, alors $\Phi \bullet \mu \leftrightarrow \Phi \wedge \mu$,	82
(R*3) Si μ est consistant, alors $\Phi \circ \mu$ est consistant,	59
(R*3) Si μ est consistant, alors $\Phi \bullet \mu$ est consistant,	82
(R*4) Si $\Phi_1 = \Phi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\Phi_1 \circ \mu_1 \leftrightarrow \Phi_2 \circ \mu_2$,	59
(R*4) Si $\Phi_1 = \Phi_2$ et $\mu_1 \leftrightarrow \mu_2$, alors $\Phi_1 \bullet \mu_1 \leftrightarrow \Phi_2 \bullet \mu_2$,	82
(R*5) $(\Phi \circ \mu) \wedge \varphi \vdash \Phi \circ (\mu \wedge \varphi)$,	59
(R*5) $(\Phi \bullet \mu) \wedge \varphi \vdash \Phi \bullet (\mu \wedge \varphi)$,	82
(R*6) Si $(\Phi \circ \mu) \wedge \varphi$ est consistant, alors $\Phi \circ (\mu \wedge \varphi) \vdash (\Phi \circ \mu) \wedge \varphi$,	59
(R*6) Si $(\Phi \bullet \mu) \wedge \varphi$ est consistant, alors $\Phi \bullet (\mu \wedge \varphi) \vdash (\Phi \bullet \mu) \wedge \varphi$,	82
(R1) $\varphi \circ \mu \vdash \mu$,	27
(R2) Si $\varphi \wedge \mu$ est consistant alors $\varphi \circ \mu \leftrightarrow \varphi \wedge \mu$,	27
(R3) Si μ est consistant alors $\varphi \circ \mu$ est consistant,	27
(R4) Si $\varphi_1 \leftrightarrow \varphi_2$ et $\mu_1 \leftrightarrow \mu_2$ alors $\varphi_1 \circ \mu_1 \leftrightarrow \varphi_2 \circ \mu_2$,	27
(R5) $(\varphi \circ \mu) \wedge \phi \vdash \varphi \circ (\mu \wedge \phi)$,	27
(R6) Si $(\varphi \circ \mu) \wedge \phi$ est consistant alors $\varphi \circ (\mu \wedge \phi) \vdash (\varphi \circ \mu) \wedge \phi$,	27
(S1) Si $S, S' \in \mathbf{S}$, alors $S \subseteq S'$ ou $S' \subseteq S$,	37
(S2) $[K] \in \mathbf{S}$,	37
(S3) Si $S \in \mathbf{S}$, alors $[K] \subseteq S$,	37
(S4) $M_{\mathcal{L}} \in \mathbf{S}$,	37
(S5) Si A est une formule et si $[A]$ intersecte une sphère de $\mathbf{S}$, alors il existe une sphère minimale qui intersecte $[A]$,	37
(Sym) $f_{\varphi'}^\circ(\varphi') = f_{\varphi'}^\circ(\varphi)$ (symétrie),	169
(U1) $\varphi \diamond \mu \vdash \mu$,	55
(U2) Si $\varphi \vdash \mu$, alors $\varphi \diamond \mu \leftrightarrow \varphi$,	55
(U3) Si φ et μ sont consistants alors $\varphi \diamond \mu$ est consistant,	55
(U4) Si $\varphi_1 \leftrightarrow \varphi_2$ et $\mu_1 \leftrightarrow \mu_2$ alors $\varphi_1 \diamond \mu_1 \leftrightarrow \varphi_2 \diamond \mu_2$,	55
(U5) $(\varphi \diamond \mu) \wedge \phi \vdash \varphi \diamond (\mu \wedge \phi)$,	55
(U6) Si $\varphi \diamond \mu_1 \vdash \mu_2$ et $\varphi \diamond \mu_2 \vdash \mu_1$, alors $\varphi \diamond \mu_1 \leftrightarrow \varphi \diamond \mu_2$,	55
(U7) Si φ est une formule complète, alors $(\varphi \diamond \mu_1) \wedge (\varphi \diamond \mu_2) \vdash \varphi \diamond (\mu_1 \vee \mu_2)$,	55
(U8) $(\varphi_1 \vee \varphi_2) \diamond \mu \leftrightarrow (\varphi_1 \diamond \mu) \vee (\varphi_2 \diamond \mu)$,	55
(W1) $\tilde{\Psi} \triangleright \mu$ implique μ ,	128
(W2) Si $form(\tilde{\Psi}) \wedge \mu$ est consistant, alors $\tilde{\Psi} \triangleright \mu \leftrightarrow form(\tilde{\Psi}) \wedge \mu$,	128
(W3) Si μ est consistant, alors $\tilde{\Psi} \triangleright \mu$ est consistant,	128

(W4)	Si $\mu \leftrightarrow \phi$, alors $\tilde{\Psi} \triangleright \mu \leftrightarrow \tilde{\Psi} \triangleright \phi$,	128
(W5)	$(\tilde{\Psi} \triangleright \mu) \wedge \phi$ implique $\tilde{\Psi} \triangleright (\mu \wedge \phi)$,	128
(W6)	Si $(\tilde{\Psi} \triangleright \mu) \wedge \phi$ est consistant alors $\tilde{\Psi} \triangleright (\mu \wedge \phi)$ implique $(\tilde{\Psi} \triangleright \mu) \wedge \phi$,	128
(W7)	$(\tilde{\Psi}_1 \triangleright \mu) \wedge (\tilde{\Psi}_2 \triangleright \mu)$ implique $(\tilde{\Psi}_1 \uplus \tilde{\Psi}_2) \triangleright \mu$,	128
(W8)	Si $(\tilde{\Psi}_1 \triangleright \mu) \wedge (\tilde{\Psi}_2 \triangleright \mu)$ est consistant, alors $(\tilde{\Psi}_1 \uplus \tilde{\Psi}_2) \triangleright \mu$ implique $(\tilde{\Psi}_1 \triangleright \mu) \wedge (\tilde{\Psi}_2 \triangleright \mu)$,	128
(WA)	$\Delta_\mu(\Psi_1 \sqcup \Psi_2) \vdash \Delta_\mu(\Psi_1 \sqcup \Delta_\mu(\Psi_2))$ (associativité faible),	138
(WEE1)	Si $A \leq B$ et $B \leq C$, alors $A \leq C$ (transitivité),	71
(WEE2)	Si $A \vdash B$, alors $A \leq B$ (domination),	71
(WEE3)	$A \leq A \wedge B$ ou $B \leq A \wedge B$ (conjonction),	71
(WEE4)	Si $\exists C$ t.q. $\perp < C$, alors si $B \leq A \forall B$, alors $\vdash A$ (maximalité),	71
(WMI)	$\forall \varphi' \exists \varphi \varphi' \not\vdash \varphi \forall n \Delta_\mu(\varphi' \sqcup \varphi^n) = \Delta_\mu(\varphi' \sqcup \varphi)$ (indépendance faible de la majorité),	136
	$\alpha \vdash \beta \quad \alpha \vdash \gamma$	
AND	$\alpha \vdash \beta \wedge \gamma$,	38
	$\alpha \vdash \beta \quad \alpha \vdash \gamma$	
CM	$\alpha \wedge \gamma \vdash \beta$,	38
	$\alpha \wedge \beta \vdash \gamma \quad \alpha \vdash \beta$	
CUT	$\alpha \vdash \gamma$,	39
	$\alpha \vee \beta \vdash \gamma \quad \alpha \not\vdash \gamma$	
DR	$\beta \vdash \gamma$,	39
	$\alpha \vdash \beta \quad \vdash \alpha \leftrightarrow \gamma$	
LLE	$\gamma \vdash \beta$,	38
	$\alpha \vdash \beta \quad \alpha \wedge \gamma \not\vdash \beta$	
NR	$\alpha \wedge \neg \gamma \vdash \beta$,	39
	$\alpha \vdash \gamma \quad \beta \vdash \gamma$	
OR	$\alpha \vee \beta \vdash \gamma$,	38
REF	$\alpha \vdash \alpha$,	38
	$\alpha \vdash \beta \quad \alpha \not\vdash \neg \gamma$	
RM	$\alpha \wedge \gamma \vdash \beta$,	38
	$\alpha \vdash \beta \quad \vdash \beta \rightarrow \gamma$	
RW	$\alpha \vdash \gamma$,	38
	$\alpha \wedge \beta \vdash \gamma$	
S	$\alpha \vdash \beta \rightarrow \gamma$,	39
•	β, α ssi $\varphi \circ \alpha \vdash \beta$ (Test de Ramsey),	50
•	$K * A = (K \div \neg A) + A$ (Identité de Levi),	28
•	$K * A = (K + A) \div \neg A$ (Identité de Hansson),	48
•	$K \div A = K \cap (K * \neg A)$ (Identité de Harper),	28

Index

- P -consistance, 108
- $\Delta^{d,Max}$, 153
- $\Delta^{d,\Sigma}$, 155
- Δ^{d,Σ^n} , 157
- équivalence
 - ensemble de connaissance, 126
 - flock, 109, 188
- état épistémique
 - égalité, 59
 - équivalence, 59, 80
- absence de dictateur, 183
- acyclique, voir relation acyclique
- adéquation sémantique, 127
- agenda, 180
- ajustement, 70
- alternatives, 180
- anti-monotone en p , 94
- anti-symétrique, voir relation anti-symétrique
- approche
 - cohérente, 46
 - fondamentaliste, 46
- arbitrage, 127, 135
- assignement
 - commutatif, 130
 - conservatif, 83
 - fort, 84
 - fidèle, 35, 59
 - commutatif, 178
 - normalisé, 88
 - loyal, 127
 - loyal pondéré, 128
 - quasi-synchrétique, 143
 - synchrétique, 142
 - juste, 143
 - majoritaire, 143
 - pur, 173
 - pur juste, 173
 - pur majoritaire, 173
- avoir des explications transitives, 182
- base de P , 111
- base de connaissance, 20
 - close déductivement, 20
 - théorie, voir théorie
- belief base, 46
- Belief Base Dynamics, 47
- belief set, 46
- bien fondé, voir pré-ordre
- chainage avant
 - conséquences, 108
- choix maximal, 31
- combinaison, 188
- condition forte de Pareto, 182
- conditionnalisation, 68
- conditionnels, 49
- conséquence
 - opération de conséquence, 19
 - par chainage avant, 108
- consistance
 - ensemble de connaissance, 126
 - formule, 20
 - programme, 108
- consolidation, 53
- continue $\vdash$ vers le bas, 34
- continue $\vdash$ vers le haut, 34
- contraction, 27
 - à choix maximal, 31
 - effacement, 44
 - par intersection partielle, 31
 - par intersection totale, 31
 - protégeant ξ , 52
 - sûre, 34
- contrainte du domaine standard, 182
- contre-modèle, 20
- degré
 - d'incrédulité, 68

- de confiance, 68
- distance, 152
 - de Dalal, 153
 - topologique, 152
- effacement, 44, 55
- enracinement épistémique, 32
 - faible, 71
 - ordinal, 71
- ensemble conditionnel, 79
- ensemble de connaissance, 125
 - équivalence, 126
 - consistant, 126
 - pondéré, 128
- espace épistémique, 81
- expansion, 25
- flock, 109, 188
 - équivalence, 109, 188
 - implication, 109, 188
- fonction de choix, 181
- fonction de fusion, 172
 - juste, 172
 - majoritaire, 172
- fonction de rang, 111
- fonction de sélection, 31
 - relationnelle, 32
 - relationnelle transitive, 32
 - sensée, 120
 - totale, 32
- fonction ordinale conditionnelle, 68
 - ajustement, 70
 - conditionnalisation, 68
 - restructuration, 70
 - transmutation, 69
- formule
 - anti-monotone en p , 94
 - consistante, 20
 - monotone en p , 94
- fusion
 - basique, 152
 - combinaison, 188
 - par intersection totale, 151
 - pure, 171
 - arbitrage pur, 171
 - majoritaire, 171
- fusion contrainte, 134
 - arbitrage, 135
 - fusion majoritaire, 135
 - quasi-fusion, 135
- fusion majoritaire, 135
- fusion majoritaire (Lin-Mendelzon), 131
- Identité
 - de Hansson, 48
 - de Harper, 28
 - de Levi, 28
- indépendance des alternatives non disponibles, 182
- inférence naturelle, 94
- interprétation, 20
 - contre-modèle, 20
 - modèle, 20
- itération, 135
- leximin, 159
- littéral, 108
 - exceptionnel pour P , 111
- maxichoice, 31
- mesure
 - de nécessité, 42
 - de possibilité, 42
- minimax, 155
- mise à jour, 55
 - factuelle, 110
- modèle, 20
 - préférentiel, 39
 - rangé, 40
 - rangé emboité, 65, 66
- model fitting, 127
 - pondéré, 128
- modulaire, voir relation modulaire
- monde possible, 36
- monotone en p , 94
- opérateur
 - $\triangle^{d,Max}$, 153
 - $\triangle^{d,\Sigma}$, 155
 - $\triangle^{d,\Sigma^n}$, 157
 - fusion basique, 152
 - fusion par intersection totale, 151
 - mise à jour factuelle, 110
 - révision basique à mémoire, 90
 - révision Dalal à mémoire, 93

- révision de Ryan, 95
- révision gangue, 112
- révision gangue étendue, 112
- révision gangue sélective, 120
- révision hiérarchique, 111
- révision OTP_1 , 95
- révision OTP_2 , 96
- opération de conséquence, 19
- ordre, 19
 - strict, 19
- partielle, voir relation partielle
- Possible Models Approach, 55
- pré-ordre, 19
 - bien fondé, 19
- présentation ordonnée de théories (OTP), 94
- primauté de la nouvelle information, 52
- principe de primauté forte, 77
- profil, 181
- programme, 108
 - P -consistant, 108
 - base, 111
 - consistant, 108
- quasi-fusion, 135
- réflexive, voir relation réflexive
- révision, 26
 - états épistémiques, 59
 - à mémoire, 83
 - révision à mémoire par état épistémique, 102
 - révision basique à mémoire, 90
 - révision Dalal à mémoire, 93
 - révision dynamique à mémoire, 99
 - révision OTP_1 , 95
 - révision OTP_2 , 96
 - semi-révision à mémoire, 104
- commutative, 129
- définie à partir d'une distance, 170
- externe, 48
- filtrée, 52
 - relationnelle, 53
- gangue, 112
- gangue étendue, 112
- gangue sélective, 120
- hiérarchique, 111
- interne, 48
- naturelle, 63
 - par intersection partielle externe, 48
 - par intersection partielle interne, 48
 - sélective, 54
 - semi-révision, voir semi-révision
- règle, 108
 - exceptionnelle pour P , 111
- règle de choix social, 181
- relation
 - acyclique, 19
 - anti-symétrique, 19
 - continue $\vdash$ vers le bas, 34
 - continue $\vdash$ vers le haut, 34
 - d'équivalence, 19
 - de nécessité qualitative, 41
 - de préférence, 180
 - inférence naturelle, 94
 - modulaire, 19
 - ordre, 19
 - strict, 19
 - partielle, 19
 - pré-ordre, voir pré-ordre
 - réflexive, 19
 - symétrique, 19
 - totale, 19
 - transitive, 19
 - virtuellement totale, 34
- relevage, 175
- restructuration, 70, 71
- sélection, voir fonction de sélection
- schéma de transmutation, 69
- semi-révision, 53
 - à mémoire, 104
- smooth, 39
- supersélecteur, 47
 - unifié, 47
- symétrique, voir relation symétrique
- système de sphères, 37
- test de Ramsey, 50
- théorème d'impossibilité d'Arrow, 183
- théorie, 20
 - sous-théorie minimale, 33
- totale, voir relation totale
- transitive, voir relation transitive

transmutation, 69, 70
truth-maintenance systems, 46
two-place selection functions, 47

virtuellement totale, 34

Bibliographie

- [AGM85] C. E. Alchourrón, P. Gärdenfors, and D. Makinson. On the logic of theory change: Partial meet contraction and revision functions. *Journal of Symbolic Logic*, 50:510–530, 1985.
- [ALP⁺98] J. J. Alferes, J. A. Leite, L. M. Pereira, H. Przymusinska, and T. C. Przymusinski. Dynamic logic programming. In *Proceedings of the Sixth International Conference on Principles of Knowledge Representation and Reasoning (KR'98)*, pages 98–109. Morgan Kaufmann, 1998.
- [AM82] C. E. Alchourrón and D. Makinson. The logic of theory change: Contraction functions and their associated revision functions. *Theoria*, 48:14–37, 1982.
- [AM85] C. E. Alchourrón and D. Makinson. On the logic of theory change: Safe contraction. *Studia Logica*, 44:405–422, 1985.
- [AM86] C. E. Alchourrón and D. Makinson. Maps between some different kinds of contraction function: the finite case. *Studia Logica*, 45:187–198, 1986.
- [Arr63] K. J. Arrow. *Social choice and individual values*. Wiley, New York, second edition, 1963.
- [AV91] S. Abiteboul and V. Vianu. Datalog extensions for database queries and updates. *Journal of Computer and System Sciences*, 43:62–124, 1991.
- [BCD⁺93] S. Benferhat, C. Cayrol, D. Dubois, J. Lang, and H. Prade. Inconsistency management and prioritized syntax-based entailment. In *Proceedings of the Thirteenth International Joint Conference on Artificial Intelligence (IJCAI'93)*, pages 640–645, 1993.
- [BDL⁺98] S. Benferhat, D. Dubois, J. Lang, H. Prade, A. Saffioti, and P. Smets. A general approach for inconsistency handling and merging information in prioritized knowledge bases. In *Proceedings of the Sixth International Conference on Principles of Knowledge Representation and Reasoning (KR'98)*, pages 466–477, 1998.
- [BDP97] S. Benferhat, D. Dubois, and H. Prade. Some syntactic approaches to the handling of inconsistent knowledge bases: a comparative study, part 1: the flat case. *Studia Logica*, 58:17–45, 1997.
- [BDP99] S. Benferhat, D. Dubois, and O. Papini. A sequential reversible belief revision method based on polynomials. In *Proceedings of the Sixteenth National Conference on Artificial Intelligence (AAAI'99)*, 1999.
- [BFH98] C. Boutilier, N. Friedman, and J. Y. Halpern. Belief revision with unreliable informations. In *Proceedings of the Fifteenth National Conference on Artificial Intelligence (AAAI'98)*, 1998.
- [BGMS98] B. Bessant, E. Grégoire, P. Marquis, and L. Sais. Combining nonmonotonic reasoning and belief revision: a practical approach. In *Proceedings of the Interna-*

- tional Conference on Artificial Intelligence: Methodology, Systems, Applications (AIMSA'98)*, pages 115–128. Springer Verlag, 1998. Lecture Notes on Artificial Intelligence 1480.
- [BGMS99] B. Bessant, E. Grégoire, P. Marquis, and L. Sais. *Frontiers of belief revision*, chapter Iterated syntax-based revision in a nonmonotonic setting. Kluwer, 1999.
 - [BJKP97] H. Bezzazi, S. Janot, S. Konieczny, and R. Pino Pérez. Forward chaining and change operators. In *Proceedings of DYNAMICS'97, Workshop on (Trans)Actions and change in logic programming and deductive databases*, Port Jefferson, N.Y., 1997.
 - [BJKP98] H. Bezzazi, S. Janot, S. Konieczny, and R. Pino Pérez. Analysing rational properties of change operators based on forward chaining. In B. Freitag, H. Decker, M. Kifer, and A. Voronkov, editors, *Transactions and Change in Logic Databases*, Lecture Notes in Computer Science, pages 317–339. Springer Verlag, 1998.
 - [BKM91] C. Baral, S. Kraus, and J. Minker. Combining multiple knowledge bases. *IEEE Transactions on Knowledge and Data Engineering*, 3(2):208–220, 1991.
 - [BKMS92] C. Baral, S. Kraus, J. Minker, and V. S. Subrahmanian. Combining knowledge bases consisting of first-order theories. *Computational Intelligence*, 8(1):45–71, 1992.
 - [BKPP] S. Benferhat, S. Konieczny, O. Papini, and R. Pino Pérez. Révision itérée par état épistémique. submitted.
 - [BKPP99] S. Benferhat, S. Konieczny, O. Papini, and R. Pino Pérez. Révision itérée basée sur la primauté forte des observations. In *Journées Nationales sur les Modèles de raisonnement (JNMR'99)*, Paris, 1999. Electronic proceedings <http://www.irit.fr/GDRI3-ModRais/articlesJNMR.html>.
 - [BMP97] H. Bezzazi, D. Makinson, and R. Pino Pérez. Beyond rational monotony: Some strong non-horn rules for nonmonotonic inference relations. *Journal of Logic and Computation*, 7:605–631, 1997.
 - [Bor85] A. Borgida. Language features for flexible handling of exceptions in information systems. *ACM Transactions on Database Systems*, 10:563–603, 1985.
 - [Bou92a] C. Boutilier. *Conditional logics for default reasoning and belief revision*. PhD thesis, University of Toronto, 1992.
 - [Bou92b] C. Boutilier. Normative, subjunctive and autoepistemic defaults: Adopting the ramsey test. In *Proceedings of the Third International Conference on Principles of Knowledge Representation and Reasoning*, pages 685–696, 1992.
 - [Bou93] C. Boutilier. Revision sequences and nested conditionals. In *Proceedings of the Thirteenth International Joint Conference on Artificial Intelligence (IJCAI'93)*, 1993.
 - [Bou96] C. Boutilier. Iterated revision and minimal change of conditional beliefs. *Journal of Philosophical Logic*, 25(3):262–305, 1996.
 - [BP96] H. Bezzazi and R. Pino Pérez. Rational transitivity and its models. In *Proceedings of the Twenty Sixth International Symposium on Multiple-Valued Logics*, pages 160–165. IEEE Computer Society Press, 1996.
 - [Bre89] G. Brewka. Preferred subtheories: an extended logical framework for default reasoning. In *Proceedings of the Eleventh International Joint Conference on Artificial Intelligence (IJCAI'89)*, pages 1043–1048, 1989.

- [Cho93] L. Cholvy. Proving theorems in a multi-sources environment. In *Proceedings of the Thirteenth International Joint Conference on Artificial Intelligence (IJCAI'93)*, pages 66–71, 1993.
- [Cho95] L. Cholvy. Automated reasoning with merged contradictory information whose reliability depends on topics. In *Proceedings of the Third European Conference on Symbolic and Quantitative Approaches to Reasoning and Uncertainty (ECSQARU'95)*, 1995.
- [Cho98] L. Cholvy. Reasoning about data provided by federated deductive databases. *Journal of Intelligent Information Systems*, 10(1), 1998.
- [CMH⁺94] S. Chawathe, H. Garcia Molina, J. Hammer, K. Ireland, Y. Papakonstantinou, J. Ullman, and J. Widom. The tsimmi project: Integration of heterogeneous information sources. In *Proceedings of IPSJ Conference*, pages 7–18, October 1994.
- [Dal88a] M. Dalal. Investigations into a theory of knowledge base revision: preliminary report. In *Proceedings of the National Conference on Artificial Intelligence (AAAI'88)*, pages 475–479, 1988.
- [Dal88b] M. Dalal. Updates in propositional databases. Technical report, Rutgers University, 1988.
- [dC85] Marquis de Condorcet. *Essai sur l'application de l'analyse à la probabilité des décisions rendues à la pluralité des voix*. Paris, 1785.
- [DdSCP95] D. Dubois, F. Dupin de Saint-Cyr, and H. Prade. Update postulates without inertia. In *Proceedings of the Third European Conference on Symbolic and Quantitative Approaches to Reasoning and Uncertainty (ECSQARU'95)*, LNAI 946, pages 162–170, 1995.
- [del97] A. del Val. Belief revision and non-monotonic reasoning: Syntactic, semantic, foundational and coherence approaches. *Journal of Applied Non-Classical Logics*, 7:213–240, 1997.
- [DFP96] D. Dubois, H. Fargier, and H. Prade. Refinements of the max-min approach to decision-making in fuzzy environment. *Fuzzy Sets and Systems*, 81:103–122, 1996.
- [dK86] J. de Kleer. An assumption-based TMS. *Artificial Intelligence*, 28(2):127–162, 1986.
- [DLP94] D. Dubois, J. Lang, and H. Prade. *Handbook of Logic in Artificial Intelligence and Logic Programming*, volume 3: Nonmonotonic Reasoning and Uncertain Reasoning, chapter Possibilistic Logic. Oxford Science Publications, 1994.
- [DNP94] C. Damasio, W. Nedjdl, and L. M. Pereira. Revise: An extended logic programming system for revising knowledge bases. In *Proceedings of the Fourth International Conference on Principles of Knowledge Representation and Reasoning*, pages 607–618. Morgan Kaufmann, 1994.
- [Doy79] J. Doyle. A truth maintenance system. *Artificial Intelligence*, 12(3):231–272, 1979.
- [Doy92] J. Doyle. *Belief revision*, chapter Reason maintenance and belief revision: foundations versus coherence theories, pages 29–51. Cambridge University Press, 1992.
- [DP91] D. Dubois and H. Prade. Epistemic entrenchment and possibilistic logic. *Artificial Intelligence*, 50:223–239, 1991.

- [DP94] A. Darwiche and J. Pearl. On the logic of iterated belief revision. In Morgan Kaufmann, editor, *Theoretical Aspects of Reasoning about Knowledge: Proceedings of the 1994 Conference (TARK'94)*, pages 5–23, 1994.
- [DP97] A. Darwiche and J. Pearl. On the logic of iterated belief revision. *Artificial Intelligence*, 89:1–29, 1997.
- [dSC96] F. Dupin de Saint-Cyr. *Gestion des croyances et de l'évolutif en logiques pondérées*. PhD thesis, Université Paul Sabatier, Toulouse, 1996.
- [dSCLS94] F. Dupin de Saint-Cyr, J. Lang, and T. Schiex. Gestion de l'inconsistance dans les cases de connaissances: une approche syntaxique basée sur la logique des pénalités. In *Actes du neuvième congrès Reconnaissance des Formes et Intelligence Artificielle*, 1994.
- [Dub86] D. Dubois. Belief structures, possibility theory and decomposable confidence measures on finite sets. *Computers and Artificial Intelligence (Bratislava)*, 5(5):403–416, 1986.
- [EG92] T. Eiter and G. Gottlob. On the complexity of propositional knowledge base revision, updates, and counterfactuals. *Artificial Intelligence*, 57(2-3):227, 270 1992.
- [Eth88] D. Etherington. *Reasoning with incomplete information*. Morgan Kaufmann, 1988.
- [Fer99a] E. Fermé. Five faces of recovery. In H. Rott and M. A. Williams, editors, *Frontiers of Belief Revision*. Kluwer, 1999.
- [Fer99b] E. Fermé. *Revising the AGM Postulates*. PhD thesis, University of Buenos Aires, 1999.
- [FH96] N. Friedman and J.Y. Halpern. Belief revision: a critique. In *Proceedings of the Fifth International Conference on Principles of Knowledge Representation and Reasoning (KR'96)*, pages 421–431, 1996.
- [FH99] E. Fermé and S. O. Hansson. Selective revision. *Studia Logica*, 63(3):331–342, 1999.
- [FKUV86] R. Fagin, G. M. Kuper, J. D. Ullman, and M. Y. Vardi. Updating logical databases. *Advances in Computing Research*, 3:1–18, 1986.
- [FL94] M. Freund and D. Lehmann. Belief revision and rational inference. Technical Report TR-94-16, Institute of Computer Science, The Hebrew University of Jerusalem, 1994.
- [Fre93] M. Freund. Injective models and disjunctive relations. *Journal of Logic and Computation*, 3:231–247, 1993.
- [Fuh91] A. Fuhrmann. Theory contraction through base contraction. *Journal of Philosophical Logic*, 20:175–203, 1991.
- [FUV83] R. Fagin, J. D. Ullman, and M. Y. Vardi. On the semantics of updates in databases. In *Proceedings of the 2nd ACM SIGACT-SIGMOD Symposium on the Principles of Database Systems*, pages 352–365, 1983.
- [Gab85] D. M. Gabbay. Theoretical foundations for nonmonotonic reasoning in experts systems. In K. Apt, editor, *Logic and Models of Concurrent Systems*. Springer Verlag, 1985.
- [Gär88] P. Gärdenfors. *Knowledge in flux*. MIT Press, 1988.

- [Gär90] P. Gärdenfors. Belief revision and nonmonotonic logic: two sides of the same coin? In *Proceeding of the Ninth European Conference on Artificial Intelligence (ECAI'90)*, pages 768–773, 1990.
- [Gef92] H. Geffner. *Default reasoning: Causal and Conditional Theories*. MIT Press, 1992.
- [GP96] M. Goldszmidt and J. Pearl. Qualitative probabilities for default reasoning, belief revision, and causal modeling. *Artificial Intelligence*, 84:57–112, 1996.
- [Gro88] A. Grove. Two modellings for theory change. *Journal of Philosophical Logic*, 17(157-180), 1988.
- [Han89] S. O. Hansson. New operators for theory change. *Theoria*, 55:114–132, 1989.
- [Han91a] S. O. Hansson. *Belief base dynamics*. PhD thesis, Uppsala University, 1991.
- [Han91b] S. O. Hansson. Belief contraction without recovery. *Studia Logica*, 50(2):251–260, 1991.
- [Han93] S. O. Hansson. Reversing the levi identity. *Journal of Philosophical Logic*, 22:637–669, 1993.
- [Han97] S. O. Hansson. Semi-revision. *Journal of applied non-classical logic*, 7:151–175, 1997.
- [Han98a] S. O. Hansson. *Handbook of Defeasible Reasoning and Uncertainty Management Systems*, volume 3, chapter Revision of belief sets and belief bases, pages 17–75. Kluwer, 1998.
- [Han98b] S. O. Hansson. What's new isn't always best. *Theoria*, pages 1–13, 1998. Special issue on non-prioritized belief revision.
- [Har86] G. Harman. *Change in view: principles of reasoning*. MIT Press, 1986.
- [HR98] A. Herzig and O. Rifi. Update operations: a review. In *Proceedings of the Thirteenth European Conference on Artificial Intelligence (ECAI'98)*, pages 13–17, 1998.
- [Kel78] J. S. Kelly. *Arrow impossibility theorems*. Series in economic theory and mathematical economics. Academic Press, New York, 1978.
- [Kel88] J. S. Kelly. *Social Choice Theory: An Introduction*. Springer-Verlag, 1988.
- [KLM90] S. Kraus, D. Lehmann, and M. Magidor. Nonmonotonic reasoning, preferential models and cumulative logics. *Artificial Intelligence*, 44:167–207, 1990.
- [KM91a] H. Katsuno and A. O. Mendelzon. On the difference between updating a knowledge base and revising it. In *Proceedings of the Second International Conference on Principles of Knowledge Representation and Reasoning (KR'91)*, pages 387–394, 1991.
- [KM91b] H. Katsuno and A. O. Mendelzon. Propositional knowledge base revision and minimal change. *Artificial Intelligence*, 52:263–294, 1991.
- [Kon96a] S. Konieczny. Révision de la connaissance et coopération. Technical report, Laboratoire d'Informatique Fondamentale de Lille, juillet 1996. Mémoire de DEA.
- [Kon96b] S. Konieczny. Vers une modélisation de la coopération. In F. Anceaux and J. M. Coquery, editors, *Actes du sixième Colloque de l'Association pour la Recherche Cognitive*, pages 227–231, décembre 1996.
- [Kon98] S. Konieczny. Operators with memory for iterated revision. Technical report, LIFL, 1998. IT-314.

- [KP97] S. Konieczny and R. Pino Pérez. On the difference between arbitration and majority merging. In *4th Workshop on Logic, Language, Information and Computation*, Fortaleza, Brasil, 1997.
- [KP98] S. Konieczny and R. Pino Pérez. On the logic of merging. In *Proceedings of the Sixth International Conference on Principles of Knowledge Representation and Reasoning (KR'98)*, pages 488–498, 1998.
- [KP99a] S. Konieczny and R. Pino Pérez. Fusion avec contraintes d'intégrité. In *Journées Nationales sur les Modèles de raisonnement (JNMR'99)*, Paris, 22-23 Mars 1999. Electronic proceedings <http://www.irit.fr/GDRI3-ModRais/articlesJNMR.html>.
- [KP99b] S. Konieczny and R. Pino Pérez. Merging with integrity constraints. In *Proceedings of the Fifth European Conference on Symbolic and Quantitative Approaches to Reasoning with Uncertainty (ECSQARU'99)*, Lecture Notes in Artificial Intelligence 1638, pages 233–244, 1999.
- [KW85] A. M. Keller and M. Winslett. On the use of an extended relational model to handle changing incomplete information. *IEEE Transaction on Software Engineering*, SE-11(7):620–633, 1985.
- [Lan91] J. Lang. *Logique possibiliste : aspects formels, déduction automatique et applications*. PhD thesis, Université Paul Sabatier Toulouse, 1991.
- [Leh95] D. Lehmann. Belief revision, revised. In *Proceedings of the Fourteenth International Joint Conference on Artificial Intelligence (IJCAI'95)*, pages 1534–1540, 1995.
- [Lev88] I. Levi. Iteration of conditionals and the ramsey test. *Synthese*, 76:49–81, 1988.
- [Lew73] D. Lewis. *Counterfactuals*. Basil Blackwell, 1973.
- [Lib97] P. Liberatore. The complexity of iterated belief revision. In *Proceedings of the Sixth International Conference on Database Theory (ICDT'97)*, pages 276–290, 1997.
- [Lin95] J. Lin. *Frameworks for dealing with conflicting information and applications*. PhD thesis, University of Toronto, 1995.
- [LM92] D. Lehmann and M. Magidor. What does a conditional knowledge base entail? *Artificial Intelligence*, 55:1–60, 1992.
- [LM98] J. Lin and A. O. Mendelzon. Merging databases under constraints. *International Journal of Cooperative Information System*, 7(1):55–76, 1998.
- [LM99] J. Lin and A. O. Mendelzon. Knowledge base merging by majority. In *Dynamic Worlds: From the Frame Problem to Knowledge Management*. Kluwer, 1999.
- [LR91] S. Lindström and W. Rabinowicz. Epistemic entrenchment with incomparabilities and relational belief revision. In A. Furmann and M. Morreau, editors, *The logic of theory change*, pages 93–126. 1991.
- [LS95] P. Liberatore and M. Schaerf. Arbitration: A commutative operator for belief revision. In *Proceedings of the Second World Conference on the Fundamentals of Artificial Intelligence*, pages 217–228, 1995.
- [LS98] P. Liberatore and M. Schaerf. Arbitration (or how to merge knowledge bases). *IEEE Transactions on Knowledge and Data Engineering*, 10(1):76–90, 1998.
- [Mak87] D. Makinson. On the status of the postulate of recovery in the logic of theory change. *Journal of Philosophical Logic*, 16:383–394, 1987.

- [Mak89] D. Makinson. General theory of cumulative inference. In *Non-Monotonic Reasoning*, volume 346 of *Lectures Notes in Artificial Intelligence*, pages 1–18, 1989.
- [Mak94] D. Makinson. *Handbook of Logic in Artificial Intelligence and Logic Programming*, volume III, chapter General Pattern in nonmonotonic reasoning, pages 35–110. Clarendon Press, Oxford, 1994.
- [Mak98] D. Makinson. Screened revision. *Theoria*, pages 14–23, 1998. Special issue on non-prioritized belief revision.
- [MG89] D. Makinson and P. Gärdenfors. Relations between the logic of theory change and nonmonotonic logic. In *The Logic of Theory Change, Workshop, Konstanz, FRG*, volume 465 of *LNAI*, pages 185–205, 1989.
- [MIR94] R. J. Miller, Y. E. Ioannidis, and R. Ramakrishnan. Schema equivalence in heterogeneous systems: bridging theory and practice. In *Proceedings of the International Conference on Extending Database Technology*, 1994.
- [Mou88] H. Moulin. *Axioms of cooperative decision making*. Monograph of the Econometric Society. Cambridge University Press, 1988.
- [MT94] V. Marek and M. Truszczyński. Revision specifications by means of programs. In *Proceedings of JELIA'94*. Springer Verlag, 1994. Lecture Notes in Artificial Intelligence.
- [MT95] V. Marek and M. Truszczyński. Revision programming, database updates and integrity constraints. In *Proceedings of the Fifth International Conference of Database Theory*, pages 368–382, 1995. Lecture Notes in Computer Science 893.
- [Neb89] B. Nebel. A knowledge level analysis of belief revision. In *Proceedings of the First International Conference on Principles of Knowledge Representation and Reasoning (KR'89)*, pages 301–311, 1989.
- [Neb91] B. Nebel. Belief revision and default reasoning: syntax-based approaches. In *Proceedings of the Second International Conference on Principles of Knowledge Representation and Reasoning (KR'91)*, pages 417–428, 1991.
- [Neb94] B. Nebel. Base revision operations and schemes: semantics, representation, and complexity. In *Proceedings of the Eleventh European Conference on Artificial Intelligence (ECAI'94)*, pages 341–345, 1994.
- [Neb96] B. Nebel. *How hard is it to revise a belief base?* PhD thesis, University of Freiburg, 1996.
- [NFPS94] A. C. Nayak, N. Y. Foo, M. Pagnucco, and A. Sattar. Entrenchment kinematics 101. In *Proceedings of the Seventh Australian Joint Conference on Artificial Intelligence*, pages 157–164, 1994.
- [NFPS96] A. C. Nayak, N. Y. Foo, M. Pagnucco, and A. Sattar. Changing conditional beliefs unconditionally. In *Proceedings of the Sixth Conference of Theoretical Aspects of Rationality and Knowledge (TARK'96)*, pages 119–135. Morgan Kaufmann, 1996.
- [Nie91] R. Niederée. Multiple contraction. a further case against gärdenfors' principle of recovery. In A. Furmann and M. Morreau, editors, *The Logic of Theory Change*, pages 322–334. 1991.
- [Par97] V. Pareto. *Cours d'économie politique*. Rouge, Lausanne, 1897.
- [Pea90] J. Pearl. System z: a natural ordering of defaults with tractable applications to nonmonotonic reasoning. In *Proceedings of the Third Conference on Theoretical Aspects of Reasoning About Knowledge (TARK'90)*, 1990.

- [PU99] R. Pino Pérez and C. Uzcategui. On representation theorems for nonmonotonic inference relation. *Journal of Symbolic Logic*, 1999. To appear.
- [Ram31] F. P. Ramsey. *Foundations of Mathematics and Other Logical Essays*, chapter General Propositions and Causality, pages 237–257. 1931.
- [Rei80] R. Reiter. A logic for default reasoning. *Artificial Intelligence*, 13:81–132, 1980.
- [Rev93] P. Z. Revesz. On the semantics of theory change: arbitration between old and new information. In *Proceedings of the 12th ACM SIGACT-SIGMOD-SIGART Symposium on Principles of Databases*, pages 71–92, 1993.
- [Rev97] P. Z. Revesz. On the semantics of arbitration. *International Journal of Algebra and Computation*, 7(2):133–160, 1997.
- [Rya91] M. D. Ryan. Defaults and revision in structured theories. In *Proceedings of the Sixth IEEE Symposium on Logic in Computer Science (LICS'91)*, pages 362–373. Morgan Kaufmann, 1991.
- [Rya92] M. D. Ryan. *Ordered Presentations of Theories*. PhD thesis, Imperial College, London, 1992.
- [Rya94] M. D. Ryan. Belief revision and ordered theory presentations. In A. Fuhrmann and H. Rott, editors, *Logic, Action and Information*. De Gruyter Publishers, 1994. Also in *Proceedings of the Eighth Amsterdam Colloquium on Logic*, University of Amsterdam, 1991.
- [SAB⁺] V.S. Subrahmanian, Sibel Adali, Anne Brink, Ross Emery, James J. Lu, Adil Rajput, Timothy J. Rogers, Robert Ross, and Charles Ward. Hermes: Heterogeneous reasoning and mediator system.
- [Sav71] L. J. Savage. *The foundations of statistics*. Dover Publications, New York, 1971. Second revised edition.
- [Sch98] K. Schlechta. Non-prioritized belief revision based on distances between models. *Theoria*, pages 34–53, 1998. Special issue on non-prioritized belief revision.
- [Sen79] A. K. Sen. *Collective Choice and Social Welfare*. Advanced Textbooks in Economics. Elsevier, 1979.
- [SK94] P. Smets and R. Kennes. The transferable belief model. *Artificial Intelligence*, 66(2):191–234, 1994.
- [SLM96] K. Schlechta, D. Lehmann, and M. Magidor. Distance semantics for belief revision. In *Proceedings of the Sixth Conference on Theoretical Aspects of Rationality and Knowledge (TARK'96)*, pages 137–145, 1996.
- [Spo87] W. Spohn. Ordinal conditional functions: a dynamic theory of epistemic states. In W. L. Harper and B. Skyrms, editors, *Causation in Decision, Belief Change, and Statistics*, volume 2, pages 105–134. 1987.
- [Spo90] W. Spohn. A general non-probabilistic theory of inductive reasoning. In *Uncertainty in Artificial Intelligence*, pages 149–158. Elsevier Science, 1990.
- [SSU91] A. Silberschatz, M. Stronebraker, and J. D. Ullman. Database systems: Achievements and opportunities. *Communications of the ACM*, 34(10):110–120, 1991.
- [Sta68] R. C. Stalnaker. *Ifs*, chapter A theory of conditionals, pages 41–55. D. Reidel, Dordrecht, 1968.
- [Sub94] V. S. Subrahmanian. Amalgamating knowledge bases. *ACM Transactions on Database systems.*, 19(2):291–331, 1994.

- [Tar56] A. Tarski. *Logic, Semantics, Metamathematics*. Clarendon Press Oxford, 1956. Papers from 1923 to 1938. Translated by J. H. Woodger.
- [Ull88] J. D. Ullman. *Principles of databases and knowledge-base systems*. Computer Science Press, Rockville, MD, 1988.
- [WFPS95] M. A. Williams, N. Foo, M. Pagnucco, and B. Sims. Determining explanations using transmutations. In *Proceedings of the 11th International Joint Conference on Artificial Intelligence (IJCAI'95)*, pages 822–830, 1995.
- [Wil94] M. A. Williams. Transmutations of knowledge systems. In *Proceedings of the Fourth International Conference on the Principles of Knowledge Representation and Reasoning (KR'94)*, pages 619–629, 1994.
- [Win88] M. Winslett. Reasoning about action using a possible model approach. In *Proceedings of the Seventh National Conference on Artificial Intelligence (AAAI'88)*, pages 89–93, 1988.
- [Win90] M. Winslett. *Updating logical databases*. Cambridge University Press, 1990.
- [Zad78] L. A. Zadeh. Fuzzy sets as a basis of a theory of possibility. *Fuzzy sets and Systems*, 1:3–28, 1978.



Sur la logique du changement : révision et fusion de bases de connaissance

Résumé

Un des problèmes majeurs dans la formalisation et la modélisation du raisonnement est la prise en compte de son aspect dynamique. C'est-à-dire la capacité à raisonner à partir de connaissances imparfaites ou incomplètes, pouvant être remises en doute ultérieurement. La théorie de la révision étudie les changements qui interviennent dans une base de connaissance lorsqu'une nouvelle information, plus fiable, est incorporée. Pour l'étude logique de la révision c'est le cadre AGM (d'après ses trois auteurs Alchourrón, Gärdenfors, Makinson) qui s'est imposé. Cette thèse est une contribution à ce cadre.

Le cadre AGM ne contraint pas suffisamment la révision pour assurer un bon comportement lors de l'itération. Nous définissons une classe d'opérateurs basés sur la notion de primauté forte de la nouvelle information. Nous montrons que ces opérateurs ont un bon comportement pour l'itération.

D'un point de vue pratique, les opérateurs de révision usuels posent des problèmes de représentation et de calcul. Nous définissons et étudions des opérateurs syntaxiques basés sur une logique induite par le chaînage avant.

Un problème similaire à celui de la révision se pose lorsque l'on veut extraire une connaissance d'un ensemble d'informations contradictoires provenant de sources différentes. On parle alors de fusion de bases de connaissance. Nous proposons une caractérisation logique pour ces opérateurs, généralisation du cadre AGM pour la révision.

Cette caractérisation logique peut être utilisée pour comparer différentes méthodes de fusion ou pour suggérer des adaptations de méthodes existantes. Pour illustrer ceci nous étudions des opérateurs de fusion syntaxique basés sur la sélection de sous-ensembles maximaux consistants de la réunion des bases de connaissance à fusionner. Nous montrons que les opérateurs existants ne satisfont que très peu de propriétés logiques et que l'utilisation de fonctions de sélection permet d'augmenter le nombre de propriétés vérifiées.

Mots-clés: Intelligence artificielle, logiques non monotones, révision, AGM, itération, bases de connaissances, fusion, agrégation de préférences

On the Logic of Theory Change : Belief Revision and Knowledge Base Merging

Abstract

One of the main issue concerning the formalization and the modelization of reasoning concerns its dynamic nature. That is the ability of reasoning with imperfect or incomplete knowledge that can be questioned later on. Belief revision is concerned with the change produced in a knowledge base when a new (more reliable) evidence is incorporated. In the logical context, the most successful framework is the AGM one (from its authors Alchourrón, Gärdenfors, Makinson). This thesis is a contribution to this framework.

The AGM framework do not constrained sufficiently the revision operators to ensure good properties for the iteration process. We define a family of operators based on the notion of strong primacy of the new information. We show that those operators have good properties with respect to iteration.

From a practical point of view, classical revision operators are not easily computable and faced representation problems. We define some syntactical revision operators based on a logic induced by the forward chaining.

A problem close to belief revision arises when one wants to extract some knowledge from a set of information coming from different sources. We propose a logical characterization of merging operators. This characterization is a generalization of the AGM framework for belief revision.

This logical characterization can be used to compare various merging methods or to suggest improvements of some of them. We illustrate it by studying syntactical merging operators. Those operators select some maximal consistent subsets in the union of the knowledge bases. We show that the operators proposed in the literature do not satisfy a lot of logical properties. Then we show how the use of suitable selection functions can improve the number of satisfied logical properties.

Keywords: Artificial intelligence, nonmonotonic logics, belief revision, AGM, iteration, knowledge bases, merging, aggregation.