

UNIVERSITÉ DU QUÉBEC

THÈSE PRÉSENTÉE À
L'UNIVERSITÉ DU QUÉBEC À TROIS-RIVIÈRES

COMME EXIGENCE PARTIELLE
DU DOCTORAT EN PHILOSOPHIE

PAR
ANA LÊDA DE ARAÚJO

LA PROBLÉMATIQUE DE LA PERTINENCE PRAGMATIQUE

NOVEMBRE 2000

Université du Québec à Trois-Rivières

Service de la bibliothèque

Avertissement

L'auteur de ce mémoire ou de cette thèse a autorisé l'Université du Québec à Trois-Rivières à diffuser, à des fins non lucratives, une copie de son mémoire ou de sa thèse.

Cette diffusion n'entraîne pas une renonciation de la part de l'auteur à ses droits de propriété intellectuelle, incluant le droit d'auteur, sur ce mémoire ou cette thèse. Notamment, la reproduction ou la publication de la totalité ou d'une partie importante de ce mémoire ou de cette thèse requiert son autorisation.

À

Monsieur le professeur Matias Francisco Dias qui, dès le baccalauréat, m'a initiée et guidée dans l'étude de la Logique, et qui, avec une affection presque paternelle, a pris une très grande part dans ma formation académique et professionnelle.

REMERCIEMENTS

Une thèse de doctorat n'est jamais le fruit d'un travail réalisé dans l'isolement. Tout chercheur, et c'est inévitable, profite de multiples rencontres qui laissent des traces, je veux dire, des suggestions, des idées, des pistes, sans lesquelles personne n'arriverait à philosopher avec pertinence et universalité. Je tiens ici à remercier tous ceux et toutes celles qui m'ont aidé, directement ou indirectement, à élaborer ce travail de Philosophie Analytique, plus particulièrement en Philosophie du Langage et de la Logique.

Toutefois, je me dois de mentionner ici certaines institutions et certains individus qui ont directement contribué à la réalisation de cette recherche.

Au CNPq, organisme de soutien à la recherche du Gouvernement Brésilien, qui m'a octroyé une bourse d'études sans laquelle ce travail n'aurait jamais pu être réalisé.

À Monsieur le professeur Daniel Vanderveken, mon directeur de thèse, qui, le premier, m'a montré toute l'importance d'une étude approfondie des actes de discours, pour la philosophie en général, pour la philosophie du langage aussi, et plus particulièrement pour un traitement adéquat du

problème de la pertinence.

À Monsieur le professeur André Leclerc, de l'Université Fédérale de la Paraíba - Brésil, qui a accepté la codirection de cette étude entièrement rédigée au Brésil. En plus d'accompagner, pas à pas, la rédaction de mon travail avec patience et sérieux, présentant ici et là plusieurs critiques et suggestions, il a dû également corriger "mon français" afin de le rendre plus académique. Sans lui, j'en suis sûre, cette thèse n'aurait jamais vu le jour.

Je remercie également les professeurs du Département de philosophie de l'UQTR qui m'ont accueillie à Trois-Rivières, en particulier Mme Suzanne Foisy, M. Nicolas Kaufmann, M. Claude Panaccio et M. Daniel Vanderveken; je remercie aussi les professeurs M. Alain Voizard de l'UQAM et M. Elias Humberto Alves de l'UNICAMP – Brésil. Sans leur soutien, je n'aurais jamais pu réaliser ce travail philosophique où la rigueur et la précision sont indispensables.

Je dois aussi remercier le Département de philosophie de l'Université Fédérale de la Paraíba, qui m'a accordé le détachement nécessaire à la réalisation de ma recherche, et plus particulièrement le professeur Giovanni da Silva de Queiróz qui m'a consacré beaucoup de son temps en discussions fort riches, surtout pour le chapitre III de notre travail.

À tous ceux et toutes celles qui, au Canada ou au Brésil, m'ont toujours soutenu au cours de ma recherche, grand merci! Sans pouvoir, malheureusement citer tous le monde, je me dois cependant d'exprimer ma gratitude à Mme Hélène Boisclair, Mme Cécile Juneau, Mme Myrta Nogueira Dias, Mme Maria Vilma de Lucena, Mme Gesuina Elias Leclerc, Mme Cândida Jacy Melo, et en particulier. à M. Dr. Grinberg Medeiros Botelho.

Et finalement, je remercie les membres de ma famille (spécialement mon père, *in memoriam*), ma mère, ma sœur Roseane Araújo Montenegro et mon beau-frère Geraldo Montenegro, pour l'appui émotionnelle et l'affection qu'ils m'ont sans cesse prodiguées, ainsi que ma fille Janaína pour sa patience. Je remercie de tout mon cœur Antonio Rufino Vieira, pour ses encouragements constants, son appui, et en particulier sa patience et sa tendresse, qualités qu'il a toujours possédées, et qu'il m'a prodiguées en double pendant toute la durée de mes études doctorales.

TABLE DES MATIÈRES

Introduction	9
1- Importance de la problématique	9
2- Intérêt de la perspective adoptée	12
 Chap. I - Présentation de la théorie de la pertinence de Sperber et Wilson	25
Introduction	25
1- Le modèle du code, le modèle inférentiel et la communication	26
2- Le modèle inférentiel de Sperber et Wilson	51
3- Définition de la pertinence chez Sperber et Wilson	69
 Chap. II - Appréciations de quelques critiques sur le dispositif déductif de Sperber et Wilson	77
Introduction.....	77
1-Quelques critiques adressées au dispositif déductif de Sperber et Wilson	79
2- Quelques considérations critiques à propos de la réponse de Sperber et Wilson à la critique de Hinkelman.....	94
 Chap. III - Définitions de la pertinence et la logique de la pertinence P	115
Introduction	115
1- Considérations historiques	116
2- Définitions	127
3- Langage et axiomatique pour la logique P	141
 Chap. IV - Quelques applications de la logique de la pertinence P à la théorie de la pertinence de Sperber et Wilson	157
Introduction	157
1- Le problème de la réitération <i>ad infinitum</i> dans la définition de l'implication triviale chez Sperber et Wilson	158
2-La définition de degrés de pertinence chez Sperber et Wilson versus la définition de degrés de pertinence dans la Logique P	167
3-Convergence entre la seconde clause de la présomption de pertinence optimale chez Sperber et Wilson et la maxime de quantité de Grice	194

Chap. V - Présentation de la Théorie des Actes de Discours de Searle et Vanderveken	205
Introduction	205
1- Considérations historiques	206
2- La sémantique des actes de discours	213
3- La pragmatique des actes de discours	222
4- L'analyse de la conversation	224
Chap. VI - Considérations critiques et applications de la logique de la pertinence <i>P</i> à la Théorie des Actes de Discours de Searle et Vanderveken	233
Introduction	233
1-Considérations critiques à la pragmatique des actes de discours à la lumière de la logique de la pertinence <i>P</i>	235
2-La logique de la pertinence <i>P</i> versus la logique propositionnelle dans la sémantique formelle de la théorie des actes de discours.....	257
3-Applications de la logique de la pertinence <i>P</i> aux forces illocutoires des actes de discours	267
Conclusion	283
Bibliographie	292

INTRODUCTION

L'objectif de notre thèse est de discuter de la problématique de la pertinence pragmatique et de présenter une bonne définition de la notion de *pertinence* (en anglais: *relevance*) dans la pragmatique en général, notamment en théorie de la conversation.

Notre intérêt pour ce thème vient du fait que jusqu'à présent il n'y a pas de définition précise et rigoureuse de la notion de pertinence qui puisse servir à expliquer les *jugements de pertinence* que font les participants d'une conversation pour comprendre la signification des locuteurs et pour décider si ce qu'ils disent est pertinent ou non dans le contexte de leurs énonciations.

1. Importance de la problématique

La notion de pertinence est essentielle pour toute communication humaine, orale ou écrite. C'est à partir des jugements de pertinence que les interlocuteurs déterminent la signification du locuteur, qu'elle soit littérale, non

littérale, ou plus riche que celle de l'énoncé que le locuteur utilise, et qu'ils planifient leur réponse ou répliques dans une conversation ordinaire.

Dans le cas d'énonciations littérales, la notion de pertinence est fondamentale, car c'est à partir d'elle que se mettent en branle les processus mentaux des individus qui leur permettent de déterminer, dans certains contextes ou situations, s'il y a correspondance entre la signification du locuteur et la représentation sémantique de l'énoncé qu'il utilise. Dans le cas d'énonciations non littérales telles que les *actes de discours indirects*, les *métaphores*, l'*ironie* et dans le cas des *implicatures conversationnelles*, la notion de pertinence est également fondamentale: elle permet d'identifier un écart entre la signification du locuteur et celle de l'énoncé qu'il utilise, et dès lors, elle incite les interlocuteurs à faire usage des capacités ou habiletés qui leur permettront de comprendre ce que veut dire le locuteur. Nous faisons donc, tout au long de notre thèse, l'hypothèse que dans les cas d'énonciations littérales aussi bien que non littérales, la pertinence joue un rôle décisif pour le succès de la compréhension: lorsqu'il s'agit d'expliquer les processus de communication, la notion de pertinence doit venir en premier lieu puisqu'elle est présupposée par n'importe quelle description des normes implicites de la communication verbale.

Discuter de la pertinence conversationnelle nous place sur le terrain

de la philosophie pratique, puisqu'elle a tout à voir avec l'atteinte de certaines fins; de ce point de vue, la compréhension est une activité dans la mesure où elle vise à interpréter l'usage que font les agents des *règles structurales* et/ou *stratégiques* du dialogue.

À ce sujet, le contexte d'usage est fondamental car c'est en lui que s'établit la pertinence de ce qui est en question. Soit un énoncé du type: "Je vais à l'université le dimanche"; dépendamment du contexte, cet énoncé peut être une promesse (je te promets d'aller à l'université le dimanche), une confession (je t'avoue que je vais à l'université ce dimanche), une excuse (je vais manquer ton anniversaire ce dimanche), etc. L'interprétation de mots comme *ici*, *maintenant*, *je*, etc., dépend toujours du contexte de leur énonciation. Par exemple, quand je dis à quelqu'un: "il fait beau" alors qu'il pleut très fortement au dehors, il suffit de vérifier dans ce contexte que mon énonciation ne correspond pas au fait qu'il pleut au dehors. Ainsi, on pourrait dire que la notion de pertinence est une notion contextuelle, c'est-à-dire, dans le sens où elle fonctionne toujours dans la situation globale où se situe le dialogue et où sont respectées préalablement, les règles linguistiques, grammaticales, le sens de la phrase, ainsi qu'une correspondance avec un état de choses existant. Cela étant dit, c'est dans les recherches pragmatiques sur la communication que nous chercherons un sens précis pour le concept de

pertinence.

2. Intérêt de la perspective adoptée

De nombreux auteurs ont déjà discuté des phénomènes pragmatiques dans la communication, en particulier, du problème de la pertinence pragmatique et son rôle en conversation¹. Parmi eux, nous citons le célèbre livre *Relevance* (1986), de Dan Sperber et Deirdre Wilson (traduction française, *La Pertinence*, (1989)), dont l'objectif principal est d'élaborer une théorie de la cognition et de la communication qui intègre un *Principe de Pertinence*. La notion centrale de cette théorie cognitive est la notion de pertinence. Il s'agit d'une théorie pragmatique postgricéenne, développée dans le même courant d'idées que Grice dans le sens où Sperber et Wilson soutiennent eux aussi que la récupération du sens communiqué implicitement par un énoncé se fait par un calcul inférentiel - quoique, comme les auteurs eux-

¹ Pour plus d'informations à ce sujet, voir: Carston et Uchida (1998); Cherniak (1986); Cooper (1974); Dascal (1982); Fillmore (1971); Gazdar (1979); Ghiglione et Trognon (1993); Grice (1975); Harder et Kock (1976); Holderoit (1987); Jacques (1985); Keeane. (1971); Kasher (1976); Leech (1983); Levinson (1979); (1983); Moeschler et Reboul (1996); Parret (1988); Récanati (1981); Searle et al. (1980, 1990); Sperber et Wilson

mêmes le soulignent, il existe des divergences fondamentales entre leur théorie de la pertinence et l'approche de Grice - divergences que nous indiquerons dans le développement de notre thèse.

Dans l'approche de Grice (1975, 1979), le calcul inférentiel qui est fait pour récupérer la signification intentionnelle d'une énonciation est déclenché, parmi d'autres recours, par le respect ou non des règles pragmatiques (principe de coopération et maximes conversationnelles). Grice a fait l'hypothèse que toute conversation rationnelle est régie par un *principe général de coopération conversationnelle* qu'il a formulé de la façon suivante: à chaque étape d'une conversation, faites en sorte que votre contribution soit telle que requise par le but ou la direction acceptée de l'échange verbal dans lequel vous participez. Grice a développé ce principe en quatre catégories de *maximes conversationnelles*: la maxime de *quantité* (que votre contribution soit aussi informative qu'il est requis), la maxime de *qualité* (essayez de dire ce que vous croyez être vrai), la maxime de *relation* (soyez pertinent) et la maxime de *manière* (soyez clair, bref, ordonné). Grice a introduit la notion d'*implicature*² *conversationnelle* en tant que signification impliquée d'un énoncé qui peut être

(1987); Vanderveken (1988, 1990, 1990a, 1991b, 1992, 1997a, 1999a, 1999b); Verschueren (1979); Ziv (1988).

² Dans Sperber et Wilson (1989), ils introduisent l'expression "implication" au lieu d'*implicature*. Pour cette raison, nous ferons dès maintenant l'usage du mot "implication"

calculée au moyen de maximes conversationnelles. Quand quelqu'un m'invite à faire du ski de fond et que je lui réponds que je n'aime pas les sports d'hiver, ma réponse donne lieu à l'implication conversationnelle que je ne veux pas faire du ski de fond (sans quoi la maxime de quantité serait violée). Mais il se peut que des maximes soient violées, et le cas échéant, en supposant que le communicateur rationnel respecte le principe de coopération, l'interlocuteur cherchera, à partir de son énonciation, à calculer des implications conversationnelles.

À la suite de Austin (1962), s'est développée aussi la *théorie des actes de discours* de John R. Searle et Daniel Vanderveken (1985)³, devenue aujourd'hui une branche importante de la philosophie contemporaine du langage. La contribution de la Théorie des Actes de Discours (T.A.D.) consiste à reconstruire formellement les procédures que le communicateur utilise pour donner les indices nécessaires à l'auditeur afin qu'il puisse interpréter de tels

pour signifier la même chose que *implicature*.

³ La Théorie des Actes de Discours est une élaboration à partir de certaines idées de John L. Austin (1962) et Grice (1975, 1979). Les *actes de discours* qui sont importants pour la théorie de la signification sont du type appelé par Austin actes *illocutoires*. Il a considéré que par leurs énonciations, les locuteurs entendent, d'abord, accomplir des actes de discours du type illocutoire, tels que des assertions, des excuses, des questions, etc. Son idée est que de tels actes illocutoires sont les unités premières de signification et de communication dans l'usage et la compréhension du langage en général. Plus tard, Searle (1969, 1972) a formulé des conditions et des règles pour les actes de discours, et postérieurement, Searle et Vanderveken (1985) élaborèrent les fondements de la logique illocutoire. Nous reviendrons sur ce sujet au chapitre V de notre thèse.

indices et faire des inférences pour calculer la signification du communicateur. Cette théorie rend compte des actes de discours non littéraux tels que les actes de discours indirects, et les métaphores où ce que le locuteur signifie est différent de ce qu'il dit (dans les cas d'énonciations littérales, la signification du locuteur est identique à celle de l'énoncé utilisé). Dans le cas des actes de discours indirects, par exemple, ce que le locuteur veut dire dépasse ce que signifie l'énoncé qu'il utilise. Dans l'approche de Searle et Vanderveken, les actes illocutoires élémentaires complets sont composés d'une *force illocutoire* F en plus d'un contenu propositionnel P . Ainsi, les énoncés qui sont utilisés pour accomplir de tels actes illocutoires élémentaires sont composés à la fois d'un marqueur de force illocutoire (assertive, directive, exclamative, force de promesse, etc.) et d'une clause exprimant un contenu propositionnel. Par exemple, comme on peut le vérifier dans les énoncés suivants, ces énoncés expriment le même contenu propositionnel (référence et prédication), et comportent des indicateurs de force illocutoire différents: *Jean sortira de la salle* indique que cette énonciation est une assertion et *Sortez de la salle, Jean!* exprime un ordre. Cette théorie formule des conditions et des règles pour les actes de discours, comme par exemple: lorsque le locuteur accomplit un acte illocutoire du type donner un ordre à quelqu'un, il faut que le locuteur soit en position d'autorité ou de force relativement à l'interlocuteur (ce sont les

conditions préparatoires); pour promettre quelque chose à quelqu'un, il faut que le locuteur se rapporte à une action future, par exemple, *Je te promets que je viendrai* (il s'agit des *conditions sur le contenu propositionnel*). Mais pour que beaucoup d'actes de langage soient réussis et sans défaut, il faut que certaines conditions soient réalisées et que le locuteur présuppose qu'elles sont réalisées; lorsque le locuteur donne un ordre à son interlocuteur, cet ordre a pour condition préparatoire que l'interlocuteur puisse lui obéir; en accomplissant un acte illocutoire, certaines conditions stipulent l'état psychologique exprimé dans cet acte lorsqu'il est sincère; par exemple, en faisant l'assertion *Montréal est une ville du Québec*, le locuteur exprime la croyance que Montréal est une ville du Québec.

Retournons maintenant aux divergences techniques qui existent entre la théorie de la pertinence de Sperber et Wilson (1986, 1989) et l'approche de Grice. Différemment de Grice, la théorie de la pertinence de Sperber et Wilson, traite les processus de compréhension dans la communication en faisant appel à des principes non linguistiques pour en donner une explication cognitive. Ainsi, dans leur approche, le calcul inférentiel réalisé dans l'interprétation du sens communiqué implicitement par un énoncé est régi par un *principe de pertinence* non linguistique selon lequel tout acte de communication ostensive véhicule une garantie de pertinence optimale. C'est

pourquoi le locuteur réalise l'acte le plus pertinent dans un certain contexte. Lorsque la signification intentionnelle communiquée d'une énonciation (ou acte de communication) est pragmatique, le calcul inférentiel comporte non pas seulement des informations linguistiques mais aussi contextuelles. Dans le but de présenter leur modèle de la communication (qu'ils appellent le "modèle inférentiel"), Sperber et Wilson formulent quelques définitions techniques telles que *implication contextuelle*, *effets contextuels*, *pertinence*, *degrés de pertinence*, etc. Ils définissent une *implication contextuelle* comme étant une implication obtenue à partir des informations linguistiques (l'énoncé) et des informations contextuelles. En ce sens, l'implication contextuelle est une implication pragmatique. Elle est, aussi, considérée comme appartenant à la classe des *effets contextuels*. Dans les termes de Sperber et Wilson, il existe trois types d'*effets contextuels*: les implications contextuelles, le *renforcement* d'hypothèses qui nous sont déjà disponibles (il se peut que des informations nouvelles confirment ces hypothèses et les renforcent), et l'*effacement* d'hypothèses préalables (il se peut que des informations nouvelles affaiblissent ces hypothèses, voire les suppriment). À partir de la notion d'effets contextuels, Sperber et Wilson formulent la définition de la pertinence de façon comparative en termes d'effets contextuels et d'efforts cognitifs: plus une hypothèse entraîne d'effets contextuels dans un contexte, plus elle est pertinente dans ce contexte;

et inversement, plus une hypothèse exige d'efforts de traitement dans un contexte, moins elle est pertinente dans ce contexte. Donc, selon Sperber et Wilson, il est plus facile d'être optimalement pertinent que de respecter les maximes de Grice: le degré de coopération exigé dans la communication est plus grand chez Grice qu'il ne l'est dans leur théorie.

Comme nous venons de le voir, Sperber et Wilson soulignent bien les divergences entre leur théorie de la pertinence (communication et cognition) et l'approche de Grice. Contrairement à ce que font Sperber et Wilson, nous aimerions souligner qu'au lieu d'insister sur les divergences entre leurs théories, les linguistes contemporains devraient plutôt chercher des convergences entre leurs théories. Et c'est là notre objectif. En prenant pour base des aspects importants des approches de Grice, Sperber et Wilson, Searle et Vanderveken, etc. nous formulerons notre propre définition de la notion de pertinence.

Cette idée a été inspirée à partir d'un article "À la recherche d'une pragmatique unifiée", Jef Verschueren (1979) où il souligne que la pragmatique est un des domaines les plus confus de la linguistique contemporaine:

... c'est un domaine où on peut construire des théories et mouler à volonté de nouveaux concepts sans rendre compte et même sans prendre connaissance des théories et des concepts qui ont été proposés par les autres linguistes. (Verschueren, 1979, p. 274)

Bien entendu, dans notre thèse nous nous limitons à proposer une bonne définition de pertinence pragmatique qui pourra être prise comme un point de départ dans l'explication des jugements de pertinence que font les interlocuteurs engagés dans une conversation. Puisque notre définition de la pertinence sera traitée du point de vue linguistique (et pragmatique), c'est-à-dire, comme une relation entre des énoncés exprimés dans un contexte donné, il nous faudra prendre en considération les aspects liés à la signification des énoncés ainsi que les aspects illocutoires de la signification et les aspects vériconditionnels des propositions. Pour en tenir compte, nous faisons appel à la théorie des actes de discours, de Searle et Vanderveken, dans laquelle, comme on le sait, de tels aspects sont bien précisés. Ainsi, en utilisant notre définition de la pertinence et en prenant en considération des aspects des théories de la présupposition, des actes de discours, des implications conversationnelles et de la pertinence de Sperber et Wilson, nous tenterons d'illustrer certains jugements de pertinence par des exemples extraits de la vie quotidienne.

Notre thèse est composée de six chapitres. Dans le Chapitre I, nous présentons la théorie de la Pertinence de Sperber et Wilson dont l'idée centrale est, comme nous l'avons déjà dit, de proposer un modèle explicatif

pour la pragmatique linguistique. Selon ces auteurs l'esprit humain cherche toujours à "maximiser" la pertinence. Un individu traite les informations de façon à améliorer au maximum sa représentation globale du monde, ou en d'autres mots, de façon à obtenir un maximum d'effets dans le contexte donné. Toutes choses étant égales, plus une information a d'effets contextuels dans un contexte, plus elle est pertinente dans ce contexte; et toutes choses étant égales, plus l'effort requis pour traiter l'information est grand, moins elle est pertinente dans ce contexte. L'idée de Sperber et Wilson est de généraliser leur modèle communicationnel à toutes les activités mentales et cognitives humaines.

Dans le Chapitre II, nous faisons quelques considérations critiques à la Théorie de la Pertinence de Sperber et Wilson concernant quelques-unes de leurs définitions techniques, plus précisément la définition d'*implication non-triviale* et de *degrés de pertinence* lesquelles nous permettent d'indiquer quelques lacunes existantes. Or, pour Sperber et Wilson, une information est pertinente dans un contexte si et seulement si elle a un effet contextuel dans ce contexte. Il nous semble raisonnable de discuter le problème de l'implication non-triviale car il s'agit d'une notion centrale, laquelle est utilisée dans la définition même d'implication contextuelle; nous rappelons que les implications contextuelles constituent une partie de la classe des effets contextuels, et ceux-

ci sont déterminants pour établir la pertinence d'une hypothèse relativement à un contexte donné. Donc, il s'agit de définitions fondamentales puisqu'elles sont déterminantes pour la propre définition de pertinence formulée par Sperber et Wilson.

Dans le Chapitre III, nous proposons notre définition de la *pertinence* en plus d'une logique, "*la logique de la pertinence P* ", que nous avons élaborée dans le but d'analyser plusieurs jugements de pertinence que nous faisons dans le quotidien.

Au Chapitre IV, nous allons nous concentrer sur les critiques que nous faisons, dans le Chapitre II, à la Théorie de la Pertinence de Sperber et Wilson, pour les analyser à la lumière de notre *logique de la pertinence P* . Nous y montrons que notre définition de la pertinence est plus précise et plus riche que celle de Sperber et Wilson: elle rend compte des lacunes de leur théorie. En outre, nous faisons une analyse de la définition de Sperber et Wilson concernant les *degrés de pertinence* et nous en concluons que notre définition de *degrés de pertinence* (chapitre III), est plus rigoureuse.

Même si Sperber et Wilson disent qu'une analyse de la compréhension des énoncés n'a pas besoin de la théorie des actes de discours dans une pragmatique générale, nous croyons que pour faire une analyse adéquate de la pertinence (pertinence conçue comme une relation entre des

énoncés exprimés par des interlocuteurs dans un contexte donné), toute approche doit y intégrer un traitement des aspects liés à la signification des énoncés. Pour cette raison, dans notre approche, nous avons recouru à la théorie des actes de discours (laquelle sera présentée dans le chapitre V) pour procéder à l'analyse de la pertinence dans le chapitre III.

Cela étant dit, dans le Chapitre V, nous exposons la Théorie des Actes de Discours de John Searle et Daniel Vanderveken dans le but de montrer ses principaux résultats. La T.A.D. intègre la pragmatique de Grice en reformulant et généralisant les maximes conversationnelles de qualité et de quantité proposées par Grice. Comme nous venons de le dire, dans le programme de la T.A.D., il existe une sémantique formelle qui traite de la signification des énoncés; il y a aussi une logique illocutoire qui traite des aspects illocutoires de la signification et une logique propositionnelle qui rend compte des aspects vériconditionnels des propositions.

Finalement au Chapitre VI, nous faisons quelques considérations critiques à la T.A.D. en ce qui concerne la généralisation des maximes conversationnelles de Grice. En utilisant la T.A.D., Vanderveken propose de généraliser et d'expliquer les maximes de qualité et de quantité de Grice, de façon à traiter de tous les types d'énonciations (pas seulement les assertions). Selon Vanderveken, ces deux maximes conversationnelles sont les principales

maximes que les locuteurs respectent dans leurs discours. L'idée de Vanderveken est de généraliser la théorie des actes illocutoires à l'analyse de discours. Mais d'après nous, pour ce faire, il faut enrichir la logique illocutoire actuelle en y intégrant un traitement formel de la *pertinence*. Néanmoins, la T.A.D. ne reformule pas, ne généralise pas la maxime de pertinence de Grice (qui selon nous est la plus importante de toutes les maximes). Évidemment, la maxime de pertinence est fondamentale dans le modèle qui explique les procédures des interlocuteurs pour communiquer une signification autre ou plus riche que celle des énoncés qu'ils utilisent.

Encore dans le Chapitre VI, nous proposons une nouvelle manière d'intégrer un traitement formel de la pertinence à la T.A.D. afin de compléter l'analyse des maximes conversationnelles. Notre suggestion est d'adopter une logique différente dans la logique propositionnelle de Vanderveken, c'est-à-dire, une logique non classique, une logique qui pourrait s'apparenter à notre logique de la pertinence *P* (chapitre III). Nous proposons en outre, d'étendre notre logique *P* à tous les types de forces illocutoires d'actes de discours tels que les assertions, les interrogations, les exclamations, les promesses, etc.

Remarquons que, dans notre thèse, il y a plusieurs usages du terme *pertinence* : il y a la notion intuitive de pertinence de la philosophie du langage et de la linguistique que nous essayons d'expliquer; il y a la notion formelle de

pertinence que nous définissons dans notre logique ainsi que le connecteur d'implication pertinente qui est la constante logique de notre langue formelle.

CHAPITRE I

PRÉSENTATION DE LA THÉORIE DE LA PERTINENCE DE SPERBER ET WILSON

1. Introduction

Les recherches contemporaines sur la communication humaine en général, la communication linguistique en particulier et les fondements de la pragmatique s'inspirent amplement des travaux de Grice pour décrire les normes implicites de la communication verbale et la façon dont elles sont utilisées pour l'interprétation des énoncés.

Le *modèle du code* selon lequel la communication se fait par codage et décodage de messages, gouverne la plupart des théories de la communication (surtout en linguistique structurale et en sémiologie) depuis Aristote. Certains philosophes contemporains, comme Grice, ont proposé un autre modèle applicable à la description de la communication, soit le *modèle inférentiel*, selon lequel communiquer, c'est produire intentionnellement et

interpréter des indices.

1- Le modèle du code, le modèle inférentiel et la communication

Dans leur analyse du modèle du code, Sperber et Wilson (1986, et 1989 pour la traduction française) le définissent comme un système qui associe des messages à des signaux et qui explique comment les symboles sont émis, transmis et interprétés pour rendre la communication possible; une relation de symétrie entre les processus de décodage (relation décodeur-destination) et d'encodage (relation source-codeur) est à la base de ce modèle.

Le modèle du code et le modèle inférentiel sont au centre d'une discussion entre philosophes du langage, théoriciens de la littérature, linguistes, psychologues et anthropologues, dont la principale question est de savoir lequel des deux modèles il faut considérer pour fonder une théorie de la communication. Discussion inutile peut-être puisqu'on ne peut ignorer que, dans le processus de communication, en général, l'un et l'autre modèles y sont appliqués pour rendre compte de modes distincts de communication: la communication peut être réalisée en codant et en décodant des messages aussi bien qu'en produisant et en interprétant des indices. Mais, comme nous venons

de le dire, quelques auteurs cherchent à déterminer lequel des deux modèles est adéquat pour bien fonder une théorie générale de la communication. Pour Grice, le codage/décodage n'est ni nécessaire ni suffisant pour expliquer la communication verbale; pour comprendre le fonctionnement de la communication et expliquer son succès, il faut recourir à un autre modèle basé sur *la reconnaissance des intentions du locuteur* qui justement communique en rendant manifeste certains indices qui révèlent en partie ses intentions. Dans le modèle inférentiel, il est clair que la Grammaire ne permet de rendre compte que d'une partie seulement de la signification *communiquée*.

En fait, qu'est-ce qui définit les langues naturelles? Premièrement, les règles de syntaxe qui régissent la combinaison des signaux dans les séquences de symboles (des unités linguistiques en phrases) et, deuxièmement, les règles sémantiques qui concernent l'interprétation des séquences de symboles et qui rendent possible l'association des messages aux séquences de signaux (unités linguistiques et leurs combinaisons). Alors, même si des sémioticiens, par exemple, justifient les grammaires génératives comme les meilleurs modèles des langues humaines et insistent pour dire que c'est le modèle du code qui permet d'expliquer la possibilité de la communication - bien entendu, la communication verbale, puisque la grammaire générative est elle-même un code, c'est-à-dire, un système d'appairement <représentations

phonétiques, représentations sémantiques > - Grice, et plus récemment Sperber et Wilson s'opposent à cela en disant que la représentation sémantique qu'une grammaire générative assigne à une phrase ne coïncide pas avec les pensées qui peuvent être communiquées par l'énonciation de cette phrase; elle ne tient pas compte des éléments du contexte qui permettent de déterminer l'acte illocutoire exprimé par l'énonciation littérale de cet énoncé dans un contexte d'énonciation. Selon ces auteurs, la représentation sémantique d'une phrase ne fournit pas à elle seule l'interprétation complète d'une énonciation de cette phrase car elle peut être utilisée en diverses circonstances soit pour exprimer la même pensée, soit pour manifester notre propre attitude vis-à-vis de la pensée exprimée (attitudes propositionnelles), soit pour exprimer un nombre illimité de pensées à cause de facteurs comme des indéterminations référentielles, des ambiguïtés et incomplétudes sémantiques. Dans les exemples suivants l'interprétation des énoncés dépend d'autres aspects qui ne sont pas grammaticaux:

(1) Hier il faisait froid.

Si Pierre énonce (1) le dimanche, (1) exprime la pensée qu'il faisait froid samedi. Si Marie énonce (1) le mercredi, (1) exprime la pensée qu'il

faisait froid mardi.

Alors, en prenant une grammaire générative comme base, on ne peut pas déterminer la référence de “hier”, par exemple, car une telle grammaire ne tient compte que des propriétés linguistiques d’un énoncé et dans ce cas, elle ne détermine pas à quel moment a lieu l’énonciation. Elle ne détermine pas non plus la référence de “je” dans l’exemple suivant:

(2) Je suis un professeur

Dans ce cas, la grammaire générative nous dit que “je” est le locuteur, mais pas qui est le locuteur (ce qui ne peut être déterminé qu’à partir de ressources extra-linguistiques).

En bref, selon l’argumentation de Sperber et Wilson, on ne passe pas de la représentation sémantique à la pensée communiquée au moyen des énoncés par un accroissement de codage, car dans la communication verbale la signification linguistique d’une phrase énoncée n’encode pas entièrement ce que veut dire le communicateur; la signification permet au destinataire d’inférer ce *vouloir-dire* et le processus de décodage lui sert comme un indice des intentions du locuteur, c’est-à-dire, ce qui permet de passer de la représentation sémantique à la pensée communiquée est un processus

inférentiel de reconnaissance des intentions du locuteur (du type décrit par Grice).

Remarquons que ce qui distingue les processus inférentiels des processus de décodage, c'est que les premiers ont pour *input* un ensemble de prémisses et pour *output* un ensemble de conclusions - qui sont logiquement impliquées par les prémisses ou justifiées par celles-ci - alors que les processus de décodage ont pour *input* un signal et pour *output* la reconstitution du message associé au signal par un code donné.

C'est dans son article intitulé *Meaning* que Grice (1957) définit ce qu'il appelle la "signification *non-naturelle*" à partir des intentions d'un communicateur⁴. Selon cette définition, le communicateur *C* provoque intentionnellement l'effet *E* chez l'auditeur *A* du fait que cet auditeur reconnaît l'intention du communicateur.

(3) "*C* signifie *non-naturellement* quelque chose par l'énonciation d'un énoncé⁵ *x*" équivaut (en gros) à:

⁴ Strawson (1971, p. 155) et par la suite Schiffer (1972, p. 11) ont attiré l'attention sur le fait qu'il y a là trois sous-intentions.

⁵ Grice considère comme "énoncé" non seulement les énoncés linguistiques et signaux codés mais aussi toute forme de comportement du locuteur pour communiquer quelque chose à son destinataire.

- (i) C a l'intention que l'énonciation de x produise un certain effet E sur un auditoire A .
- (ii) C a l'intention que A reconnaisse l'intention (i).
- (iii) C a l'intention que la reconnaissance par A de l'intention (i) soit au moins en partie la raison pour laquelle l'énonciation de x produise l'effet E chez A .

Appelons (i) “intention informative” ou “primaire”; (ii) “intention communicative” ou “secondaire” et (iii) “intention méta-informative”. Alors, chez Grice, la communication ainsi décrite en termes d'intentions et d'inférences est réalisée par la manifestation et la reconnaissance d'intentions.

En s'inspirant d'un exemple de Sperber et Wilson (1986 et 1989, chapitre I), voyons comment le modèle gricéen peut être appliqué à la description de la communication en termes d'intentions et d'inférences.

Supposons que Pierre est à la cuisine en train de préparer son dîner et que Marie vient d'arriver chez lui. Supposons aussi que Marie ait l'intention de l'informer du fait qu'il neige au dehors dans un contexte où Pierre ne peut savoir le temps qu'il fait (les volets étant fermés). Elle passe devant lui et son manteau entièrement couvert de neige est un indice patent pour convaincre Pierre qu'il neige au dehors. Dans ce cas, Pierre pourrait s'apercevoir de ce fait sans avoir conscience de l'intention de Marie. Cette façon d'informer en fournissant des indices directs de l'information qu'on veut transmettre n'est

pas, selon l'analyse de Grice, considérée comme une forme de communication puisqu'elle n'est pas intentionnelle (elle pourrait l'être; si oui, alors il y a communication). Supposons maintenant que Marie avant de rencontrer Pierre à la cuisine enlève son manteau et le range dans l'armoire. Elle passe devant lui, mais cette fois-ci elle ne produit pas un indice *direct* du fait qu'il neige au dehors. Cependant Marie peut donner à Pierre un indice direct de son intention de l'en informer en faisant, par exemple, l'énonciation (laquelle servira d'indice *indirect* du fait qu'elle veut l'en informer):

(4) "Il neige au dehors".

Supposons que Marie ait précisément l'intention d'informer Pierre de cette façon et que celui-ci croit ce qu'elle lui dit. Cette façon d'informer est selon Grice, une forme de communication puisque:

- (i) Marie a l'intention que l'énonciation de (4) produise un certain effet (la conviction qu'il neige au dehors) sur Pierre.
- (ii) Marie a l'intention que Pierre reconnaisse l'intention (i).
- (iii) Marie a l'intention que la reconnaissance par Pierre de l'intention (i) soit au moins en partie la raison pour laquelle l'énonciation de (4) produise cette conviction chez Pierre.

Ainsi, pour que la communication réussisse, il suffit, d'après Grice, qu'il y ait un moyen, n'importe lequel, de reconnaître les intentions informatives du communicateur; l'auditeur ne doit pas simplement reconnaître le sens linguistique d'une phrase énoncée mais en inférer ce que veut dire le communicateur. Celui-ci doit transmettre une information de façon intentionnelle et manifester ouvertement à son destinataire le caractère intentionnel de sa transmission d'information; à cette analyse, Grice a ajouté la condition (iii) qu'il faut que la transmission de l'information primaire soit subordonnée à la transmission de l'information secondaire, ce qui veut dire que (i) ne doit pas être récupérable intégralement sans passer par (ii); par définition, (iii) n'est pas réalisable si l'intention (i) ne se réalise pas.

Mais comment fait l'auditeur pour reconnaître les intentions informatives du locuteur? Sperber et Wilson ont raison de considérer qu'une analyse descriptive de la communication est nécessaire mais qu'il faudrait tout d'abord expliquer son succès, c'est-à-dire, identifier les mécanismes sous-jacents dans la psychologie humaine que les interlocuteurs mettent en oeuvre afin de communiquer. Alors, préoccupés de trouver une explication concernant un tel mécanisme, Sperber et Wilson proposent une analyse de la communication basée sur un modèle de la communication ostensive qui a été

inspiré de Grice (1975; 1979). Il s'agit de son article "Logique et conversation" où Grice introduit la notion d'implication (*implicature*) - ce qui au plan théorique a été considéré comme une grande contribution à la pragmatique. Il a montré comment certaines énonciations d'énoncés communiquent plus que ce que les mots constituants de la phrase ne signifient, c'est-à-dire, plus que ce qui est dit littéralement; il existe donc une divergence entre la signification de la phrase et le sens transmis par l'énonciation - qui ne dépend pas des conditions de vérité de la phrase entière - ce sens Grice l'appelle une *implication (implicature)*. Nous reviendrons sur ce sujet dans le chapitre V où nous présenterons la Théorie des Actes de Discours (T.A.D.), de John Searle et Daniel Vanderveken, qui traite de cette question.

Comme nous avons déjà mentionné dans l'introduction de notre thèse, au plan de la communication, Grice a soutenu l'idée que l'usage du langage est régi par le principe général (tacitement accepté par les interlocuteurs) qu'il a appelé *Principe de Coopération*:

(5) que votre contribution conversationnelle corresponde à ce qui est exigé de vous, au stade atteint par celle-ci, par le but et la direction acceptés de l'échange parlé dans lequel vous êtes engagé. (1979, p. 61)

Quatre types de maximes conversationnelles sont subordonnées au

Principe de Coopération:

La maxime de quantité détermine que la contribution du locuteur ne doit être ni plus ni moins informative que ce qui est exigé.

La maxime de qualité (de véridicité) détermine que le locuteur doit dire la vérité et être sincère.

La maxime de relation détermine que le locuteur doit être pertinent.

La maxime de manière détermine que le locuteur doit être bref, clair, ordonné et éviter les ambiguïtés.

L'exemple du dialogue suivant montre comment le respect du principe de coopération (et des quatre maximes conversationnelles), ainsi que l'observation du comportement du locuteur et du contexte, permettent à l'auditeur d'inférer l'intention informative du locuteur:

Pierre: Pars-tu en vacances ce matin?

Marie: J'ai encore du travail à l'université aujourd'hui.

Comme on peut le voir, le contenu explicite de l'énoncé de Marie ne répond pas directement à la question de Pierre et dans ce cas, cet énoncé est un usage de la maxime de qualité. Si l'on retient l'hypothèse que Pierre

suppose que Marie dit le vrai, que le fait de l'arrière-plan existe (quelqu'un qui a du travail à l'université toute la journée, ne partira pas en vacances cette journée) et que Marie a l'intention de lui faire inférer de sa réponse qu'elle ne peut partir aujourd'hui puisqu'elle doit rester à l'université toute la journée, cela permettra à Pierre d'inférer dès lors, que Marie ne part pas en vacances ce matin.

Ces hypothèses qui sont inférées contextuellement, ou plus précisément, ces implicatures, sont communiquées par le fait que Marie a l'intention de les transmettre.

Ainsi, en s'inspirant des idées de Grice, Sperber et Wilson proposent de développer l'analyse de la communication inférentielle, suggérée par Grice (1957) en un modèle explicatif. Selon eux, de telles idées peuvent être prises comme point de départ dans leur analyse dont le but est d'enrichir la définition et l'explication que donne Grice de la communication.

Étant donné l'importance et l'originalité du programme gricéen pour un modèle inférentiel de la communication, quelques auteurs, parmi lesquels Strawson (1964, 1971), Searle (1969, 1979), Schiffer (1972), Grice lui-même (1975, 1978), Bach et Harnish (1979), Ducrot (1972, 1980a, b), Leech (1983), Sperber et Wilson (1986) et Vanderveken (1991, 1997a), ont suggéré des reformulations dans le but d'enrichir ce programme en ce qui

concerne la définition de la “signification” et la définition de la communication qui, selon eux, souffraient de quelques défauts. Malgré l’effort de ces auteurs pour décrire les processus de la communication, leurs suggestions suscitent encore des objections qui méritent d’être présentées.

À cet égard, Sperber et Wilson ont essayé de montrer comment décrire la communication comme un processus inférentiel de reconnaissance des intentions du locuteur. Nous reviendrons sur ce sujet dans le développement de notre thèse.

Searle (1969) affirme que l’analyse de la signification de Grice est purement inférentielle et qu’elle ne fait pas ressortir le rapport entre le vouloir-dire du locuteur et la signification linguistique. L’auteur propose donc de combiner la communication codée et la communication inférentielle en soulignant que les individus ont en commun une langue (ou code) qui leur permet, en général, de produire des indices de leurs intentions. Searle attire notre attention sur le fait que dans la définition de la signification non naturelle, Grice pensait principalement à des effets perlocutoires. Mais, comme Searle lui-même le remarque, dire quelque chose et vouloir signifier ou communiquer ce que l’on dit c’est accomplir un acte illocutoire en visant un effet illocutoire: la compréhension de ce qu’a dit le communicateur. Ici, l’analyse searlienne comprend juste la signification “littérale”. Elle exige que le communicateur ait

l'intention que le destinataire le comprenne en vertu de sa connaissance des règles auxquelles la phrase énoncée est subordonnée. Searle reformule la notion gricéenne de signification non naturelle de la façon suivante:

Dire que le communicateur *C* fait une énonciation *E* d'une phrase *P* avec l'intention de signifier (littéralement) ce qu'il dit, c'est dire que:

- (i) Par l'énonciation *E* de *P*, *C* a l'intention *i-l* de faire connaître / reconnaître à un auditoire *A* que la situation spécifiée par les règles de *P* est réalisée. (Appelons cet effet, l'effet illocutoire *EI*)
- (ii) Par l'énonciation *E* de *P*, *C* a l'intention de produire *E* par la reconnaissance de *i-l*.
- (iii) *C* a l'intention que la reconnaissance par *A* de l'intention *i-l* se fasse en vertu de la connaissance qu'a *A* des règles qui gouvernent les éléments de la phrase *P*.

Dès lors, le caractère intentionnel de la signification du locuteur n'est pas suffisant pour la communication humaine. Cette reformulation tient compte du fait que la communication humaine repose pour une large part sur l'utilisation de codes.

Il est clair que la communication sans code est possible: je demande à Pierre s'il pense voyager. Son comportement se limite à me montrer le billet d'avion qu'il vient d'acheter. Ça me permet de reconnaître son

intention de m'en informer - ce qui est suffisant pour la réussite de la communication. Mais, comme Searle le soutient, ça ne fonctionne pas en règle générale: malgré le fait que certains de nos actes illocutoires soient effectués sans faire appel à des conventions, comme le montre l'exemple ci-dessus, cependant, "...si l'on ne dispose pas d'un langage il est impossible de demander à quelqu'un d'entreprendre des recherches sur le diagnostic et le traitement de la mononucléose infectieuse chez les étudiants américains" (Searle, 1969, 38).

Alors, qu'est-ce que c'est que communiquer? Est-ce qu'on peut définir un acte de communication en termes de réalisation de toutes les "intentions gricéennes"? Comme le souligne Searle, un locuteur peut vouloir dire quelque chose et arriver à le communiquer sans passer par toutes ces intentions. En énonçant (4) "Il neige au dehors", Marie veut que son énonciation produise sur Pierre la conviction qu'il neige au dehors. Mais il peut arriver que Pierre reconnaisse cette intention (primaire) de Marie sans croire ce qu'elle dit. Dans ce cas, Marie réussit à communiquer son message à Pierre par la réalisation d'une seule intention, l'intention secondaire. Chez Pierre, l'intention primaire n'a pas été satisfaite, mais on ne peut pas nier que Marie n'ait pas telle intention; Marie a l'intention primaire, bien que Pierre ne la reconnaisse pas. Une fois admise cette possibilité, on voit que l'analyse de Grice montre seulement que la reconnaissance d'une intention informative peut

amener à sa réalisation.

Selon Sperber et Wilson (1989), la révision de Searle revient à transformer l'analyse de Grice en une modification triviale du modèle du code. Ils affirment qu'il ne serait pas possible de combiner la théorie inférentielle forte de la communication avec une théorie du code modifiée sans que cette réconciliation ne combine les défauts de chacune de ces théories: elle ne tiendrait pas compte du rôle de l'inférence non codée dans la communication; qui plus est, l'aspect non inférentiel d'une partie du décodage serait négligé. Selon ces auteurs, des formes complexes de communication peuvent reposer sur une combinaison des deux modes de communication, le mode inférentiel et le mode codage-décodage; un processus inférentiel peut être utilisé à l'intérieur d'un processus de décodage; il y a des cas où la communication inférentielle utilise des signaux codés, mais les intentions du communicateur ne sont pas codées entièrement par de tels signaux. Ils soutiennent, malgré que le modèle du code ait un fort pouvoir explicatif de la bonne communication (dont l'exigence est le partage d'un code commun), que cette explication ne peut pas être retenue car la description de la communication sur laquelle elle se fonde est basée sur l'hypothèse que les humains communiquent en codant et en décodant des pensées. Par conséquent, elle n'explique pas comment fait l'interlocuteur pour recouvrer l'intention informative du locuteur.

Différemment de Grice, Sperber et Wilson (1989) soutiennent qu'il existe une continuité entre les cas où le locuteur parvient à informer l'allocataire de quelque chose en l'informant de son intention (ce qui est rendu plus simple par la transmission de son intention primaire) et les cas où la transmission d'une information ne met pas en jeu la reconnaissance par l'auditeur de l'intention informative du communicateur. On a déjà fait remarquer que chez Grice la transmission de l'information primaire doit être subordonnée à la transmission de l'information secondaire. Selon Sperber et Wilson, l'intention communicative est une intention informative de deuxième ordre. Le résultat d'un acte de communication est la réalisation simultanée de ces deux types d'intention étant donné que l'information secondaire facilite l'accès à l'information primaire. Ces auteurs suggèrent la modification de l'analyse gricéenne en proposant la réformulation des notions d'intention informative et d'intention communicative et en se bornant à ces deux types d'intention. D'après eux, cela permet de résoudre les problèmes de régression à l'infini que l'analyse de Grice suscite. Avant de présenter cette réformulation, voyons d'abord un de leurs exemples qui illustre bien ce problème:

Supposons, par exemple, que Marie souhaite que Pierre répare son sèche-

cheveux, mais qu'elle ne souhaite pas le lui demander directement. Elle entreprend alors de démonter son sèche-cheveux et en étale les pièces comme si elle était en train de le réparer. Elle ne pense pas que Pierre sera dupe de cette mise en scène; en fait, s'il croyait vraiment qu'elle était en train de réparer toute seule son sèche-cheveux, il n'interviendrait probablement pas. Marie espère que Pierre sera assez perspicace pour comprendre qu'il s'agit d'une mise en scène dont le but est de l'informer du fait qu'elle a besoin de son aide pour réparer son sèche-cheveux. Elle espère cependant qu'il ne sera tout de même pas perspicace au point de comprendre que le raisonnement qu'il se tient est précisément celui qu'elle voulait qu'il se tienne. Puisque ainsi Marie n'a à proprement parler rien demandé à Pierre, s'il ne l'aide pas, elle n'aura pas essuyé un refus. (Sperber et Wilson, 1989, p. 52)

On ne peut nier que dans cet exemple Marie possède une intention informative ainsi qu'une intention communicative (qu'elle veut cacher à Pierre). Sperber et Wilson attirent notre attention sur le fait que nous hésitons à dire, dans ce cas, que Marie a communiqué avec Pierre puisqu'elle dissimule à Pierre son intention de second ordre. À cet égard, Sperber et Wilson remarquent que la solution présentée par Strawson (1964) et développée par Schiffer (1972) contient un problème en ce qui a trait à la plausibilité psychologique - ou en d'autres mots, elle n'est pas psychologiquement plausible - soit le problème de la régression à l'infini. Or, comment préciser l'idée que les intentions communicatives doivent être entièrement publiques (*overts*), c'est-à-dire, entièrement accessibles, comme le suggéraient Strawson et Schiffer, puisqu'il s'agit d'une notion assez vague? Strawson suggérerait de considérer une intention (*overt*) comme étant la première dans une série d'intentions; dans cette série chaque intention précédente doit être reconnue.

La suggestion de Schiffer (1972) était de considérer une intention publique (*overt*), l'intention qui "fait partie du savoir mutuel" (ce que Lewis (1969) appelait *savoir commun*), c'est-à-dire, quand elle est mutuellement connue des interlocuteurs. Strawson soutenait que la solution serait d'ajouter aux intentions informative et communicative, une intention de troisième ordre, à savoir, l'intention méta-communicative (l'intention qui permet au destinataire de reconnaître l'intention de deuxième ordre). Mais cela ne suffit pas. Car il est possible d'imaginer un autre exemple où le communicateur veuille dissimuler son intention méta-communicative. Si on ajoute une autre intention méta-métacomcommunicative de quatrième ordre, on pourrait imaginer encore un autre exemple ... et ainsi de suite à l'infini. Comme le soulignent Sperber et Wilson, cette solution conduit encore à postuler une infinité d'intentions et conséquemment, le savoir mutuel en forme de régression à l'infini ne peut pas être intégré dans une explication psychologiquement plausible. Puisque d'un point de vue psychologique une intention est une représentation mentale réalisable sous forme d'action, les psychologues ne s'intéresseraient pas à l'analyse d'un énoncé impliquant la réalisation d'une infinité d'intentions au sens que nous venons de définir. Comme nous le verrons plus loin, pour solutionner ce problème Sperber et Wilson proposent de repenser les notions d'intention informative et d'intention communicative. Ils se débarrassent du

réquisit du savoir mutuel en introduisant le concept de *manifesteté mutuelle* (pour remplacer la notion vague d'intention *overt*) en faisant ressortir que l'intention informative doit être qualifiée en plus de manifeste, de *mutuellement manifeste*. Ils donnent la définition suivante d'un fait manifeste:

- (7) Un fait est *manifeste* à un individu à un moment donné si et seulement si cet individu est capable à ce moment-là de représenter mentalement ce fait et d'accepter sa représentation comme étant vraie ou probablement vraie. (p. 65)

Conséquemment, l'environnement cognitif d'un certain individu dans une certaine situation comprend non pas seulement les faits qui sont représentés par des propositions qu'il peut inférer à partir des informations qui lui sont présentes mais aussi les faits qu'il peut percevoir (dans le sens qu'il est capable de percevoir et qu'il a l'opportunité de percevoir: un handicapé par exemple, qui n'est pas capable de tourner sa tête, a besoin de quelqu'un pour tourner sa chaise roulante afin qu'il puisse percevoir ce qui est derrière lui). Mais le fait que je sois *capable* de percevoir un fait dans mon environnement cognitif n'implique pas que je le perçoive effectivement; en d'autres mots, une chose peut être manifeste sans *uptake* (appréhension effective du stimulus), c'est-à-dire, sans que je l'appréhende effectivement. Ainsi, la définition de manifesteté n'implique pas qu'il y ait toujours uptake, car une telle implication

serait très forte. Pour cette raison, Sperber et Wilson ont affaibli la définition de manifesteté en utilisant une notion de capacité (ici il s'agit donc d'une définition modale). Mais bien sûr, l'uptake entre en jeu toujours dans les actes bien réussis.

En ce qui concerne la notion de *connaissance mutuellement manifeste* celle-ci se définit à partir des notions de connaissance manifeste et de la notion d'*environnement cognitif mutuel*:

(8) Un *environnement cognitif* d'un individu est un ensemble de faits qui lui sont manifestes. (p. 65)

Sperber et Wilson appellent

(9) *environnement cognitif mutuel* tout environnement cognitif partagé dans lequel est manifeste l'identité des individus qui partagent cet environnement. (p. 70)

À partir de là, dans un environnement cognitif mutuel, toute hypothèse manifeste pour les individus est *mutuellement manifeste*.

Revenons maintenant à la reformulation de la notion d'intention informative par Sperber et Wilson:

Un communicateur produit un stimulus⁶ avec:

(10) l'*intention informative* de rendre manifeste ou plus manifeste à l'auditoire un ensemble d'hypothèses I. (p. 93)

Par "intention", Sperber et Wilson entendent un état psychologique dont le contenu doit être mentalement représenté; par hypothèses, ils entendent "des pensées que l'individu traite comme des représentations du monde réel (par opposition à des fictions, des désirs ou des représentations de représentations)" (p. 12). Ils soulignent que le communicateur doit avoir en tête une représentation de l'ensemble d'hypothèses I; il n'est pas nécessaire qu'il y ait une représentation de chaque hypothèse dans cet ensemble, mais seulement une description qui identifie l'ensemble.

Sperber et Wilson redéfinissent donc la notion d'intention communicative: communiquer par ostension, c'est produire un stimulus pour parvenir à réaliser une intention informative avec

(11) l'*intention communicative* de rendre mutuellement manifeste au destinataire et au communicateur que le communicateur a cette intention informative. (p. 97)

⁶ Terme emprunté aux psychologues qui l'utilisent pour se référer aux modifications perceptibles de l'environnement.

Ainsi, l'exemple de Marie étalant les pièces de son sèche-cheveux devant Pierre afin qu'il le lui répare, montre bien qu'elle désirait que son intention informative soit manifeste à Pierre, mais non pas *overt* (ou, dans la terminologie de Sperber et Wilson, non pas mutuellement manifeste). Dans ce cas, on a l'intuition que Marie n'a pas vraiment communiqué avec Pierre. Ce que Sperber et Wilson soutiennent est que leur redéfinition de l'intention communicative rend compte de cette intuition. Ils nous offrent, donc, une définition de la *communication ostensive-inférentielle*:

(12) Le communicateur produit un stimulus qui rend mutuellement manifeste au communicateur et au destinataire que le communicateur veut, au moyen de ce stimulus, rendre manifeste ou plus manifeste au destinataire un ensemble d'hypothèses I. (p. 101)

À cette reformulation de Sperber et Wilson, Marc Dominicy (1991) adresse quelques objections. Premièrement, en ce qui concerne la définition (9) d'environnement cognitif mutuel, il remarque que la notion d'identité utilisée par Sperber et Wilson échappe à un traitement extensionnel car "...il s'agit d'une identité de rôle applicable, par exemple, au lecteur potentiel d'un livre ou à l'auditeur potentiel d'une émission radiophonique" (p. 86). Dans son exposé, Dominicy souligne, en ce qui concerne la définition de manifesteté mutuelle, son option de fonder sa critique sur une propriété logique d'un ensemble

d'hypothèses E et non pas sur son statut d'environnement cognitif mutuel:

une hypothèse est mutuellement manifeste à deux individus A et B si et seulement si, elle appartient à un ensemble d'hypothèses E tel que, si H appartient à E , alors l'hypothèse que H est (à un certain degré) manifeste à A et à B appartient également à E .(p. 86)

Pour lui, la stipulation qu'une hypothèse H est mutuellement manifeste "équivalait à poser l'existence d'un ensemble infini d'hypothèses dont le degré de manifesteté 'tend asymptotiquement vers zéro' " (p. 87). Il nous rappelle l'anecdote du sèche-cheveux pour illustrer ce qui est nécessaire pour que Marie puisse vraiment communiquer à Pierre son désir qu'il répare son sèche-cheveux: Marie devrait avoir l'intention que son intention informative devienne mutuellement manifeste à elle-même et à Pierre. Alors, elle devrait avoir eu l'intention qu'il soit manifeste à Pierre qu'il est manifeste à Marie qu'il est manifeste à Pierre que Marie a l'intention de l'en informer.

L'objection de Dominicy s'adresse au contenu même de l'intention communicative. En supposant que l'intention communicative de C représente un "état de choses désirable" comme l'illustrent Sperber et Wilson (1989, p. 347) et en considérant que, dans cet état de choses, l'intention informative de C est mutuellement manifeste à C et à D , l'hypothèse que C a cette intention informative fait partie d'un ensemble infini d'hypothèses E (comme

celui décrit plus haut). Sa critique se résume à demander si l'intention communicative de C représente la structure logique de l'ensemble d'hypothèses E et les degrés décroissants de manifesteté qui leur sont assignés. Car, comme lui-même le remarque, dans l'approche psychologique, une intention est une représentation mentale réalisable sous forme d'action et conséquemment, il n'est pas clair qu'une intention se réaliserait sous forme d'action:

Il ne servirait à rien de transformer le degré de manifesteté de chaque hypothèse en un degré de force intentionnelle (c'est-à-dire de réinterpréter ' C a l'intention que $H1$ soit manifeste à C et D au degré $d1$, et que $H2$ soit manifeste à C et D au degré $d2$ ' en ' C a l'intention de force $f1$ que $H1$ soit manifeste à C et D , et l'intention de force $f2$ que $H2$ soit manifeste à C et D '); car cela nous amènerait à introduire une infinité d'intentions. En d'autres termes, la notion d'intention communicative ne semble pas devoir échapper à une intensionnalité irréductible qui compromet évidemment sa crédibilité psychologique. (p. 87)

Dominicy attire encore notre attention sur l'observation que font Sperber et Wilson relativement à la définition (12) de la communication ostensive-inférentielle. Selon ces auteurs, cette définition admet que la communication puisse être aussi non intentionnelle; par exemple, lorsqu'un stimulus visant seulement à informer vient, pour une raison quelconque, rendre mutuellement manifeste l'intention d'informer, nous sommes quand même devant un cas de communication. Ce passage est bien illustré par l'exemple

que nous offrent Sperber et Wilson: Supposons que Marie veuille informer Pierre qu'elle est fatiguée. Elle se prend à bâiller de façon à ce que son bâillement paraisse naturel. Mais Pierre s'aperçoit que son bâillement est volontaire et ainsi, l'intention informative de Marie devient mutuellement manifeste. L'objection de Dominicy à cet égard est bien évidente. Elle repose sur le fait que l'intention informative de Marie soit manifeste à Pierre et à Marie ne suffit pas à rendre cette intention mutuellement manifeste à Pierre et Marie. En effet, il est possible qu'il ne soit pas manifeste à Marie que son intention informative est manifeste à Pierre; "...rien ne garantit, dans un tel exemple, que l'itération récursive du préfixe 'Il est manifeste à Pierre et à Marie que' ne nous donnera pas, à un stade aussi lointain que ce soit, une hypothèse non manifeste à Pierre ou à Marie" (p. 87).

Ainsi, conclut-il, la théorie sur laquelle se fondent Sperber et Wilson en replaçant la manifesteté mutuelle dans le champ des faits qui peuvent réaliser une intention réelle, se heurte à la même catégorie de contre-exemples que l'analyse de Grice suscite.

2- Le modèle inférentiel de Sperber et Wilson

Dans la section antérieure nous avons introduit quelques définitions techniques de la théorie de la pertinence de Sperber et Wilson concernant le modèle de la communication ostensive-inférentielle, lequel a été inspiré des idées de Grice. Dans cette première étape, Sperber et Wilson se concentrent plus sur la nature ostensive du comportement du communicateur. Pour ce qui est plus spécifiquement du caractère inférentiel du processus de compréhension, nous allons présenter dans cette section le modèle proposé par Sperber et Wilson des principales capacités inférentielles à l'oeuvre dans la compréhension verbale.

Comme on l'a vu plus haut, dans le processus de compréhension verbale, le destinataire ne peut décoder l'intention informative du communicateur ni la déduire non plus; il peut interpréter un comportement ostensif en préservant et confirmant une hypothèse sur l'intention communicative du locuteur; plus précisément, l'hypothèse que le destinataire forme est fournie à dessein par des indices ostensifs du communicateur; en fait, c'est le communicateur qui produit des stimuli afin de rendre la compréhension possible. Une telle hypothèse retenue par le destinataire peut être confirmée mais non pas démontrée.

Le modèle esquissé par Sperber et Wilson concernant le processus de compréhension inférentielle repose sur un seul type de processus inférentiel, à savoir, l'inférence non démonstrative, ou plus précisément, l'inférence non-démonstrative *spontanée* que les humains sont capables de maîtriser en utilisant plusieurs techniques: l'ensemble des indices fournis par le communicateur n'équivaut pas à la démonstration de son intention informative.

Soit dit en passant, l'inférence démonstrative est un processus par lequel on applique des *règles* déductives à un ensemble d'hypothèses de départ. Les raisonnements déductifs permettent de traiter une hypothèse, soit d'un point de vue syntaxique (ou computationnel selon Sperber et Wilson) soit d'un point de vue sémantique. Syntaxiquement, dans un système déductif une hypothèse Q est une *conséquence logique* de l'hypothèse P si, et seulement si, les règles déductives de ce système permettent de déduire Q à partir de P . Cette relation syntaxique ne tient compte que des propriétés formelles des hypothèses, c'est-à-dire, leurs propriétés sémantiques ne sont pas prises en considération dans le processus computationnel. Sémantiquement, Q est une *conséquence nécessaire* de P (*entailment*, en anglais) si, et seulement si, il n'existe aucun état de choses qui rende P vraie sans rendre Q également vraie. La relation de conséquence nécessaire est une relation sémantique puisqu'elle est définie à partir des états de choses qui constituent

l'interprétation sémantique des hypothèses particulières: leur interprétation sémantique dépend de tels états de choses. La fonction des règles d'inférence consiste à garantir la *validité* logique des inférences qu'elles gouvernent. Dans l'inférence démonstrative on applique des règles déductives à un ensemble de prémisses prises au départ. De telles règles déductives s'appliquent à des hypothèses en vertu de leurs formes logiques. Dans le cas d'une inférence déductive valide, la vérité des prémisses garantit la vérité des conclusions. Dans l'inférence inductive il est probable que si les prémisses sont vraies, les conclusions le sont également. Étant donné une prémisses P , les règles d'inférence non démonstratives ne garantissent pas qu'une certaine conclusion Q soit obligatoirement inférée; qui plus est, elles ne rendent pas possible l'*engendrement* de toutes les conclusions qui sont d'intérêt pour confirmer un certain ensemble de prémisses.

Comme le soulignent Sperber et Wilson, de telles règles d'inférence non démonstratives permettent d'arriver à des conclusions non démonstratives valides qui, du "point de vue logique" mettent en oeuvre deux processus distincts: le processus représentationnel, à savoir la représentation d'un état de choses, ou la formation d'hypothèses, et le processus computationnel, soit la représentation de la probabilité ou valeur de confirmation de ces hypothèses. La formation d'hypothèses est liée à

l'imagination créatrice alors que le processus computationnel considéré, comme un processus purement logique, est régi par des règles d'inférence. L'un et l'autre processus sont vus comme des aspects de l'interprétation qui combinent, d'une part, la constitution du contexte⁷ et, d'autre part, le processus inférentiel. L'inférence est dans son tout une forme de fixation de croyances dans le sens où une hypothèse est admise comme vraie ou probablement vraie sur la base d'autres hypothèses admises au départ comme vraies ou probablement vraies. De même, la perception est une autre forme de fixation de croyances: une hypothèse est acceptée comme vraie ou probablement vraie sur la base de données cognitives non conceptuelles.

À ce sujet, Sperber et Wilson remarquent que quelques auteurs comme Leech, Levinson, Bach et Harnish, Brow and Yule, de Beaugrande et Dressler, soutiennent qu'une inférence non démonstrative ne peut pas comporter une déduction comme une de leurs sous-parties. Le cas plus simple d'inférence non démonstrative sont les implications, lesquelles selon Leech (1983) sont identifiées par un processus de stratégie rationnelle informelle de résolution de problèmes, et qui, selon Levinson (1983), ne peuvent être directement modélisées au moyen d'une relation sémantique telle que

⁷ Un tel contexte n'est pas donné par la situation elle-même, mais comme nous le verrons dans la section 3, il dépend largement de la capacité cognitive du destinataire à construire un

l'implication. De même, Bach et Harnish (1979) soutiennent que les implications sont identifiées non pas par une inférence déductive mais par une inférence qui vise une explication plausible; même s'il manque encore une explication satisfaisante de la nature du processus d'inférence en cause, indépendamment de ces processus en cause, des facteurs qui les déclenchent, des principes ou stratégies utilisées, ils reconnaissent que de tels processus, quels qu'ils soient, fonctionnent quand même très bien. Pour Brown et Yule (1983) le discours ordinaire utilise une forme bien plus relâchée d'inférence informelle que, selon Beaugrande et Dressler (1981), les logiques traditionnelles ne peuvent pas expliquer. Dans ce cas, ce n'est pas la validité logique de la procédure qui est en jeu, mais le fait qu'une telle procédure donne des résultats satisfaisants dans le processus de raisonnement informel. Dans les termes de Sperber et Wilson,

Il est donc tentant de considérer que l'inférence non-démonstrative consiste à appliquer des règles d'inférence non-déductives. Mais cette tentation repose plus sur une analogie que sur des arguments. En fait, on a de bonnes raisons de douter que l'inférence non-démonstrative spontanée, telle que la pratiquent les humains, mette en jeu des règles d'inférence non-déductives. (1989, p. 107)

D'après Sperber et Wilson, si l'inférence non-démonstrative

contexte qui soit suffisant pour la pertinence dans la mesure où il produit une interprétation

spontanée des humains est “moins un processus logique qu’une forme d’imagination réaliste soumise à des contraintes adéquates”, le mieux serait de la considérer comme étant réussie (efficace) ou non, au lieu de logiquement valide ou non. En effet, une telle inférence n’est pas dans sa totalité un processus déductif. De toute façon, Sperber et Wilson ne partagent pas le point de vue des pragmaticiens ci-dessus; leur modèle de la communication ostensive-inférentielle ne comporte pas, dans une telle inférence non-démonstrative, que des règles *déductives*, lesquelles, selon eux, sont les seules règles logiques dont dispose spontanément l’esprit humain. En d’autres mots, la formation par déduction d’hypothèses spontanées et non conscientes est un processus basique de l’inférence non-démonstrative.

Dès lors, ces auteurs suggèrent comment des inférences non conscientes peuvent être exploitées dans les raisonnements conscients. Les nombreuses données et les hypothèses interagissent dans le processus de compréhension, c’est-à-dire, le processus inférentiel, de façon automatique et non consciente; mais seules les données immédiatement accessibles entrent en ligne de compte.

Une telle habileté déductive non consciente est décrite par un dispositif déductif qui prend pour input un ensemble d’hypothèses et qui déduit

cohérente avec le principe de pertinence.

toutes les conclusions que l'on peut inférer des prémisses. Premièrement, Sperber et Wilson assument que les règles disponibles dans le dispositif déductif ne sont pas les règles d'une logique standard. S'il le pouvait, le dispositif déductif dériverait un nombre infini de conclusions (parmi lesquelles un grand nombre de conclusions triviales) à partir d'un ensemble quelconque de prémisses. Par exemple, on pourrait avoir la règle standard *Introduction de* et:

- (13) a) Input: P , Q .
 Output: P et Q .
- b) Input: P et Q , Q .
 Output: $((P$ et $Q)$ et $Q)$ et ainsi de suite à
 l'infini (par répétition).

Cependant, du point de vue psychologique, il n'y a aucune raison pour accepter une telle règle dans le système déductif (si on fait l'hypothèse que l'esprit humain est orienté vers la pertinence) car elle permet la redondance des informations et produit ainsi des implications triviales.

D'autre part, les systèmes logiques usuels n'ont pas de règle déductive permettant la dérivation suivante (Sperber et Wilson, 1989, p. 133):

- (14) (a) Tous les célibataires sont heureux.
- (b) Toutes les personnes qui n'ont jamais été mariées
sont heureuses.

Comme on peut le voir, (b) est une conséquence nécessaire de (a) car il n'existe aucun état de choses concevable dans lequel l'hypothèse (a) serait vraie et l'hypothèse (b) fausse. Comme le remarquent Sperber et Wilson, les systèmes logiques usuels ne tiennent compte que des conséquences nécessaires associées à des "particules logiques" comme *et*, *ou*, *non*, *un*, et *tous*. Selon eux,

Les logiciens s'intéressent à la nature de tous les systèmes déductifs concevables, qu'ils aient ou non une réalité psychologique. Par contre, cette question est d'un intérêt majeur pour la psychologie cognitive, et pour la théorie pragmatique en particulier. Nous supposons, comme le font la plupart des chercheurs dans ce domaine, qu'il existe un ensemble de règles déductives qui sont spontanément mises en oeuvre dans le traitement déductif de l'information. (p. 133)

Conformément à Sperber et Wilson, l'ensemble de règles déductives sur lequel se fonde l'habileté humaine à faire des inférences démonstratives spontanées sont des opérations qui prennent en considération les propriétés sémantiques des hypothèses qui sont reflétées dans leurs formes; les éléments qui composent les formes logiques (en particulier les formes

propositionnelles) des hypothèses sont appelés *concepts* - la forme des hypothèses est la forme en vertu de laquelle des règles déductives peuvent s'y appliquer, et les concepts régissent l'application de ces règles déductives.

Sperber et Wilson remarquent que d'un point de vue formel, ils considèrent un concept comme une adresse qui remplit deux fonctions différentes (mais complémentaires): premièrement, un concept est une adresse dans la mémoire du système dans lequel figurent des informations de trois types - des informations logiques, des informations encyclopédiques et des informations lexicales; deuxièmement, cette adresse peut figurer comme constituant dans une forme logique (la forme logique de l'énoncé est une suite structurée de concepts) que permettra l'application de règles déductives propres. Comme on peut le voir, ces fonctions sont complémentaires: "en effet, lorsque l'adresse d'un certain concept figure dans une forme logique en cours de traitement, sa présence donne accès aux informations de divers types contenues à cette adresse dans la mémoire" (p. 135).

L'entrée logique d'un concept est composée de règles déductives avec input consistant en une ou deux prémisses et output consistant en une conclusion. Sperber et Wilson supposent que les règles disponibles pour le dispositif déductif ne sont pas des règles des systèmes logiques usuels: elles ne

sont que des règles d'*élimination* liées à des concepts⁸. Selon eux, une règle d'élimination d'un concept est une règle qui s'applique à des prémisses où le concept figure et qui fournit une conclusion (où ce concept ou du moins une occurrence de ce concept ne figure plus). Par exemple, les règles d'*élimination de "et"*:

(15) Input: $(P \text{ et } Q)$.
 Output: P , ou Output: Q .

Cette règle s'applique à des prémisses contenant au moins une occurrence du "concept" *et* et elle produit des conclusions où cette occurrence ne figure plus.

(16) les règles de *modus tollendo ponens* (élimination de "ou"):

⁸ Ici il nous semble y avoir une certaine confusion terminologique relativement à l'usage que font Sperber et Wilson du mot "concept". Ils utilisent "concept" en deux sens différents. Premièrement, comme "éléments" des hypothèses - ce qui est déjà dans l'usage traditionnel. Deuxièmement, comme "constantes logiques" (*et, ou, si... alors*) - ce qui nous semble étrange vu que les constantes logiques n'ont pas une signification par elles-mêmes, elles sont *syncatégorématiques* et, pour cette raison, pourquoi les traiter comme des concepts? Même dans le cas du dispositif cognitif décrit par Sperber et Wilson, il nous semble étrange d'appeler "constantes logiques" des concepts seulement parce qu'elles pourraient avoir une certaine représentation mentale. Les constantes logiques ne font aucune contribution au contenu, elles contribuent seulement à la détermination des conditions de vérité. Et quant aux concepts, ils ont à voir avec le contenu. Nous reviendrons sur ce sujet au Chapitre II pour mieux préciser notre divergence sur l'usage que font Sperber et Wilson de mots

Input: $(P \text{ ou } Q), \text{ non-}P$

Output: Q , ou

Input: $(P \text{ ou } Q), \text{ non-}Q$.

Output: P .

Cette règle s'applique à des prémisses contenant au moins une occurrence du "concept" *ou* et elle engendre une conclusion où cette occurrence de *ou* ne figure plus.

(17) les règles de *modus ponens* (élimination de "si... alors"):

Input: $(\text{si } P \text{ alors } Q), P$.

Output: Q .

Cette règle s'applique à des prémisses contenant au moins une occurrence du "concept" *si...alors*, et elle engendre une conclusion où cette occurrence de *si...alors* ne figure plus.

(18) et les règles de concept logique qui permettent de déduire la conclusion (b) à partir de la prémisse (a) - dans l'exemple (14) ci-dessus.

comme "concept" et "forme logique" des hypothèses.

L'*entrée encyclopédique* d'un concept contient des informations à propos de l'extension ou de la dénotation du concept, des objets, des événements ou des propriétés des concepts qui ne sont pas seulement logiques ou lexicales mais aussi empiriques. Bien entendu, l'interprétation de l'énoncé dépend toujours d'un contexte; le contexte fournit d'autres informations (sous forme propositionnelle) stockées sous les adresses des concepts qui auront une influence sur l'énoncé. Finalement, l'*entrée lexicale* est l'ensemble des informations qui portent sur le mot ou la phrase qui exprime le concept dans une langue. Sperber et Wilson supposent qu'elle contient des informations phonologiques et syntaxiques.

À partir d'un ensemble fini de prémisses, un tel dispositif déductif humain, dont les seules règles sont celles d'élimination liées à des "concepts", permet de déduire un ensemble fini de conclusions non triviales. Selon la définition d'*implication logique non triviale*,

(19) Un ensemble d'hypothèses *P* implique logiquement et non trivialement une hypothèse *Q* si et seulement si, lorsque *P* est l'ensemble des thèses initiales d'une dérivation qui ne fait appel qu'à des règles d'élimination, *Q* appartient à l'ensemble des thèses finales. (p. 152)

Comme on l'a déjà remarqué plus haut, les règles d'inférence

déductives permettent de dériver les conclusions intéressantes logiquement impliquées par les prémisses. Mais du point de vue cognitif, certaines conclusions ne sont pas dérivées déductivement, comme par exemple: P peut impliquer, par la règle d'introduction du "ou", un nombre infini de conclusions de la forme $(P \text{ ou } Q)$ où Q serait dans cette dérivation, une hypothèse arbitraire et, par conséquent, cognitivement non réalisée. De même, les règles d'introduction du "et": Input: P . Output: $(P \text{ et } P)$; les règles d'introduction de la négation double: Input: P . Output: non-non P , etc. (dont l'output contient tous les concepts contenus dans l'input, et au moins un concept en plus) ne sont pas d'intérêt d'un point de vue cognitif puisqu'elles n'ont aucune influence dans le processus de compréhension. Sperber et Wilson soutiennent que dans des conditions normales, une hypothèse est traitée dans le contexte d'autres hypothèses, mais seules les implications non triviales sont calculées (une implication est triviale si elle est produite par une règle d'introduction).

En plus de la distinction entre les implications triviales et non-triviales, Sperber et Wilson font encore une distinction entre les implications analytiques et les implications synthétiques (une règle analytique prend pour input une seule prémisses, alors qu'une règle synthétique en prend deux):

(20) Un ensemble d'hypothèses P *implique analytiquement* une hypothèse

Q si et seulement si Q est l'une des thèses finales d'une déduction dont l'ensemble des thèses initiales est P , et dans laquelle ne sont appliquées que des règles analytiques. (p. 161)

Comme conséquence de cette définition nous avons que toute hypothèse s'implique analytiquement elle-même. Une *implication synthétique* est définie de la façon suivante:

(21) Un ensemble d'hypothèses P *implique synthétiquement* une hypothèse Q si et seulement si Q est l'une des thèses finales d'une déduction dont l'ensemble des thèses initiales est P , et Q n'est pas une implication analytique de P . (p. 161)

L'exemple suivant est un cas d'implication analytique:

(a) Pierre va à l'université et Marie part en vacances.

La règle analytique d'élimination du *et* permet de dériver à partir de la seule prémisse (a) les conclusions suivantes:

(b) Pierre va à l'université.

(c) Marie part en vacances.

Dans ce cas, (a) implique analytiquement (b - c).

Pour ce qui est des implications synthétiques, considérons

l'ensemble d'hypothèses suivant:

- (d) Pierre va à l'université.
- (e) Si Pierre va à l'université, les cours ont commencé.
- (f) Si les cours ont commencé, Pierre doit arriver à l'heure.

La règle synthétique *modus ponens* permet de dériver à partir de (d - f) les conclusions (g - h) suivantes:

- (g) Les cours ont commencé.
- (h) Pierre doit arriver à l'heure.

Et dans ce cas, (d - f) implique synthétiquement (g - h).

Sperber et Wilson ne font pas de distinction entre termes “logiques” et “non logiques”; ils traitent globalement les règles déductives “logiques” et “sémantiques”. Sous ce rapport, nous présentons un exemple “sémantique” d'une implication synthétique. Considérons les hypothèses suivantes:

- (i) La photo de Bertrand Russell est dans le Laboratoire de Philosophie analytique.
- (j) Le Laboratoire de Philosophie Analytique est dans le Pavillon Ringuet.

Une règle synthétique étant quelque chose comme la “règle d’inclusion” permet de dériver:

- (l) La photo de Bertrand Russell est dans le Pavillon Ringuet.

Ainsi, (i - j) implique synthétiquement (l).

Un troisième type d’implication est aussi présentée par Sperber et Wilson. Dans leur terminologie, une inférence qui a pour prémisses l’énoncé P (l’information nouvelle) et une hypothèse contextuelle C (l’information ancienne), est appelée une *contextualisation de P dans le contexte C* . Le résultat ou combinaison conjointe de ces nouvelles et anciennes informations est appelé une *implication contextuelle de P dans C* :

- (22) Un ensemble d’hypothèses P implique contextuellement une

hypothèse Q dans le contexte C si et seulement si:
 (a) l'union de P et de C implique non-trivialement Q ,
 (b) P n'implique pas non-trivialement Q , et
 (c) C n'implique pas non-trivialement Q . (p. 166)

Ainsi, une implication contextuelle est synthétique et non triviale.

Nous reviendrons sur l'*implication non triviale* au chapitre IV, dans le but d'indiquer une lacune concernant cette définition, et de l'analyser à la lumière de notre logique de la pertinence P (qui sera présentée au chapitre III), pour montrer que notre définition de la *pertinence* rend compte d'une telle lacune.

On voit bien que Sperber et Wilson introduisent une distinction entre règles analytique et synthétique dans le but de restreindre les règles d'élimination impliquées dans les implications contextuelles aux seules règles synthétiques. L'hypothèse que les règles d'élimination (de "ou" et de "si ... alors") sont les seules règles de déduction impliquées dans le processus inférentiel nous permet de dériver uniquement les implications qui sont d'intérêt du point de vue communicationnel, à savoir les implications non triviales. Comme la fonction du dispositif déductif est d'analyser et de manipuler le contenu conceptuel des hypothèses, Sperber et Wilson montrent que cette fonction n'est accomplie que par de telles règles d'élimination contenues dans

les entrées logiques des concepts. La notion d'implication contextuelle y joue un rôle décisif puisque une fonction centrale du dispositif déductif est de calculer de façon spontanée, automatique et non consciente les implications contextuelles produites par la synthèse de nouvelles et anciennes informations.

C'est principalement en termes d'implications contextuelles qu'il convient d'analyser l'effet du contexte sur l'interprétation des énoncés et d'expliquer pourquoi une information donnée est traitée dans tel contexte plutôt que dans tel autre (...).

Toutes choses étant égales d'ailleurs, plus cette information nouvelle entraînera d'implications contextuelles, plus elle améliorera la représentation du monde de l'individu. (p. 167)

Dans la terminologie de Sperber et Wilson, une contextualisation de P dans le contexte C donne lieu à ce qu'ils appellent des *effets contextuels*: modifier et améliorer la représentation du monde de l'individu c'est produire un certain effet sur cette représentation à partir de certaines prémisses et des propositions qui constituent le contexte. Ils distinguent trois types d'effets contextuels:

(23) (i) les *implications contextuelles*: le produit obtenu de l'énoncé et du contexte conjointement qui jouent le rôle de prémisses d'une implication synthétique;

(ii) les *renforcements d'hypothèses*: des informations

nouvelles peuvent confirmer des hypothèses anciennes dont on dispose déjà et donc renforce certaines croyances du système;

(iii) les *effacements de prémisses*: des informations nouvelles peuvent affaiblir (ou éliminer) des hypothèses anciennes et donc diminuer la force de certaines croyances du système.

À cet égard, la notion d'effet contextuel est, comme nous le verrons, essentielle à la définition de la pertinence. Sperber et Wilson se servent de cette notion pour caractériser ce qu'est la pertinence d'une hypothèse: "une hypothèse est *pertinente* dans un contexte si et seulement si elle a un effet contextuel dans ce contexte" (p. 187); par ailleurs, "toutes choses étant égales d'ailleurs, plus les effets contextuels sont grands, plus grande est la pertinence" (p. 182).

3- Définition de la pertinence chez Sperber et Wilson

Nous venons de remarquer que la notion d'effet contextuel joue un rôle important dans la théorie de la pertinence de Sperber et Wilson. En effet, deux propriétés indispensables à la compréhension des énoncés peuvent être

décrites à partir de cette notion. La première propriété est que la compréhension met à exécution une synthèse conjointe de certaines hypothèses. La deuxième propriété est que quelques-unes de ces hypothèses constituent une information nouvelle qui est traitée dans un contexte formé par des informations qui y sont déjà traitées.

Mais, comment procédons-nous pour choisir, entre plusieurs, les informations qui formeront le contexte? Selon Sperber et Wilson la sélection de cet ensemble d'informations s'effectue par un *principe de pertinence* selon lequel,

(24) Tout acte de communication ostensive communique la présomption de sa propre pertinence optimale. (p. 237)

Nous avons mentionné dans la section 2 que selon la terminologie de Sperber et Wilson, une information est pertinente dans un contexte quand elle a au moins un effet contextuel dans ce contexte. Mais pour définir la pertinence d'une information, il faut prendre en considération des degrés de pertinence et préciser la façon dont ces degrés sont déterminés. Sperber et Wilson analysent cette situation en termes de comparaison de coûts et de bénéfices. Ils soutiennent que les effets contextuels d'une information dans un contexte sont produits par des processus mentaux qui (de même que les processus biologiques) requièrent le moins d'effort possible. Ainsi, ils

définissent:

(25) *La Pertinence*

Condition comparative 1: une hypothèse est d'autant plus pertinente dans un contexte donné que ses effets contextuels y sont plus importants.

Condition comparative 2: une hypothèse est d'autant plus pertinente dans un contexte donné que l'effort nécessaire pour l'y traiter est moindre. (p. 191)

Comme conséquence de cette définition, nous avons la condition nécessaire et suffisante suivante:

(26) *La pertinence*

Une hypothèse est pertinente dans un contexte si et seulement si elle a un effet contextuel dans ce contexte. (p. 167)

La définition (25) est prise au sens comparatif.⁹ Elle veut dire que la pertinence est susceptible de degrés en fonction de deux facteurs principaux (l'effet et l'effort). Telle définition de la pertinence ne permet des comparaisons claires que dans certains cas:

⁹ Sperber et Wilson soulignent que c'est ce sens comparatif qui intéresse les psychologues puisqu'une définition de la pertinence au sens quantitatif aurait plus d'intérêt pour ceux qui font de la logique et de l'intelligence artificielle. En effet, il serait difficile de calculer le nombre d'effets et d'efforts d'un individu. Alors, Sperber et Wilson soutiennent que de tels effets existent indépendamment du fait de leur représentation conceptuelle. Une fois représentés, on procède au moyen de jugements (intuitifs) comparatifs - ce qui vaut aussi pour la pertinence. En ce sens, les effets contextuels et l'effort d'interprétation d'une information sont des aspects non représentationnels des processus mentaux. Nous reviendrons sur ce sujet dans les chapitres III et IV où nous proposons une évaluation différente (du point de vue logique) des degrés de pertinence.

toutes choses étant égales d'ailleurs, une hypothèse ayant plus d'effets contextuels est plus pertinente; et toutes choses étant égales d'ailleurs, une hypothèse demandant moins d'effort de traitement est plus pertinente. (p. 191)

Ainsi, pour traiter une hypothèse de la façon la plus pertinente possible, l'esprit humain cherche un contexte - vu qu'il n'est pas donné - qui maximise la pertinence.

Le traitement d'un énoncé dans un contexte exige d'un individu un certain effort visant à calculer ses effets contextuels. Un tel traitement commence dans un contexte initial (formé par l'ensemble des croyances qui y sont activées) et qui peut, si nécessaire, être élargi dans le but de calculer plus d'effets contextuels de l'énoncé dans ce contexte élargi; cela conduira dès lors à un coût de traitement (l'effort cognitif) de ces effets et à un accroissement des croyances. De même, l'accès à un contexte exige d'un individu un certain effort. "Moins un contexte est accessible, plus l'effort requis pour y accéder est grand et inversement" (Sperber et Wilson, 1989, p. 216). Son effort consiste donc à trouver un contexte qui garantit un plus grand nombre d'effets contextuels; ce contexte doit être accessible à partir du contexte initial. À la fin de chaque processus déductif, l'individu compte sur un ensemble de contextes accessibles partiellement ordonné.

L'ensemble des contextes accessibles est donc partiellement ordonné par la relation d'inclusion. Cette relation formelle a un analogue psychologique: l'ordre d'inclusion correspond à l'ordre d'accessibilité. Le contexte initial, minimal, est immédiatement donné; viennent ensuite dans l'ordre d'accessibilité les contextes auxquels on peut accéder en une seule étape, c'est-à-dire les contextes qui n'incluent pas d'autres contextes que le contexte initial; puis les contextes auxquels on peut accéder en deux étapes, c'est-à-dire ceux qui n'incluent comme autres contextes que le contexte initial et une extension de ce contexte initial obtenue en une étape; et ainsi de suite. (p. 216)

Alors, choisir un contexte qui maximise la pertinence d'une information est choisir le meilleur contexte pour avoir une relation optimale entre effort et effet. Dans la terminologie de Sperber et Wilson, cela veut dire *traiter optimalement* l'information.

(27) *La pertinence pour un individu* (définition comparative)

Condition comparative 1: une hypothèse est d'autant plus pertinente pour un individu que les effets contextuels qu'elle entraîne lorsqu'elle est traitée optimalement sont importants.

Condition comparative 2: une hypothèse est d'autant plus pertinente pour un individu que l'effort requis pour la traiter optimalement est moindre. (p. 219)

On a vu dans la section 1, définition (12) de la communication ostensive-inférentielle, que le locuteur choisit de rendre manifeste son intention de rendre manifeste un ensemble d'hypothèses *I* ainsi que cet ensemble d'hypothèses *I* lui-même. Mais pour qu'un acte de communication ostensive réussisse, le locuteur doit procéder de façon à attirer l'attention de l'auditeur. Comment le stimulus ostensif du locuteur peut-il rendre manifeste son intention informative? Or, le locuteur a l'intention de rendre manifeste à l'auditeur que

son stimulus est pertinent. En y supposant un certain degré de mutualité, on pourrait dire qu'il est mutuellement manifeste pour les deux, communicateur et destinataire, que le locuteur a l'intention qu'il soit manifeste à l'auditeur que son stimulus est pertinent. Donc, la définition de la communication ostensive-inférentielle implique que l'acte de communication ostensive-inférentielle du communicateur s'accompagne d'une garantie de pertinence optimale, c'est-à-dire un communicateur ostensif communique nécessairement que le stimulus qu'il produit est pertinent pour l'auditeur. Selon Sperber et Wilson, un acte de communication ostensive communique automatiquement une *présomption de pertinence optimale*:

- (28) (a) L'ensemble d'hypothèses I que le communicateur veut rendre manifestes au destinataire est suffisamment pertinent pour que le stimulus ostensif mérite d'être traité par le destinataire.
 (b) Le stimulus ostensif est le plus pertinent de tous ceux que le communicateur pouvait utiliser pour communiquer I . (p 237)

Le principe de pertinence selon lequel "tout acte de communication ostensive communique la présomption de sa propre pertinence optimale" est une généralisation sur la communication ostensive-inférentielle: pour expliquer qu'un énoncé est digne de l'attention de l'auditeur et mérite un traitement interprétatif, il doit comporter une garantie de pertinence.

Il est important de remarquer que le communicateur fait de son mieux pour rendre manifeste à l'auditeur que le stimulus ostensif qu'il utilise est aussi pertinent qu'il le faut pour retenir l'attention: s'il ne le réussit pas, il peut malgré cela réussir à rendre manifeste son effort d'être optimalement pertinent.

Le "degré suffisant pour mériter l'attention de l'auditeur" dépend non pas seulement de la façon dont l'information se présente à la compréhension, mais aussi du degré de "réceptivité intellectuelle" de l'interlocuteur.

Mais comment un stimulus produit par le communicateur peut-il rendre son intention informative mutuellement manifeste et conduire à la réalisation de son intention communicative? En communication ostensive, l'effet communicatif vise à la reconnaissance de l'intention informative; généralement, l'effet informatif prétendu se produira et sera observé seulement à partir de la reconnaissance de cette intention informative. Ainsi, pour que l'intention informative devienne mutuellement manifeste, il faut passer par des étapes inférentielles comme: le stimulus doit rendre manifeste, dans l'environnement cognitif mutuel du communicateur et du destinataire, des hypothèses qui permettent d'inférer l'intention informative (tout en considérant qu'il est manifeste que le stimulus est ostensif). Cela se fait en produisant un

stimulus manifestement intentionnel qui attire l'attention et qui est pertinent - dans la mesure où il est traité comme source d'information relativement aux intentions du communicateur.

Comme la nature ostensive d'un stimulus est mutuellement manifeste aux interlocuteurs, de même il est mutuellement manifeste que le communicateur a l'intention informative de rendre manifeste au destinataire un ensemble particulier d'hypothèses I. Ainsi, le problème de l'identification de l'intention informative est au fond le problème de l'identification de cet ensemble d'hypothèses I. Dans cet ordre d'idées, le principe de pertinence permet d'identifier un élément de I, soit la présomption de pertinence, qui peut être confirmée ou infirmée par le contenu de I - cette possibilité de confirmation ou d'infirmerie permet au destinataire d'identifier l'intention informative du communicateur, ou plus précisément l'ensemble d'hypothèses I. Il faut tout d'abord que le destinataire présuppose que le communicateur communique rationnellement. De cette façon, l'auditeur peut trouver dans l'ensemble I que le communicateur avait des raisons d'imaginer qu'il confirmerait la présomption de pertinence. En bref, cette présupposition de rationalité est une condition *sine qua non* pour l'identification des intentions informatives et pour l'identification inférentielle des intentions en général.

CHAPITRE II

APPRÉCIATIONS DE QUELQUES CRITIQUES SUR LE DISPOSITIF DÉDUCTIF DE SPERBER ET WILSON

Introduction

Ainsi que nous venons de le présenter dans la section 2 au chapitre précédent, Sperber et Wilson (1989) proposent un modèle explicatif pour la pragmatique linguistique en généralisant ce modèle communicationnel aux activités mentales et cognitives humaines. Pour eux, ces activités sont orientées vers la pertinence.

Selon Sperber et Wilson, leur dispositif déductif peut automatiquement calculer l'ensemble complet d'implications contextuelles d'une hypothèse donnée dans un certain contexte. L'inférence est contrôlée par deux méthodes principales. Premièrement, la distinction entre implications triviales et non triviales impose, comme eux-mêmes le remarquent, une contrainte interne substantielle sur la classe des implications que le mécanisme

inférentiel peut calculer. Deuxièmement, la classe des inférences tirées est déterminée par des facteurs tels que la sélection par le contexte, la saillance perceptuelle et l'attention du sujet; ces facteurs sont engrenés de telle façon qu'ils contribuent à la maximisation de la pertinence.

Sperber et Wilson admettent que les processus d'inférence qui interviennent dans la compréhension sont des processus non démonstratifs et que c'est au destinataire de former une hypothèse basée surtout sur des indices fournis par le comportement ostensif du communicateur à propos de son intention communicative. Le communicateur doit fournir au destinataire des indices qui lui permettent de concevoir et de confirmer ou d'infirmer quelques-unes de ses hypothèses; quoiqu'une telle hypothèse (qui, selon eux, n'est jamais certaine) ne soit pas démontrable, elle peut néanmoins être confirmée.

Mais en quoi consiste le processus de construction et de confirmation d'une hypothèse qu'un locuteur peut avoir l'intention de communiquer?

1. **Quelques critiques adressées au dispositif déductif de Sperber et Wilson**

Comme on l'a déjà vu, l'inférence démonstrative est le processus qui consiste à appliquer des règles déductives à un ensemble de prémisses prises au départ; ce sont ces règles qui rendent possible d'engendrer toutes les conclusions qui sont logiquement impliquées par l'ensemble de prémisses; leur fonction est de garantir la validité logique des inférences qu'elles gouvernent: de la vérité de leurs prémisses suit nécessairement la vérité de leurs conclusions. En d'autres mots, une hypothèse provenant d'une inférence démonstrative dépend de règles déductives que la logique nous offre à travers plusieurs modèles d'inférence démonstrative. Cependant, dans le cas d'une hypothèse provenant d'une inférence non-démonstrative, elle ne peut qu'être confirmée par un modèle adéquat de l'inférence non-démonstrative, c'est-à-dire, un modèle constitué de règles inductives; dans ce cas, la validité des conclusions non-démonstratives dépend des deux choses suivantes: premièrement, la formation d'hypothèses et deuxièmement, la confirmation qui justifie ces hypothèses. Mais, comme Sperber et Wilson (1989) le reconnaissent, "il n'existe pas de logique inductive bien développée qui pourrait constituer un modèle plausible des processus cognitifs centraux" (p.

106), et de plus, il n'est pas sûr qu'un tel modèle adéquat de l'inférence non-démonstrative spontanée puisse être présenté par un système de logique inductive.

D'un autre côté, ils mettent en doute que l'inférence non-démonstrative spontanée (la seule qui les intéresse - par son caractère psychologique) pratiquée par les humains repose sur des règles d'inférence non-déductives. Évidemment, "...des règles d'inférence non-démonstratives ne permettent pas d'*engendrer* toutes les conclusions intéressantes que confirment un ensemble donné de prémisses" (p. 107).

Dès lors, Sperber et Wilson font l'hypothèse que les seules règles logiques dont dispose spontanément l'esprit humain sont des règles déductives. Leur objectif est de montrer que les règles déductives sont essentielles dans l'inférence non-démonstrative telle que la pratiquent spontanément les humains, quoique celle-ci n'est pas en totalité un processus logique; pour eux, l'inférence non-démonstrative ne fait que "comporter" une déduction comme sous-partie; certainement la validité d'une inférence déductive ne préserve pas la validité de l'inférence non-démonstrative dans laquelle elle est une sous-partie.

On a remarqué auparavant que, puisque l'inférence non-démonstrative pratiquée par les humains n'est pas complètement un processus

logique, au lieu de l'examiner comme étant logiquement valide ou non, Sperber et Wilson proposent de la considérer comme étant réussie ou non, efficace ou non - la formation d'une hypothèse demande l'appui des règles déductives, tandis que la confirmation d'une hypothèse n'est qu'un sous-produit des différents traitements, déductifs ou non, dont l'hypothèse a été l'objet.

Le fait de considérer comme non-démonstratifs les processus d'inférence impliqués dans la compréhension justifie l'intérêt de Sperber et Wilson pour une analyse de l'inférence non-démonstrative dans le but d'expliquer le rôle de la pertinence dans la communication et dans la cognition: comme on le sait, d'après eux, une information est pertinente pour un individu quand elle modifie et améliore sa représentation globale du monde. Une telle représentation du monde n'est qu'un ensemble d'hypothèses factuelles.

À son tour, les hypothèses factuelles sont les représentations contenues dans la mémoire et qui sont traitées par l'esprit comme des descriptions vraies du monde réel (ou de faits), sans pour autant rendre explicite le fait qu'il s'agit là simplement d'hypothèses exprimées ou représentées. En tant qu'hypothèses factuelles, elles forment en particulier le domaine auquel s'appliquent les processus d'inférence non-démonstrative spontanée. La façon dont chaque hypothèse factuelle (ou information nouvelle) participe à l'amélioration d'une telle représentation peut être décrite en prenant

pour base des opérations du dispositif déductif: chaque information nouvelle se combine avec un ensemble d'informations antérieures en servant de point de départ à un processus d'inférence.

Rappelons encore que selon Sperber et Wilson (1989), une représentation conceptuelle est à la fois un état mental et un état cérébral, mais elle doit avoir aussi des propriétés logiques qui lui permettront de se soumettre à des règles déductives, par le biais de sa forme logique (qu'est l'ensemble de ces propriétés logiques), et à des relations de contradiction ou d'implication avec d'autres représentations conceptuelles et servant de prémisses ou de conclusions à des règles déductives.

Un point qu'il faut bien préciser est la façon comment Sperber et Wilson distinguent les opérations logiques des autres opérations formelles; pour eux, les opérations logiques préservent la vérité: quand une représentation P est vraie, elle donne lieu à une autre représentation Q vraie. Dans le cas d'une représentation dans laquelle on efface le premier constituant, cela n'est qu'une opération formelle (sans être, bien entendu, une opération logique). Comme eux-mêmes nous en avisent, à cause du lien entre vérité et logique, il est tentant de penser que seulement les représentations conceptuelles capables d'être vraies ou fausses aient une forme logique.

Nous adopterons un autre point de vue. Nous considérons que, pour qu'une représentation puisse être soumise à des opérations logiques, il suffit qu'elle soit bien formée, alors que, pour être vraie ou fausse, il faut aussi qu'elle soit sémantiquement complète, c'est-à-dire qu'elle représente un état de choses, possible ou réel, dont l'existence la rendrait vraie. Une structure conceptuelle sémantiquement incomplète peut néanmoins être bien formée et peut être soumise à des opérations logiques. (p. 114)

Ainsi, ils introduisent la définition de la forme logique *propositionnelle*:

Disons qu'une forme logique est *propositionnelle* si elle est sémantiquement complète et donc susceptible d'être vraie ou fausse, et qu'elle est *non-propositionnelle* dans le cas contraire. (p. 114)

Sperber et Wilson proposent deux exemples, l'un formel, l'autre psychologique, de formes logiques non-propositionnelles:

a) une formule du calcul de prédicats contenant une variable libre est un exemple formel d'une forme logique qui est syntaxiquement bien formée sans être entièrement propositionnelle. (p. 114)

Ils soutiennent que les formes logiques non-propositionnelles, ou incomplètes, ont une fonction importante dans la cognition bien que seules les formes propositionnelles complètes représentent des états de choses précis; elles sont essentielles dans la représentation globale que l'individu a du monde.

Pour ce qui est de l'"exemple psychologique", ils remarquent que

“le sens d’une phrase constitue un exemple psychologique d’une forme logique non propositionnelle”. L’exemple donné est le suivant:

b) Elle le tenait à la main.

Selon Sperber et Wilson, du fait que tant “elle” que “le” dans l’exemple (b) ne sont pas conformes à des concepts précis, la phrase ci-dessus n’est ni vraie ni fausse, quoique elle maintient des propriétés logiques telle que le fait qu’elle implique (c) et contredit (d):

c) Elle tenait quelque chose à la main.

d) Personne n’a jamais rien tenu.

À la suite de toutes ces considérations, nous présenterons quelques commentaires adressés par Stuart J. Russell, Pieter A. M. Seuren, John Macnamara et Elizabeth Hinkelman au mécanisme déductif de Sperber et Wilson, lesquels sont présentés dans Sperber et Wilson (1987).

Stuart J. Russell remarque que dans le cadre théorique de Sperber et Wilson toutes les règles d’inférence sont explicitement représentées. Ainsi, selon lui, un tel schéma ne peut pas fonctionner: “tout ce qui combine les règles d’inférence avec leurs prémisses doit aussi être une règle d’inférence,

mais représentée de façon procédurale” (*apud* Sperber et Wilson, 1987, p. 731). À cette objection Sperber et Wilson répondent que dans leur cadre théorique, des règles d’inférences ne font pas partie des données de base et elles ne sont pas présentées en forme propositionnelle. Pour eux, la distinction entre représentation et *computation* est fondamentale et leurs règles d’inférences ont à voir avec l’aspect computationnel et non l’aspect représentationnel. De telles règles ne “combinent” pas avec les prémisses, elles opèrent sur les prémisses - ce qui n’est pas déterminé par de nouvelles règles ou représentations de quelque sorte que ce soit, mais par la façon dont le mécanisme déductif est construit.

En liaison avec cela, nous nous demandons, comme nous l’avons déjà remarqué dans la section 2, pourquoi donc Sperber et Wilson appellent “et”, “ou”, “si ... alors” des concepts? Or, *concept* est une notion éminemment représentationnelle. Nous insistons sur ce point: nous semble bien confuse cette particularité terminologique de Sperber et Wilson et surtout maintenant que, dans leur réponse à Stuart J. Russell, ils mettent en évidence la distinction entre représentation et computation. En effet, il faut bien remarquer que la représentation ne rend pas compte de la computation. D’après nous, c’est la logique qui doit en rendre compte. En affirmant que leurs règles d’inférences ont à voir avec l’aspect computationnel et non pas l’aspect représentationnel,

nous comprenons par là que c'est ce sens logique qu'il nous faut ici comprendre, pour la simple raison que la notion de computation renvoie à celle de *logique*. Au fond, la logique est computation et telle quelle, elle est distincte de la représentation. Dans ce cas, nous observons que Sperber et Wilson veulent bien distinguer tout ce qui fait partie de la logique de ce qui fait partie de leur théorie. Mais il faudrait d'abord bien éclairer cette confusion terminologique.

L'objection de Stuart J. Russell porte aussi sur le mécanisme déductif, lequel donne accès seulement aux règles d'élimination. Pour lui, tel mécanisme serait incapable d'utiliser "les lois de commutativité et les axiomes définissant le domaine inductif" (*apud* Sperber et Wilson, 1987, p.731). Mais, Sperber et Wilson répondent que dans leur cadre théorique, les axiomes sont traités comme information dans les entrées encyclopédiques plutôt que comme des règles d'inférences. Pour cette raison, le mécanisme déductif a accès à eux comme prémisses.

Rappelons que par domaine inductif Sperber et Wilson entendent tout ce qui a déjà été observé, c'est-à-dire, tout ce qui a été objet d'expérience. *Une pomme que j'ai vu tomber par terre aujourd'hui, une pomme que j'ai vu tomber par terre hier*, ce sont des objets d'expérience. Mais ces objets seraient utilisés comme des prémisses?

Bien, dans les entrées encyclopédiques nous avons les informations analytiques telles que *Les oiseaux sont vertébrés ovipares*, et les informations empiriques comme par exemple, *Les oiseaux chantent* (bien sûr, on peut trouver un oiseau qui ne chante pas). Il existe beaucoup d'informations dans les entrées encyclopédiques, et chaque entrée a beaucoup d'informations qui sont enregistrées d'une certaine façon dans la mémoire pendant une longue période; des études récentes¹⁰ montrent que, en général, la mémoire fonctionne par blocs sémantiques; en comparaison avec des ordinateurs traditionnels, notre mémoire fonctionne plutôt de la manière suivante: une conversation entre deux personnes donne lieu à un flux d'information énorme. Cette information est disponible immédiatement. Mais, pour l'ordinateur, le processus est plus compliqué; il lui faut parfois chercher une entrée, déplacer des informations, donner des réponses, bref, il s'agit ici d'un processus très long. Pour les humains, les informations sont immédiatement disponibles, mais elles ne sont pas isolées, elles viennent en blocs sémantiques qui sont "appelés à la surface" et y restent disponibles. Ce que l'on trouve là, c'est quelque chose ressemblant aux entrées encyclopédiques de Sperber et Wilson. Par exemple, quand nous pensons à un *oiseau*, nous avons diverses informations qui viennent

¹⁰ Pour des informations plus complètes à propos du problème de la mémoire, voir Chemiak (1986).

immédiatement à l'esprit et ensuite elles retournent simplement dans la "pénombre". Quand nous en avons besoin, elles reviennent à l'esprit.

Dans le fond, Sperber et Wilson semblent admettre la logique comme étant quelque chose de "tacite" qui n'est jamais représentée; les règles d'inférence ne sont jamais représentées dans le dispositif inférentiel.

Seuren (*apud* Sperber et Wilson, 1987, p. 733) regrette que le système de Sperber et Wilson exclut l'inférence suivante:

- (a) Henry et Jack sont entrés.
- (b) Deux personnes sont entrées.

Pour Sperber et Wilson cette inférence implique l'habileté à compter, à accomplir des calculs mathématiques et ils ne sont pas intéressés à inclure de telles habiletés dans le mécanisme déductif. La capacité à compter permet à l'individu de construire dans une occasion appropriée, la prémisse "Si Henry et Jack sont entrés, alors deux personnes sont entrées" laquelle combinera avec (a) pour produire (b) comme une implication contextuelle.

Une autre objection portant sur le mécanisme déductif de Sperber et Wilson, plus précisément sur l'élimination de règles d'introduction, est faite par MacNamara: "Comment pourrait-il y avoir quelque chose sur laquelle les règles d'élimination pourraient fonctionner, s'il n'y a eu auparavant aucune

règle d'introduction?" (*apud* Sperber et Wilson, 1987, p. 725). Sperber et Wilson argumentent que "les règles d'élimination travaillent sur n'importe quel ensemble d'hypothèses soumises au mécanisme déductif (...) la déduction inconsciente n'est pas la seule source d'hypothèses factuelles: la perception, le décodage linguistique, les hypothèses et schémas d'hypothèses emmagasinés en mémoire, et le raisonnement conscient sont des sources supplémentaires qui ne sont pas affectées par l'élimination de règles d'introduction" (p. 741).

Considérons quelques-uns des exemples présentés par Sperber et Wilson (1989, p. 128-131):

(a) Ceci est une orchidée.

Selon eux, la description élémentaire de ce stimulus devient une hypothèse "forte" car, comme ils le soulignent, les mécanismes de la perception (issus d'une longue évolution biologique) permettent d'associer à un stimulus sensoriel une identification conceptuelle de ce stimulus.

Relativement au décodage linguistique, les mécanismes d'input linguistique associent une forme logique à des stimuli sensoriels d'un type particulier; les formes logiques recouvrées par décodage ne sont pas en totalité propositionnelles, mais elles peuvent être complétées pour parvenir à des

formes propositionnelles - ce qui ne garantit pas encore qu'elles constituent des hypothèses factuelles. "Toutefois, il existe une procédure générale qui permet d'enchâsser la forme propositionnelle obtenue en complétant la forme logique de la phrase énoncée à l'intérieur d'une hypothèse sur ce qu'a dit le locuteur" (Sperber et Wilson, 1989, p. 128). Sur ce point, voyons leur exemple:

Si quelqu'un entend Pierre énoncer (b) à l'instant t :

(b) [zemalalat $\in t$]

Le décodage de son énoncé donnera la forme logique (c):

(c) J'ai mal à la tête.

L'on pourra compléter (c) pour parvenir à la forme propositionnelle (d):

(d) Pierre a mal à la tête à l'instant t .

L'on pourra encore enchâsser (d) dans le schéma d'hypothèse (e):

(e) Pierre dit que.....

Pour parvenir à l'hypothèse (f):

(f) Pierre dit que Pierre a mal à la tête à l'instant t .

Un autre point que nous semble étrange c'est le fait que Sperber et Wilson considèrent (c) comme étant la "forme logique" donnée à partir du

décodage de (b). Or, en (b) nous avons seulement la forme phonologique et, à notre avis, le décodage de (b) n'en donne qu'une transcription graphique telle que (c). Comme on peut le remarquer, il n'y a aucune formalisation dans (c) pour qu'on puisse la considérer comme une forme logique. Une forme logique d'un énoncé est la structure profonde qui détermine les conditions de vérité ainsi que les inférences possibles à partir de l'énoncé en question. À un certain stimulus sensoriel est associée une image acoustique; à une image acoustique est associée une interprétation sémantique qui peut être sujette à discussion: si l'énoncé est ambigu, cela rend l'interprétation sémantique plus difficile - ce qui n'arrive pas si l'énoncé est non ambigu. En (c) sont associés des concepts à l'image acoustique de (b) qui vont déterminer des conditions de vérité différentes, selon l'individu associé au pronom "je" - car il s'agit d'un exemple contenant un terme singulier indexical. Alors, nous divergeons entièrement de Sperber et Wilson sur la façon de définir la "forme logique" d'un énoncé.

Sperber et Wilson supposent que la mémoire conceptuelle contient un ensemble d'hypothèses ainsi que des schémas d'hypothèses, à savoir des "formes logiques" (conformément à leur terminologie) auxquelles peuvent être ajoutés des éléments afin d'obtenir des formes propositionnelles formant des hypothèses factuelles. Par exemple, considérons le schéma d'hypothèse (g):

(g) La température extérieure est dex..... degrés.

L'on pourrait rendre (g) complet en y ajoutant ce qui manque pour obtenir l'hypothèse (h):

(h) La température extérieure est de moins six degrés.

Ils remarquent que lorsqu'on traite des hypothèses qui correspondent à certains schémas comme (i):

(i) Si P, alors Q ,

des schémas apparentés sont normalement mis en oeuvre en vue d'envisager des hypothèses complémentaires comme, par exemple, (j) ou (l):

(j) Si non-P, alors non-Q.

(l) Si Q, alors (Q parce que P).

Les hypothèses provenant de la mémoire ont une certaine force tandis que les hypothèses complémentaires possèdent une certaine plausibilité de départ sur quoi on peut se fonder pour les prendre en considération. Comme le remarquent Sperber et Wilson, leur force ultérieure dépendra du rôle cognitif qu'elles auront joué.

Si l'on utilise un ensemble d'hypothèses comme prémisses tel que:

(a) Ceci est une orchidée.

(m) Les orchidées sont des fleurs exotiques.

L'on peut en dériver, par un processus déductif, d'autre hypothèse telle que (n):

(n) Ceci est une fleur exotique.

De même, à partir des prémisses suivantes:

(o) Lorsque la température extérieure est inférieure à moins cinq degrés, l'étang est gelé.

(h) La température extérieure est de moins six degrés.

L'on peut dériver l'hypothèse suivante:

(p) L'étang est gelé.

Nous avons souligné auparavant que Sperber et Wilson soutiennent que “la formation d'hypothèses par déduction est le processus central sur lequel repose l'inférence non-démonstrative” (1989, p. 131) et que le mécanisme déductif opère sur des représentations conceptuelles en vertu de leur structure logique et des concepts qu'elle contient. L'étude de l'inférence spontanée est, selon eux, une précondition nécessaire à toute investigation correcte de toutes les formes d'inférence humaine.

Contrairement à ce que soutiennent Sperber et Wilson, Hinkelman

remarque que le mécanisme déductif pourvu des seules règles d'élimination peut cependant engendrer un nombre infini de conclusions à partir d'un nombre fini de prémisses. À cette objection, Sperber et Wilson répondent que le dispositif déductif n'est pas "une logique"; il n'est qu'un "mécanisme computationnel limité dans ses opérations non seulement par des règles qu'il applique, mais aussi par la manière selon laquelle il les applique" (1987, p. 741).

2. Quelques considérations critiques à propos de la réponse de Sperber et Wilson à la critique de Hinkelman

Nous aimerions examiner plus intensément la réponse de Sperber et Wilson à la critique de Hinkelman par rapport aux deux points suivants: le premier point concerne l'affirmation que leur dispositif déductif n'est pas "une logique"; le deuxième point concerne ce que Sperber et Wilson affirment à propos de la manière dont leur mécanisme computationnel applique les règles.

Pour ce qui est du premier point, nous sommes en désaccord total. Sperber et Wilson répondent tout simplement que leur dispositif déductif n'est pas "une logique": il n'est qu'un "mécanisme computationnel" limité dans ses

opérations par des règles qu'il applique.

Or, y a-t-il quelque différence que ce soit entre une logique et un dispositif déductif, ou comme eux-mêmes préfèrent le dire, un “mécanisme computationnel”? “Mécanisme” est un mot très fort. “Mécanisme déductif (computationnel)” l’est encore plus. Que le mécanisme soit limité par les règles qu’il applique (et, en effet, le nombre de règles dans le mécanisme de Sperber et Wilson est bien limité car ils n’admettent que des règles d’élimination), cela ne pose aucun problème, vu qu’une logique peut, bien sûr, être limitée par les règles qu’elle applique¹¹, mais cela n’exclut pas du tout le fait qu’il existe une logique sous-jacente (une logique dont leur dispositif déductif ne peut se passer) et qui utilise des règles d’inférences de la logique classique comme les règles d’élimination du *et*, le *Modus Ponens* et le *Modus Tollens*.

Peut-être Sperber et Wilson ont-ils utilisé l’expression “mécanisme computationnel” exactement pour éviter l’expression “une logique” pour des raisons que nous présenterons plus loin; mais à vrai dire, la logique reste toujours là implicite dans l’expression “dispositif déductif”; en tout cas, nous insistons qu’un tel mécanisme est de fait une logique. Il suffit de réfléchir à

¹¹ En fait, la plupart des systèmes logiques ont deux, trois règles basiques: comme le *Modus Ponens* et la *Substitution* (c’est la règle qui préserve contradictions et validités, elle ne précise le vraie qu’en vertu de la forme : en substituant la même formule, qu’elle soit atomique ou non, à toutes les occurrences de la même lettre dans une tautologie, on obtient une autre tautologie), etc.

leurs définitions d'*implication logique non-triviale* (dans laquelle il est admis seulement les règles d'élimination que nous venons de mentionner) d'*implication contextuelle*, d'*implication analytique*, d'*implication synthétique*, de *correction*, de *complétude*, etc.

Or, comment peut-on réaliser des opérations dans un dispositif déductif qui ne tient compte que de règles d'inférences bien connues de la logique sans utiliser cette logique elle-même?

Et après tout, qu'est-ce qu'une logique? Ou plus précisément, qu'on pourrait comprendre par "*une* logique"?

Au moins en première approximation, ce serait bien de préciser tout d'abord qu'en général, toute logique peut être considérée comme une théorie (ou système) formelle à propos d'objets déterminés, et normalement présentés dans une langue artificielle que les logiciens construisent dans le but de rendre plus précises leurs analyses logiques et éviter ainsi des erreurs ou des paradoxes qui sont pratiquement inévitables lorsque le raisonnement est conduit dans une langue naturelle. Ces langues doivent tenir compte de deux dimensions à savoir, la dimension syntaxique qui traite des symboles et des signaux comme un jeu exclusivement formel gouverné par des règles de combinaison qui ne tiennent pas compte du sens de ces symboles; et la dimension sémantique, puisque les langages ne se limitent pas à des structures

uniquement syntaxiques - de fait, les symboles, les signaux, enfin les expressions ont un sens et ils se rapportent à des objets extra-linguistiques.

Pour ces motifs, dans un système de logique, il est possible d'analyser les raisonnements déductifs soit par la voie syntaxique soit par la voie sémantique. Comme on l'a déjà remarqué, la syntaxe de la logique s'occupe des propriétés et relations (par exemple, "une hypothèse Q est déductible d'une autre hypothèse P ") des expressions et formules où l'on fait abstraction des significations, tandis que la sémantique s'occupe des propriétés et relations des expressions qui tiennent compte des liens entre les expressions et les objets qu'elles désignent.

Cela étant dit, des notions comme *axiome*, *théorème*, *dérivation*, *démonstration* (ou *preuve*), *consistance syntaxique*, *inconsistance syntaxique* et *règle de transformation* (ou règle d'inférence) font partie de la syntaxe de la logique.

Supposons par exemple, une logique L contenant la seule règle d'élimination de *si ... alors* (*modus ponens*). Comme on l'a déjà remarqué auparavant, une hypothèse Q est une implication logique (bien entendu, au sens syntaxique) de l'hypothèse P dans L si, et seulement si, Q est déduite de P par l'application des règles déductives de L (dans ce cas, *modus ponens*), ou en d'autres mots, si et seulement si, il existe une démonstration

formelle dans L de Q à partir de P .

Pour vérifier si une démonstration formelle dans L est correcte (c'est-à-dire, conforme aux règles), il suffit de s'assurer que le passage de chaque ligne de la démonstration respecte les règles de manipulation des symboles de L (il ne faut pas *interpréter* ces symboles, tout se maintient au niveau syntaxique). Dans ce cas, étant donné les hypothèses (q) et (r) du langage naturel ci-dessous, (q) implique logiquement (r).

- (q) i) Si 4 est un chiffre pair, 4 est divisible par 2.
- ii) 4 est un chiffre pair.
- (r) 4 est divisible par 2.

Cette implication logique se justifie tout simplement par la règle d'élimination de *si... alors* de L , laquelle est appliquée aux hypothèses (q) et (r) en vertu uniquement de leur forme logique tout en faisant abstraction de leurs propriétés sémantiques.

D'un autre côté, des notions comme *dénotation*, *extension*, *connotation*, *intension*, *validité*, *modèle*, *interprétation*, *satisfiabilité* (*consistance sémantique*), *inconsistance sémantique*, *vérité*, *conséquence logique*, *conséquence nécessaire (entailment)* et *implication matérielle*, font partie de la sémantique de la logique.

Dans le Chapitre I, nous avons vu que la relation de conséquence nécessaire est une relation sémantique: une hypothèse Q est conséquence nécessaire de l'hypothèse P si et seulement si, dans tout monde possible où P est vraie, Q est également vraie. Dans ce cas, il n'y a pas de monde possible qui rende vraie P et Q fausse. Ainsi, l'interprétation sémantique des hypothèses P et Q dépend toujours de tels états de choses, comme nous pouvons le voir dans l'exemple suivant où (t) est une conséquence nécessaire de (s).

- (s) La lune est un satellite de la terre et le soleil est une étoile.
- (t) La lune est un satellite de la terre.

Ainsi, dans un raisonnement déductif, cette relation de conséquence nécessaire se justifie par le fait que tout état de choses dans lequel (s) est vraie, vérifie également (t).

Revenons maintenant aux langues artificielles construites par les logiciens pour tenter de préciser ce qu'on pourrait comprendre par "une logique". Bien, une logique (ou système formel) est formée par une syntaxe et/ou une sémantique qui prend en considération ce qui suit: en ce qui concerne sa dimension syntaxique, "une logique L " doit avoir en premier lieu, un langage

composé premièrement d'un vocabulaire, à savoir, des symboles primitifs de L comme par exemple, $\sim$, $\rightarrow$, $(,)$ et des lettres propositionnelles comme A_1 , A_2 , $\dots$, et deuxièmement, d'une grammaire composée de règles de formation du type: toute lettre propositionnelle est une *formule bien formée*; si A et B sont formules bien formées alors, $(\sim A)$ et $(A \rightarrow B)$ sont des formules bien formées, etc. En second lieu, "une logique L " doit présenter quelques Postulats (ou *schémas d'axiomes*) comme par exemple, $A \rightarrow (B \rightarrow A)$ et $(A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$ (quand L n'a pas d'axiomes, L est dite être non axiomatique), et avoir en plus des règles de transformation comme la règle *Modus Ponens*. Pour ce qui est de la dimension sémantique, "une logique L " explique la signification des symboles et formules bien formées de L en définissant la structure d'une interprétation possible (ou d'un modèle) pour la langue objet.

La plupart des logiciens s'intéressent surtout à certaines propriétés et relations qui existent entre les formules déductibles dans L et les formules vraies de L . Il s'agit ici de la *solidité* et de la *complétude forte* de L : une logique L est *solide* si toute implication logique dans L est en même temps une conséquence nécessaire dans L . Inversement, une telle logique est *complète* si toute conséquence nécessaire dans L est une implication logique dans L .

D'un point de vue purement logique, pour bien définir la solidité et la complétude forte d'un système de logique, il faut faire appel à la notion de *conséquence logique* qui, à son tour, repose sur la notion de *vérité dans un modèle* (relativement à une circonstance) laquelle fait appel à la notion de *satisfaction*.

Étant donné un système de logique axiomatique formalisé L , une formule A de L est vraie dans un certain modèle quand elle est satisfaite (... est vraie de...) par toutes les suites d'individus qui appartiennent à l'univers du discours donné. Elle est une *conséquence logique* d'un ensemble quelconque de formules quand tous les modèles qui rendent les formules de cet ensemble vraies, rendent également A vraie.

Nous avons vu précédemment que dans leur dispositif déductif, Sperber et Wilson admettent qu'il existe un ensemble de règles déductives qui sont spontanément mises en oeuvre dans le traitement déductif de l'information. Comme on le sait déjà, ces règles déductives (qu'ils empruntent à la logique) ne sont que des règles d'élimination. Précisons que selon Sperber et Wilson, c'est sur de telles règles que se fonde l'habileté humaine à faire des inférences démonstratives spontanées.

Dans nos remarques précédentes nous avons souligné une certaine divergence de notre part par rapport à la réponse de Sperber et Wilson à

Hinkelman lorsqu'ils soutiennent que leur dispositif n'est pas "une logique". D'ailleurs, dans l'ensemble de leur réponse, nous voyons que, Sperber et Wilson n'expliquent malheureusement ni l'aspect logique de la question concernant le premier point, ni l'aspect cognitif concernant le deuxième point. Or, si Sperber et Wilson ont raison, c'est-à-dire, si leur dispositif déductif n'est pas une logique, nous posons donc, les questions suivantes:

Par quelle procédure pourrait-on manipuler de telles règles déductives - lesquelles, selon eux, ne sont que des opérations qui prennent en considération seulement les propriétés sémantiques qui sont reflétées dans les formes des hypothèses? (1989, p. 134). On doit tenir compte uniquement des propriétés sémantiques qui ne sont que "reflétées" dans les formes des hypothèses? Mais, quelles formes des hypothèses? Leurs formes logiques?

Si oui, est-il possible d'"abstraire" les formes logiques des hypothèses en faveur uniquement des propriétés sémantiques qui s'y "reflètent"? Remarquons que nous avons souligné plus haut que du point de vue logique, une forme logique d'une hypothèse est la structure profonde qui sert à déterminer non seulement les conditions de vérité mais aussi toutes les inférences qu'on peut obtenir à partir d'une telle hypothèse.

D'autre part, même si la fonction du dispositif déductif est d'analyser et de manipuler le contenu conceptuel des hypothèses, est-il possible

d'appliquer ces règles déductives de la logique dans le dispositif déductif sans utiliser la logique elle-même?

Qui plus est, pour définir la notion de règle déductive d'un système de logique, nous avons besoin de la notion de conséquence nécessaire. Cette définition repose, à son tour, sur celle de *vérité*, laquelle présuppose l'existence d'états de choses qui correspondent aux hypothèses. Ainsi, une déduction n'est pas tout simplement un processus computationnel, comme l'insinue la réponse de Sperber et Wilson, mais elle préserve la vérité tout en y maintenant que la conclusion est toujours une conséquence nécessaire des prémisses. D'ailleurs, comme Sperber et Wilson eux-mêmes le reconnaissent (1989: 132), c'est exactement la préservation de la vérité qui distingue une déduction des autres computations. Donc, nous pouvons conclure que toutes les implications logiques sont aussi des conséquences nécessaires, et du fait que cette relation existant entre implication logique (formules déductibles) et conséquence nécessaire (formules vraies) est préservée dans le dispositif déductif de Sperber et Wilson, nous pouvons conclure que tel dispositif déductif est une logique présentée au niveau syntaxique et dans laquelle est préservée la relation de solidité.

L'inverse de cette relation n'est pas toujours vrai: comme le montre l'exemple de Sperber et Wilson, leur système formel n'a pas de règle

déductive qui permette la dérivation ci-dessus:

- (u) La neige est blanche.
- (v) L'herbe est verte.
- (x) Donc, la neige est blanche et l'herbe est verte.

Dans ce cas, nous n'avons pas que toute conséquence nécessaire est une implication logique dans leur système formel. Dans ce cas, tel dispositif n'est pas complet (comme Sperber et Wilson le soutiennent dans 1989, p. 158).

À cet égard, ce qui est plus intéressant pour la plupart des logiciens, c'est de construire des systèmes formels qui soient complets. Pour ce faire, il faut tout simplement postuler des règles déductives qui rendent possible, à partir d'un ensemble d'hypothèses, la dérivation de toutes les conséquences nécessaires concernant les implications logiques. Ainsi, du fait que le dispositif déductif humain proposé par Sperber et Wilson n'est pas complet, leurs règles ne sont pas capables de dériver les conclusions de certains arguments valides que nous faisons quotidiennement.

Malgré le fait qu'ils soutiennent que leur objectif est de construire non pas un système logique optimal, mais une sorte de système de déduction formel qui puisse constituer le modèle pour un système cognitif (qu'eux-mêmes

reconnaissent incapable de précision), il y a d'après nous une logique dans ce modèle - une logique qui sert aussi bien de fondement pour certaines de leurs systématisations rationnelles, une logique n'étant qu'un ensemble de canons d'inférence associé à un certain système de catégories par rapport auquel cette logique détermine les inférences valides. Dans leur dispositif déductif, Sperber et Wilson défendent fortement les raisonnements valides, mais il faut tenir compte du fait que là où il y a des raisonnements valides, il doit exister une logique. Or, la logique a été créée exactement pour analyser et expliquer les raisonnements valides que nous faisons spontanément dans la vie de tous les jours aussi bien que dans certains contextes plus sophistiqués.

Tout bien considéré nous aimerions observer que s'il s'agit d'une stratégie de la part de Sperber et Wilson de choisir l'expression "mécanisme computationnel" pour éviter l'autre "une logique", c'est d'après nous, une stratégie rhétorique qui peut avoir pour but d'éviter certains compromis de la part de ces deux auteurs quant aux aspects formels, techniques dont il faudrait tenir compte s'il leur fallait reconnaître leur dispositif déductif comme étant une logique. Et à vrai dire, ces aspects-là ne sont pas du tout faciles à analyser dans une théorie qui est en même temps computationnelle, inférentielle, pragmatique et cognitive (et qui traite en particulier du problème de la pertinence pragmatique dans la conversation). En tout cas, s'il s'agit d'une

stratégie rhétorique utilisée par Sperber et Wilson dans le but que nous venons de mentionner, à notre avis elle n'est pas très heureuse puisqu'un mécanisme déductif doit être associé à l'idée d'une logique.

Par contre, étant donné le caractère intrinsèque d'un tel dispositif déductif pour expliquer le mécanisme général du traitement de l'information, nous comprenons la grande difficulté qu'il y aurait à le définir optimalement, soit de façon formelle, rigoureuse, précise. Pour citer quelques-unes de ces difficultés, nous remarquons premièrement que le processus d'interprétation chez Sperber et Wilson est non démonstratif; chaque individu qui traite des informations a ses entrées encyclopédiques; ainsi, le savoir encyclopédique peut beaucoup varier d'un individu à l'autre; la connaissance des conventions de la langue change d'un individu à l'autre; les conséquences qu'un individu infère d'une hypothèse peuvent varier par rapport aux conséquences qu'un autre individu en infère, car l'interprétation peut différer en dépendant de l'entrée disponible à chaque individu, c'est-à-dire, si les entrées diffèrent, les individus n'en tireront pas les mêmes conséquences. Et nous comprenons que la pragmatique et la théorie de la communication et de la cognition ne peuvent pas prendre en considération de telles divergences car elles peuvent être si considérables qu'aucune théorie ne pourrait jamais en rendre compte. Dans la perspective de Sperber et Wilson ce qu'une théorie peut tenter d'expliquer ou

de traiter est justement le mécanisme général du traitement de l'information. Pour ces motifs, peut-être leur dispositif est loin d'être un système logique optimal, au sens qu'il lui manque la précision nécessaire pour expliquer formellement une théorie de la communication et de la cognition.

Malgré tout, on trouve dans Sperber et Wilson (1989), plus précisément dans le chapitre sur l'inférence, des commentaires en faveur de la psychologie cognitive et de la théorie pragmatique du type suivant:

La question de savoir si l'équipement mental de base des humains comporte des règles logiques - et, le cas échéant, lesquelles - n'a pas grand intérêt pour les purs logiciens. Les logiciens s'intéressent à la nature de tous les systèmes déductifs concevables, qu'ils aient ou non une réalité psychologique. Par contre, cette question est d'un intérêt majeur pour la psychologie cognitive, et pour la théorie pragmatique en particulier. (1989, p. 133)

En effet, les logiciens s'intéressent à la nature de tous les systèmes déductifs concevables mais tout en tenant compte des critères rigoureux de formation spécifique d'un domaine particulier:

On ne choisit pas une logique comme on ferait un jouet; on parvient à une logique par l'étude des conceptualisations et inférences caractéristiques d'un domaine particulier (physique: logique quantique; métaphysique: calculs modaux; mathématique: logique classique, resp. intuitionniste, resp. constructiviste). Un formalisme quelconque, une matrice ou un système de matrices de vérité comme ceux des connecteurs multivalués, n'est pas une logique; il devient une logique seulement par une interprétation qui ne fait pas abstraction du sens logique (si celui-ci existe). Cela suffit à dissiper les craintes, exprimées ici ou là, que la logique contemporaine consisterait en

“agences de symboles” dénués de sens (bref serait *nominaliste*, suivant la nomenclature commune). Généralement d’ailleurs l’interprétation logique préexiste au formalisme. Là-dessus les conceptions de Frege semblent plus raisonnables que certaines vues formalistes. (Largeault, 1993, p. 24)

On convient aujourd’hui de bien préciser la distinction entre *logique formelle* qui traite de la forme ou structure des raisonnements et *logique matérielle* qui tient compte des contenus, en abordant des problèmes concernant la classification, le raisonnement par analogie et par induction, etc.; la *logique matérielle* est considérée comme une application de la logique à des problèmes concrets, et malgré l’intérêt des logiciens classiques pour la forme, ils soutiennent qu’il y a un lien entre la logique et la réalité en supposant un isomorphisme entre ce qui est la forme purement logique et l’organisation du monde.

À partir de 1940, en théorie des modèles, s’est développée l’étude systématique de la construction de modèles. Telle théorie s’occupait de la sémantique en étudiant les rapports syntaxiques (de forme) et des propriétés en rapport avec une interprétation.

Les logiciens, depuis 1920, cherchent des algorithmes: un algorithme est une procédure mécanique qui est donné par un ensemble fini d’instructions et qui conduit à un résultat unique; il s’agit d’un objet mathématique qui produit des enchaînements de signes capables de subir

certaines interprétations conceptuelles ou propositionnelles. Cependant, ils ne se préoccupent pas seulement d'en trouver, mais de bien examiner les performances et même les limites soit théorique soit pratique, au-delà desquelles les algorithmes ne peuvent aller ou s'étendre. En ne les trouvant pas, cela sert quand même à démontrer qu'un certain ensemble de problèmes n'en comprend pas - mais telle impossibilité de solution par algorithme n'a pas pour conséquence l'inexistence d'une solution pour ces problèmes.

Soit dit en passant, il existe des systèmes formels dans la logique classique pour lesquels le problème de la décision est récursivement insoluble. Dans ce cas, il n'existe pas d'algorithme pour décider si une formule quelconque du système en question est ou non un théorème. Le Calcul de Prédicats Classiques et l'Arithmétique élémentaire (cf. théorème de Gödel¹²) par exemple, ne sont pas décidables. D'un autre côté, il existe des systèmes formels dans la logique classique pour lesquels le problème de la décision est récursivement soluble. Par exemple, le Calcul Propositionnel Classique est capable de décider si une formule quelconque d'un certain système est ou non un théorème; l'algorithme est donné par les tables de vérité.

¹² Les premier et deuxième théorèmes d'incomplétude de Gödel ont été beaucoup étudiés depuis leurs preuves à l'année 30. Le premier théorème affirme l'incomplétude de l'arithmétique élémentaire de premier ordre, si celle-ci est consistante, tandis que le deuxième théorème affirme l'impossibilité de la preuve de la consistance de l'arithmétique élémentaire.

À la suite de toutes ces constatations, revenons maintenant à la réponse de Sperber et Wilson à Hinkelman concernant le deuxième point que nous avons indiqué plus haut. Comme nous l'avons vu, Sperber et Wilson soutiennent dans leur réponse que leur mécanisme computationnel est limité dans ses opérations non seulement par les règles qu'il applique, mais aussi par la manière selon laquelle il les applique.

La manière? Que veulent-ils dire par "manière d'appliquer une règle"? Que nous n'appliquons pas toujours la règle? Ou que nous appliquons la règle du mode qu'on veut?

Mais si tel était le cas, ce ne serait plus un mécanisme! *Mécanisme* veut dire régularité, et on ne peut pas utiliser "mécanisme" pour désigner quelque chose d'irrégulier. Si appliquer une règle est une action, il existe un axiome basique de la Théorie de l'Action qui dit: Il existe diverses manières de réaliser une même action. Par exemple, on peut saisir un verre soit avec la main gauche soit avec la main droite. Mais, nous ne croyons pas que c'est ce sens là que Sperber et Wilson ont à l'esprit.

Comment donc sommes-nous limités par ces "manières"? Cela semble indiquer une bonne dose de psychologisme. S'agit-il d'une hypothèse empirique qui n'a aucune justification transcendantale, une hypothèse anthropologique sur la manière selon laquelle les gens manipulent, traitent des

informations? On pourrait concevoir la logique comme étant un mécanisme déductif, comme quelque chose de transcendantal à la Kant, mais nous ne croyons pas que ce soit le cas. Apparemment cela semble vouloir dire qu'on pourrait même changer de logique, de mécanisme inférentiel graduellement, à chaque siècle sans nous en rendre compte.

D'un autre côté, nous reconnaissons la nécessité d'explorer des possibilités de développement d'autres logiques au lieu d'admettre des hypothèses anthropologiques, psychologiques, etc., pour justifier des limitations qui s'imposent dans n'importe quel cadre théorique visant à expliquer la pertinence.

Rappelons qu'autrefois Stuart Mill a eu l'idée d'une logique complètement inductive. Par la suite, Carnap et Jeffrey ont élaboré une logique inductive. Contrairement à Stuart Mill, la logique de Frege est une logique faisant appel uniquement à des procédures effectives. Son idée était de contrôler, fermer au maximum le langage, et en ce sens de reconstruire le langage comme un calcul. Au début, Wittgenstein a donné un air transcendantal à la logique classique: la logique serait quelque chose de transcendantal et cela se montrerait en particulier dans le fait qu'on ne peut pas créer des nouvelles formes logiques. Elles sont données d'avance, et c'est ça qui permet la médiation entre le langage, la pensée et le monde. Avec

l'évolution de la logique, le développement des logiques non classiques, le développement de l'empirisme de Quine, l'idée de que la logique est quelque chose de transcendantal fut mise de côté. Dès lors, le chemin est ouvert pour explorer d'autres possibilités.

La pluralité des logiques indiquerait que la neutralité relativement au sujet (l'étoffe des raisonnements change, les règles d'inférence demeurent, ou bien la logique classique sert à n'importe quoi et à tout le monde, remarque de Quine) vaut tant que l'on considère chaque logique en elle-même, non pas si on en considère plusieurs; cette pluralité reflète la diversité d'aspects des langues et la multiplicité des langages théoriques. Tant qu'il y eut un langage unique de la science et de la vie quotidienne, le phénomène ne se présentait pas. (Largeault, 1993, p. 24)

Les techniques de la logique se sont développées beaucoup à partir des années trente. Actuellement on peut développer des logiques pour rendre compte de certaines pratiques de raisonnement. Rappelons encore que là où il existe des raisonnements valides, il doit exister une logique pour les expliquer. Aristote a développé la Syllogistique, mais il a reconnu qu'il existe des raisonnements valides dont sa logique ne rend pas compte. De nouvelles logiques furent donc créées pour rendre compte de ces raisonnements logiques.

Et c'est là peut-être une erreur de Sperber et Wilson: pour bâtir leur Théorie de la Pertinence comme modèle du dispositif déductif humain, Sperber et Wilson l'ont pensé en termes de logique classique. Ensuite ils y

introduisent leurs définitions techniques et les règles déductives d'élimination. En fait les règles d'introduction de la logique classique poseraient des problèmes dans un système formel dans la mesure où elles seraient appliquées *ad infinitum* à son *output*, produisant ainsi une dérivation sans fin. Mais tout cela pourrait être évité s'ils introduisaient des définitions techniques visant à éviter les implications triviales en adoptant une logique non classique.

À tout prendre, pour ce qui est de notre intérêt, c'est-à-dire pour expliquer les jugements de pertinence, nous croyons que c'est d'une logique de la pertinence, bien sûr, dont il faut tenir compte: une logique qui n'infère pas n'importe quoi, une logique finitaire, limitée, enfin une logique à partir de laquelle nous pouvons discuter des problèmes effectifs sans faire appel à des hypothèses anthropologiques.

Notre objectif est donc de proposer dans le chapitre III quelques définitions de la pertinence qui utilisent des notions centrales telles que *conséquence pragmatique*, *conséquence triviale*, *consistance*. Là aussi nous proposons en plus une logique de la pertinence qu'on appelle "*P*". Nous croyons que ces notions pourront contribuer à éclairer le problème de l'implication non triviale chez Sperber et Wilson ainsi que la définition de pertinence elle-même en ce qui concerne plus précisément les degrés de pertinence. Dans le chapitre IV nous ferons une analyse des problèmes que

nous venons de présenter en les confrontant avec les définitions que nous proposons dans le Chapitre III.

CHAPITRE III

DÉFINITIONS DE LA PERTINENCE ET LA LOGIQUE DE LA PERTINENCE *P*

Introduction

Notre objectif dans ce chapitre est de présenter quelques définitions de la pertinence qui puissent éclairer la notion d'implication pertinente, c'est-à-dire, une implication dans laquelle l'antécédent suffit *pertinemment* au conséquent. Ensuite, nous proposons une logique pertinente qui pourra être utilisée comme un outil pour analyser le problème suivant: quand un ensemble d'énoncés Γ est-il pertinent pour un autre ensemble d'énoncés Δ ? Dans notre analyse, ce problème est abordé du point de vue pragmatique (il ne s'agit pas effectivement d'un problème de logique du point de vue syntaxique et sémantique) car nous prenons en considération d'autres aspects des propositions au-delà de leur seule valeur de vérité. Nous considérons alors le contexte comme une hypothèse décisive pour établir quand

un énoncé A est pertinent pour un autre énoncé B .

1- Considérations historiques

Depuis les critiques dirigées par Lewis (1912) à la formalisation de propositions du type ' $\rightarrow$ ', les formules $A \rightarrow (B \rightarrow A)$ et $\sim A \rightarrow (A \rightarrow B)$ ont été connues comme les “paradoxes” de l'implication matérielle de la logique classique. Du point de vue purement classique, le Calcul Propositionnel nous dit que l'interprétation formelle de la proposition “Si les poissons vivent hors de l'eau alors les triangles ont trois côtés” est telle que cette proposition est vraie lorsque l'antécédent est faux ou le conséquent vrai. Dans ce cas, les seules propriétés d'une proposition qui importent à la logique sont sa forme et sa valeur de vérité - et cette dernière, selon l'hypothèse frégréenne, est déterminée par la forme de la proposition et par les propriétés sémantiques de leurs constituants.

Selon les “logiciens de la pertinence”, toute logique doit renoncer à ces thèses paradoxales et considérer l'usage de ' $\rightarrow$ ' comme un connecteur pertinent. Du point de vue de la logique classique, ' $\rightarrow$ ' capture seulement les aspects du “si ... alors ...” qui dépendent de la valeur de vérité des

propositions. Selon la logique de la pertinence, ‘ $\rightarrow$ ’ capture les aspects qui déterminent la validité de formules du type “Si A alors B ” *seulement* quand A est pertinente pour B , ou en d’autres mots, lorsqu’il existe une relation de pertinence entre A et B .

Mais que signifie “relation de pertinence” entre l’antécédent A et le conséquent B d’une proposition comme “Si les poissons vivent hors de l’eau, alors les triangles ont trois côtés”? En principe, aucune proposition analogue à celle-ci ne devrait être une instance de substitution vraie dans aucune logique.

Intuitivement, il semble qu’une relation de pertinence devrait s’établir par une connexion de signification (*meaning*) entre A et B . Par contre, comme l’objectent Hugues et Cresswell (1968) qui défendent la logique modale classique, et Bennett (1969), la notion d’une telle connexion ne peut pas être incorporée dans une analyse logique de “si ...alors ...”, et, conséquemment, elle ne permet pas non plus de déterminer si un système formel est ou non une logique correcte de l’implication (*entailment*). Ils défendent la logique classique et la logique modale classique dans la mesure où elles utilisent des notions bien précises telles que vérité, fausseté, possibilité et impossibilité dans le traitement formel des propositions atomiques.

D’après les logiciens de la pertinence, la logique classique est

paradoxale; la préfixation d'hypothèses dans une formule de la forme $A \rightarrow (B \rightarrow A)$ est paradoxale car dans le conséquent de cette formule il n'existe aucune pertinence de B par rapport à A , c'est-à-dire, en affirmant une proposition valide A , une quelconque proposition B serait pertinente pour A - ce que nous allons contester pragmatiquement pour la simple raison que dans la vie quotidienne on ne traite pas les propositions de cette forme sans prendre en considération d'autres aspects, et non pas seulement les valeurs de vérité des propositions atomiques. De même, $A \rightarrow (\sim A \rightarrow B)$ (et toutes ses variantes) est paradoxale car en affirmant une thèse, la négation de cette thèse implique une formule quelconque.

Comparables à ces paradoxes sont les paradoxes de l'implication stricte: $A \rightarrow (B \rightarrow B)$ et $(A \wedge \sim A) \rightarrow B$ (Duns Scot). $(A \wedge \sim A) \rightarrow B$ est paradoxale car d'une contradiction on peut dériver une formule B quelconque qui n'a aucune relation avec cette contradiction - ce qu'on peut aussi contester puisque il n'y a rien dans l'antécédent qui justifie le conséquent. Pour éviter des paradoxes de ce type, on a proposé les logiques de la pertinence, dont l'exigence principale est que le contenu des implications, c'est-à-dire, que le contenu exprimé par les propositions (antécédentes) ait une relation sémantique avec le contenu exprimé par les autres propositions (conséquentes). En d'autres termes, l'antécédent d'une implication doit être pertinent pour son

conséquent.

Dans l'approche de R. Epstein (1995) deux propositions sont apparentées si elles partagent, ou explicitement ou implicitement, quelque objet en commun. Il retient l'hypothèse qu'une proposition peut être considérée comme apparentée par sujet (*subject matter*) à quelques propositions et non apparentées à d'autres. En paraphrasant quelques-uns de ses exemples, nous voyons que "Ralph est un chien" et que "Raph aboie" partagent une référence en commun. "Don est célibataire" et "Don est un homme" partagent implicitement un prédicat en commun. Par contre, en prenant la proposition "Si la lune est faite de fromage vert alors $2 + 2 = 4$ ", on peut affirmer sans risque d'erreur que "la lune est faite de fromage vert" n'est pas apparentée à la proposition " $2 + 2 = 4$ ". Epstein construit une logique qu'il appelle *Relatedness Logic: The Subject Matter of a Proposition* où une proposition conditionnelle du type "si ... alors ..." est vraie seulement si l'antécédent est pertinent pour son conséquent, c'est-à-dire, si les deux propositions ont un sujet (matière) en commun. Selon Epstein, le contenu est ainsi considéré comme un "aspect" et non pas comme une propriété des propositions. Dans ces exemples, le sujet d'une proposition peut affecter la vérité d'une proposition de la forme "si ... alors ...", dans le cas où la proposition antécédente n'est pas apparentée à la proposition conséquente.

Considérons maintenant “Ralph est un chien” comme étant vraie. Dans ce cas, la valeur de vérité est une propriété de cette proposition, alors que “Ralph est un chien” est un “aspect” de cette proposition qui a des relations avec d’autres propositions comme, par exemple, avec la proposition “Les chiens mangent de la viande”. Dans une discussion au sujet de créatures qui vivent dans la mer, les propositions “L’orque est une baleine” et “Les requins sont dangereux” sont apparentées. Dans le cas où la discussion porte sur les mammifères, peut-être prendrait-on les deux propositions comme étant non apparentées. La proposition “ $2 + 2 = 4$ et la lune est faite de fromage vert” a un sujet en commun avec les propositions “la lune est faite de fromage vert” et “ $2 + 2 = 4$ ”. “Ralph est un chien et $2 + 2 = 4$ ” est apparentée à la proposition “Si Ralph est un chien alors $2 + 2 = 4$ ”. De même, “Ralph est un chien” est apparentée à “Ralph n’est pas un chien”. Ceci dit, l’on a que le sujet (considéré comme un “aspect”) d’une proposition est indépendant de sa valeur de vérité et, conséquemment, les connecteurs logiques sont neutres par rapport au sujet d’une proposition, ils n’ont pas de contenu, ils sont donc *syncatégorématiques*: lorsque l’on affirme que “Ralph est un chien” n’est pas apparentée à “ $2 + 2 = 4$ ”, ce n’est pas la vérité de “Ralph est un chien” qui importe. Alors, selon cette approche les tautologies ont des “sujets”.

Nous voulons contester l’approche d’Epstein car il ne nous offre

aucune définition de la pertinence; pour nous, la notion de sujet d'une proposition n'est pas suffisant pour expliquer la notion d'implication pertinente. Considérons une proposition telle que “Si la lune est faite de fromage vert, alors la lune est ronde” où les deux propositions constituantes ont un sujet (matière) en commun; comme on peut le voir, il manque encore une relation de pertinence pour justifier l'implication entre les deux propositions.

Tout bien considéré, notre motivation n'est donc pas, comme celle des logiciens de la pertinence, de trouver une logique qui soit rivale de la logique classique, mais de présenter une logique qui puisse être utilisée comme un outil dans l'analyse de plusieurs jugements de pertinence. Puisque notre discussion porte sur un problème d'ordre pragmatique, cela nous amène à prendre comme point de départ une logique de la pertinence car d'après nous elle est, à l'exemple des logiques intuitionnistes, effectivement plus proche du langage commun que quelconque autre logique que nous pourrions utiliser.

Avant de proposer notre logique de la pertinence ***P***, nous soulignons deux points importants.

Le premier point concerne ce que nous avons déjà remarqué dès l'introduction de notre thèse: pour élaborer nos définitions techniques de la pertinence, nous avons eu besoin de notions de “proposition”, “énoncé”, “hypothèse” qui sont fondamentales dans notre discussion. Nous allons préciser

ces notions en faisant appel à la Théorie des Actes de Discours (T.A.D.) de Searle et Vanderveken laquelle contient, entre autres choses, une sémantique formelle générale des langues naturelles pour rendre compte de la signification des énoncés.¹³ Notre but est de les intégrer dans nos définitions techniques de la pertinence (comme nous le verrons par la suite) car, d'après nous, la sémantique formelle est la plus riche et la plus proche du langage commun que quelconque autre que nous pourrions considérer pour discuter de la pertinence concernant les énonciations des interlocuteurs.

Qui plus est, la T.A.D. (laquelle sera présentée au chapitre V) nous sera très utile dans la justification des jugements de pertinence présentée dans le chapitre VI, quand nous montrerons comment notre définition de la pertinence permet d'expliquer les jugements de pertinence que nous faisons dans la vie de tous les jours.

Pour Vanderveken,

Le langage exerce une fonction médiatrice essentielle dans l'expression des pensées. Toute pensée conçue par un sujet humain est en principe exprimable par les moyens de son langage lors de l'accomplissement d'actes de discours. Aussi une sémantique logico-philosophique décrivant avec exactitude les universaux sémantiques relatifs à l'usage du langage est-elle transcendente en ce sens qu'elle articule les formes *a priori* de la pensée et du monde et qu'elle fixe des limites à ce qui peut être pensé et à ce qui peut exister et être

¹³ Comme nous le verrons dans les chapitres V et VI, en plus de la sémantique générale, la T.A.D. intègre aussi une logique illocutoire qui sert à rendre compte des aspects illocutoires de la signification, et une logique propositionnelle pour rendre compte des aspects vériconditionnels des propositions.

objet d'expérience. Les actes illocutoires sont les unités premières de signification dans l'usage des langues naturelles, parce qu'ils sont aussi des unités premières de pensée conceptuelle. (Vanderveken, 1988, p. 10)

Donc, notre intérêt pour la sémantique générale se doit au fait que dans son approche sémantique, Vanderveken associe la compétence linguistique d'un locuteur à sa capacité "d'accomplir des actes illocutoires et de comprendre quels actes illocutoires peuvent être accomplis par d'autres locuteurs dans les contextes actuels et possibles d'emploi de sa langue" (1988, p. 19).

Comme on le sait, les propositions introduites par Frege étaient limitées aux *sens d'énoncés déclaratifs* et *d'énoncés interrogatifs complets* (c'est-à-dire, des énoncés interrogatifs admettant une réponse par "oui" ou par "non"): une proposition était identifiée comme le sens commun à différents énoncés ayant les mêmes conditions de vérité. Ainsi, les énoncés "Jean coupe l'arbre" et "L'arbre est coupé par Jean" expriment dans chaque contexte possible d'emploi la même proposition. Comme on peut le voir, le critère d'identité propositionnelle établi par Frege se limite seulement à la stricte équivalence ou identité des conditions de vérité des propositions.

Ensuite, Searle (1969) et Strawson (1971) ont introduit les propositions comme contenus communs à des actes illocutoires différents ayant

les mêmes conditions de satisfaction. L’assertion que “Marie part en vacances ce matin” et la conjecture que “Marie part en vacances ce matin” sont deux actes illocutoires différents mais ils ont le même contenu propositionnel.

Alors, étant donné l’importance des actes illocutoires complets comme unités principales de signification des énoncés dans l’usage et la compréhension des langues naturelles, Vanderveken (1988) a proposé d’unifier ces deux conceptions des propositions dans sa sémantique générale car, selon lui, elles sont complémentaires:

En effet, c’est en vertu de leur signification que les énoncés élémentaires servent à accomplir des actes illocutoires et le *sens* d’un énoncé élémentaire (non ambigu) dans un contexte est précisément la *proposition* qui est le *contenu* de l’acte illocutoire qu’il exprime dans ce contexte. (Vanderveken, 1988, p. 82)

De cette façon, dans sa sémantique générale Vanderveken ne néglige ni les aspects véri-conditionnels ni les aspects illocutoires de la signification. Comme lui-même le remarque, chaque énoncé élémentaire qui exprime un acte illocutoire de la forme $F(P)$ dans un contexte possible d’énonciation a la proposition P pour sens dans ce contexte (indépendant de son type syntaxique et de la force illocutoire de ses énonciations).

En bref, dans la sémantique générale, les “propositions”, les “énoncés” et les “hypothèses” ne sont pas la même chose; une proposition est

considérée à la fois comme le sens commun à des énoncés de différents types syntaxiques et le contenu propositionnel commun à des actes illocutoires ayant différentes forces (elles représentent des faits). Les énoncés sont les unités syntaxiques de base pour l'accomplissement d'acte illocutoire dans la conversation. À chaque fois, quand le locuteur utilise une langue naturelle pour exprimer sa pensée, il utilise un énoncé. Il profère les sons ou il écrit les signes des mots d'un énoncé d'un certain type. La signification des mots réside (comme Frege l'a montré) seulement dans leur contribution à la signification des énoncés où ils figurent. Les énoncés déclaratifs servent à accomplir des actes illocutoires assertifs. Ainsi, la signification linguistique d'un énoncé élémentaire $f(p)$ est une fonction qui associe à chaque contexte possible d'énonciation l'acte illocutoire de la forme $F(P)$ que le locuteur tenterait d'accomplir s'il utilisait cet énoncé dans ce contexte en parlant littéralement. Cet acte illocutoire $F(P)$ est la signification particulière de cet énoncé dans ce contexte. Il a une force F et un contenu propositionnel P qui est aussi, comme nous venons de le remarquer, le sens de cet énoncé dans ce contexte. Donc, la signification d'un énoncé élémentaire en un contexte, contient une force et un contenu propositionnel (qui est le sens d'un énoncé). Quant aux hypothèses et aux informations, elles sont des actes illocutoires assertifs. Faire une hypothèse c'est faire une assertion faible d'une proposition avec la

condition préparatoire que même si cette proposition n'est pas certaine, elle est raisonnable (des raisons peuvent être données pour l'appuyer) et qu'elle est utile à considérer. Donner une information c'est **affirmer** à un interlocuteur une proposition P avec l'intention perlocutoire de lui faire savoir qu'elle est vraie (mode particulier d'atteinte du but assertif). Il y a donc la condition préparatoire que l'interlocuteur n'est pas déjà au courant de P (ne sait pas déjà que cette proposition est vraie)¹⁴.

Le deuxième point que nous avons indiqué en haut concerne la théorie de la pertinence de Sperber et Wilson. Nous soulignons que pour construire notre définition de la pertinence, nous nous inspirons de quelques définitions de leur théorie, notamment la définition d'implication contextuelle que privilégie les inférences nouvelles faites par des interlocuteurs engagés dans une conversation. Nous remarquons pour souligner une certaine relation existant entre notre approche et l'approche de Sperber et Wilson.

¹⁴ Un discours informatif est un discours du type descriptible où on donne des informations de ce type-là.

2- Définitions

Nous considérons la *pertinence* comme une relation binaire entre deux ensembles d'énoncés dans un même contexte. Dans notre logique de la pertinence il n'y a que des énoncés déclaratifs qui servent à faire des actes assertifs (assertions, hypothèses, informations). Selon notre sémantique classique, chaque énoncé exprime dans chaque contexte une proposition qui y est vraie quand le fait qu'elle représente existe dans ce contexte. Dans chaque contexte C , il y a un ensemble fini Γ d'énoncés utilisés par le locuteur qui représentent les assertions que le locuteur fait dans ce contexte. Il y a aussi l'ensemble de toutes les hypothèses que l'interlocuteur fait dans ce même contexte C (qui peuvent être anciennes ou nouvelles). Plus précisément, nous définissons la relation de pertinence comme une relation binaire R qui a lieu entre deux ensembles d'énoncés $\Gamma = \{p_1, \dots, p_n\}$, tel que $n \geq 1$, représentant l'ensemble des énoncés utilisés par le locuteur dans un contexte C , et un ensemble d'énoncés $\Delta = \{q_1, \dots, q_n\}$, tel que $n \geq 1$, représentant des hypothèses que l'interlocuteur fait dans ce même contexte. Le *domaine* de R (en abréviation, $dom(R)$) est l'ensemble Γ (énoncés du (des) locuteur(s)) dans

un contexte, et le *co-domaine* de R , ($co-dom(R)$), est l'ensemble Δ (énoncés exprimant des hypothèses de l'interlocuteur) dans ce même contexte. Dans le domaine il y a des ensembles d'énoncés et dans le contre-domaine également. Nous utilisons $\Gamma \vdash R \Delta$ comme une abréviation pour "les énoncés de Γ sont pertinents relativement aux énoncés de Δ " en ce sens qu'il existe une proposition nouvelle ϕ qui est strictement impliquée par toutes les propositions exprimées par les énoncés de l'ensemble complet $\Gamma \cup \Delta$ et non pas par les seuls énoncés de Γ ni les seuls énoncés de Δ . (En logique modale on dit qu'un ensemble Φ de propositions implique strictement une proposition ϕ quand cette proposition est vraie dans tous les contextes où toutes les propositions de cet ensemble Φ sont vraies). Pour analyser le contexte, nous nous proposons de sélectionner certains aspects contextuels qui nous sont utiles pour interpréter ce que veulent dire les interlocuteurs engagés dans une conversation particulière. Voici une explication intuitive de ce que nous entendons par "contexte".

Parmi les énoncés de l'ensemble Δ qui expriment des hypothèses de l'interlocuteur dans un contexte, certains représentent: des informations issues de la théorie de la signification linguistique qui permettent d'interpréter les énonciations du locuteur; des informations de la théorie des actes de

discours concernant les actes illocutoires exprimés et tentés dans le contexte; des informations sur le(s) participant(s), sur le temps, le moment, le lieu et le monde de l'énonciation, ainsi que sur le but conversationnel des interlocuteurs; des informations sur la culture (dans la mesure où les habitudes conversationnelles ne sont pas les mêmes), des informations sur l'histoire (ce que a été dit avant); l'arrière-plan, les formes de vie partagées par les interlocuteurs; leurs hypothèses comme leurs présuppositions, leurs suppositions, leurs états mentaux (comme par exemple, leurs croyances, désirs, intentions); leurs hypothèses scientifiques; les principes conversationnels (par exemple, le *principe de coopération* de Grice), les maximes conversationnelles, des règles, des normes, etc. Nous soulignons que dans notre approche, la relation de pertinence est, à la base, une relation entre des énoncés exprimant des assertions du locuteur et des hypothèses de l'interlocuteur dans un contexte. Pour cette raison, nous croyons que la pertinence a tout à voir avec la logique, puisque toute logique est un langage dont les énoncés sont liés par une relation de conséquence logique. De la même façon, dans une conversation, nous avons un langage commun muni d'une relation de pertinence entre énoncés sans laquelle la conversation intelligente échouerait. C'est donc cette relation de pertinence existant entre de tels énoncés qu'il faut préciser. Et c'est là notre objectif dans cette section. Pour ce faire, notre idée de base a été

de définir la relation de pertinence en analogie avec la notion de conséquence logique laquelle définit une logique particulière. Ainsi, nous commençons en énonçant le postulat suivant:

(2.1) *Postulat*

Nous soutenons la thèse que tous les traits saillants du contexte (comme par exemple, les hypothèses de l'interlocuteur, les faits mutuellement connus) peuvent être exprimés par des énoncés. Cette affirmation est basée sur le *principe d'exprimabilité* selon lequel "*l'on peut toujours dire ce que l'on veut dire*" (*Whatever can be meant can be said*):

Pour toute signification X , et pour tout locuteur L , chaque fois que L veut signifier (à l'intention de transmettre, désire communiquer, etc.) X , alors il est possible qu'existe une expression E , telle que E soit l'expression exacte ou la formulation exacte de X . Symboliquement: $(L) (X) (L \text{ signifie } X \rightarrow P (\exists E) (E \text{ est l'expression exacte de } X))$. (Searle, 1969, 20)¹⁵

Ainsi, dans une telle hypothèse, des traits saillants d'un contexte comme des faits, des changements, un événement comme par exemple, un vase qui se brise, etc. sont pris en considération dans la mesure où ces éléments sont exprimés par des énoncés. Dans le traitement d'un énoncé concernant un

¹⁵ Notre traduction.

contexte donné, il y a un nombre fini d'hypothèses saillantes du contexte à prendre en considération pour traiter cet énoncé dans ce contexte. Il existe dans chaque contexte, un nombre fini d'hypothèses saillantes dont les contenus propositionnels sont présumés vrais. Car ce qu'on veut dire, on veut le dire à partir d'un moyen fini (il existe donc un ensemble fini de faits de l'arrière-plan à prendre en considération). Les hypothèses que l'interlocuteur fait dans un contexte sont celles qu'il avait déjà à l'esprit auparavant ou celles qu'il fait sur ce contexte (ou les deux).

Pour la définition de la pertinence, nous proposons d'abord les définitions suivantes:

(2.2) *Règle inférentielle et conséquence directe* (en paraphrasant Mendelson, 1979)

Soit un ensemble fini $R_1, \dots, R_n$ de relations entre énoncés. Un tel ensemble est appelé l'ensemble de *règles d'inférence* si pour chaque R_i , il existe un seul entier positif j tel que, pour chaque ensemble j d'énoncés et chaque énoncé A , on peut effectivement décider si les j énoncés sont dans la relation R_i avec A , et dans ce cas, A est appelé une *conséquence directe* des énoncés donnés en vertu de R_i . En abréviation $A_1, \dots, A_j \vdash A$.

Dans une théorie déductive, les règles d'inférence sont celles qui permettent de conclure un énoncé à partir d'un ou de plusieurs énoncés pris au départ. Nous avons une conséquence directe quand certains énoncés permettent d'inférer par ces règles d'autres énoncés. Les règles d'inférence sont intéressantes quand elles sont valides. Dans ce cas, les conséquences directes qu'elles permettent de tirer, sont des conséquences sémantiques. Il n'est pas possible que les propositions exprimées par les prémisses soient vraies sans que la proposition exprimée par la conclusion soit également vraie.

(2.3) *Une Logique L*

Étant donné une langue L , *une logique L* (ou *système formel*) est l'ensemble des énoncés d'un langage plus un ensemble d'axiomes et de règles d'inférence formulés dans ce langage.¹⁶

¹⁶ Comme nous l'avons vu dans le chapitre précédent, cette définition concerne la dimension syntaxique d'une logique.

(2.4) *Conséquence pragmatique*

Un énoncé φ est une *conséquence pragmatique* de l'union de l'ensemble Γ (énoncés du locuteur) et de l'ensemble Δ (énoncés exprimant des hypothèses de l'interlocuteur) si et seulement si φ est une conséquence directe de l'union $\Gamma \cup \Delta$ de Γ et Δ . Puisqu'une conséquence pragmatique a toujours lieu dans une logique déterminée L , nous abrégions: $\Gamma \cup \Delta \vdash_L \varphi$.

(2.5) *Conséquence pragmatique triviale*

Nous disons qu'un énoncé φ est une *conséquence pragmatique triviale* de $\Gamma \cup \Delta$ si et seulement si φ est une conséquence directe de Γ ou une conséquence directe de Δ . En abréviation, $\Gamma \vdash_L \varphi$ ou $\Delta \vdash_L \varphi$.

(2.6) *Consistance*

Un ensemble d'énoncés $\Gamma \cup \Delta$ est *consistant* s'il n'arrive jamais que tout énoncé est une conséquence pragmatique de $\Gamma \cup \Delta$, c'est-à-dire, quand les règles d'inférence ne permettent pas de dériver

n'importe quel énoncé comme conséquence pragmatique. En abrégé, $\Gamma \cup \Delta \not\vdash_L A$ pour tout A .

Nous avons introduit cette définition afin d'éviter les trivialités.

(2.7) *La pertinence*

Un ensemble d'énoncés Γ du locuteur est *pertinent* pour un ensemble Δ d'énoncés exprimant des hypothèses de l'interlocuteur dans un contexte si et seulement si $\Gamma \cup \Delta$ est consistant et il existe une hypothèse φ telle que φ est une conséquence pragmatique non triviale de $\Gamma \cup \Delta$; et si pour chaque α et pour chaque β tel que $\Gamma \vdash_L \alpha$ et $\Delta \vdash_L \beta$, alors $\varphi \not\vdash_L \alpha \wedge \beta$.

(2.8) *Degrés de pertinence pour un interlocuteur*

Soit $\{\varphi_1, \dots, \varphi_k\}$ un ensemble d'énoncés représentant les conséquences pragmatiques nouvelles que le locuteur veut¹⁷ produire chez l'interlocuteur dans un contexte où il utilise l'ensemble d'énoncés Γ et où

¹⁷ Comme nous le verrons au Chapitre VI, la notion de pertinence a tout à voir avec la rationalité et la rationalité avec l'intentionnalité.

l'ensemble d'énoncés Δ exprime les hypothèses de l'interlocuteur dans ce contexte¹⁸. Considérons les alternatives suivantes:

1) l'interlocuteur n'infère aucune de ces conséquences pragmatiques nouvelles $\varphi_1, \dots, \varphi_k$ (dans ce cas, l'ensemble de ses hypothèses exprimées par des énoncés Δ dans le contexte C est différent de l'ensemble des hypothèses que le locuteur fait sur le contexte au moment de son énonciation). C'est ainsi que nous définissons le *degré zéro* (ou degré minimum) de pertinence pour un interlocuteur;

2) l'interlocuteur infère toutes les conséquences pragmatiques dans la séquence $\varphi_1, \dots, \varphi_k$. C'est ainsi que nous appelons le *degré 1* (ou degré maximum) de pertinence pour un interlocuteur;

3) cependant, il peut arriver que l'interlocuteur n'infère que quelques conséquences pragmatiques dans la séquence $\varphi_1, \dots, \varphi_k$. Or, la *Théorie des Ensembles* nous dit que l'ensemble de toutes les parties d'un ensemble de k éléments est 2^k . Donc, en excluant le degré zéro (minimum) de pertinence et le degré 1 (maximum) de pertinence, nous avons $2^k - 2$ degrés intermédiaires. C'est ainsi que nous définissons le *degré intermédiaire* de

¹⁸ Observons que l'ensemble $\{\varphi_1, \dots, \varphi_k\}$ est contenu dans l'ensemble de toutes les conséquences pragmatiques de $\Gamma \cup \Delta$, à savoir, $\{\varphi_1, \dots, \varphi_k, \varphi_{k+1}, \varphi_{k+2}, \dots, \varphi_{k+n}\}$. En abréviation, $\{\varphi_1, \dots, \varphi_k\} \subseteq \{\varphi_1, \dots, \varphi_k, \varphi_{k+1}, \varphi_{k+2}, \dots, \varphi_{k+n}\}$.

pertinence.

En ce qui concerne notre définition (2.7) de la pertinence, quelques commentaires sont nécessaires pour aider à éclairer cette définition qui s'appuie tant sur notre intuition que sur l'analyse logique, laquelle bien sûr doit être considérée si l'on veut s'approcher d'une analyse formelle de la pertinence. Comme nous l'avons déjà annoncé plus haut, c'est là notre objectif dans la présente section.

Tout d'abord, nous pensons prendre en considération " $\Gamma \cup \Delta$ comme consistant", car il nous semble raisonnable de supposer qu'une théorie de la pertinence généralise pratiquement toutes les logiques usuelles - ce qui montre que certains aspects de la pragmatique incluent toutes les formes d'inférences logiques valides, c'est-à-dire, les inférences pragmatiques incluent toutes les formes d'inférences logiques valides. Dans ce cas, en imposant cette condition, nous admettons la possibilité qu'apparaissent, dans l'ensemble $\Gamma \cup \Delta$ deux hypothèses telles que p et $\sim p$ et, dans ces circonstances, l'on préserve la logique paraconsistante. Dans ce cas, la notion de consistance est remplacée par la notion de non-trivialité.

Deuxièmement, nous assumons aussi qu'il existe un énoncé φ , telle que $\Gamma \not\vdash_L \varphi$, $\Delta \not\vdash_L \varphi$, mais $\Gamma \cup \Delta \vdash_L \varphi$ (c'est-à-dire, φ

est une conséquence pragmatique (mais non triviale) de $\Gamma \cup \Delta$). Ce que nous voulons dire est que l'énoncé ϕ doit être “nouveau” et qu'il doit, donc être dérivé de l'union de Γ et de Δ pris ensemble. Dans cet ordre d'idées, la première chose qui nous vient à l'esprit est qu'il faut que ϕ n'appartienne pas à Γ (dans le cas contraire, d'un quelconque Γ , on pourrait dériver un énoncé ϕ qui pourrait déjà être dans Γ lui-même). Mais au lieu de cela, nous avons pris $\Gamma \not\vdash_{\mathcal{L}} \phi$ pour la raison suivante: d'un point de vue logique nous pouvons prendre un énoncé quelconque r telle que $r \in \Gamma$, mais $r \vee s \notin \Gamma$. Alors, il est possible que $\Gamma \cup \Delta \vdash_{\mathcal{L}} r \vee s$ (où $r \vee s$ serait notre ϕ). Il en va de même pour l'ensemble d'énoncés contextuels Δ . Puisque s est un énoncé quelconque, alors n'importe quel ensemble Γ serait pertinent pour Δ (ce qui, évidemment, ne nous intéresse pas). Qui plus est, la restriction que $\Gamma \not\vdash_{\mathcal{L}} \phi$ (de même pour Δ , c'est-à-dire, $\Delta \not\vdash_{\mathcal{L}} \phi$) élimine les cas de lois et théorèmes logiques, avec Γ et Δ vides. Dans le cas contraire, l'énoncé ϕ n'apporterait rien de nouveau comme information. Suivant cette deuxième restriction: “ Γ est pertinent pour Δ si et seulement si $\Gamma \cup \Delta$ est consistant et s'il existe un énoncé ϕ telle que $\Gamma \not\vdash_{\mathcal{L}} \phi$, mais $\Gamma \cup \Delta \vdash_{\mathcal{L}} \phi$ ”. Cette solution serait, cependant, triviale pour la simple raison que nous pouvons supposer Γ, Δ quelconques où $\Gamma \vdash_{\mathcal{L}} \alpha$ et $\Delta \vdash_{\mathcal{L}} \beta$;

considérons maintenant $\alpha \wedge \beta$ comme notre φ . Alors, $\Gamma \nvdash_L \alpha \wedge \beta$, $\Delta \nvdash_L \alpha \wedge \beta$, mais $\Gamma \cup \Delta \vdash_L \alpha \wedge \beta$. Soit que $\Gamma \nvdash_L \varphi$ et $\Delta \nvdash_L \varphi$, on a $\Gamma \cup \Delta \vdash_L \varphi$; soit que $\Gamma \vdash_L \alpha$ et $\Delta \vdash_L \beta$, on a $\Gamma \cup \Delta \vdash_L \alpha \wedge \beta$ et, conséquemment, sans cette restriction n'importe quel ensemble Γ serait toujours pertinent pour Δ .

Pour terminer, remarquons que “pour chaque α et pour chaque β tels que $\Gamma \vdash_L \alpha$ et $\Delta \vdash_L \beta$, alors $\varphi \nvdash_L \alpha \wedge \beta$ ”. Notre objectif ici est justement d'éviter que l'énoncé φ ne fasse que donner une information qui était déjà déductible de Γ (ou de Δ), c'est-à-dire, il faut donner plus encore. (Remarquons que si φ est nouvelle alors $\varphi \neq \alpha \wedge \beta$, mais φ peut contenir $\alpha \wedge \beta$, parce qu'une information nouvelle peut contenir des informations déjà connues). Cette restriction paraît solutionner le problème; sinon, voyons: supposons l'ensemble Γ fini tel que $\Gamma \cup \Delta \vdash_L \varphi$; par le théorème de la déduction classique, $\Delta \vdash_L \Gamma \rightarrow \varphi$. Selon la logique classique, une formule du type ' $\varphi \rightarrow (\Gamma \rightarrow \varphi)$ ' est toujours valide et, en appliquant le théorème de la déduction à cette formule, nous avons $\varphi \vdash_L \Gamma \rightarrow \varphi$. Donc, nous avons $\Delta \vdash_L \Gamma \rightarrow \varphi$ et $\varphi \vdash_L \Gamma \rightarrow \varphi$, c'est-à-dire, la formule $\Gamma \rightarrow \varphi$ est déductible de Δ aussi bien que de φ . Par conséquent, quelles que soient les circonstances, Γ serait toujours pertinent pour Δ .

À première vue, il peut sembler que la définition de pertinence soit triviale, dans le sens qu'à partir d'un certain Γ et d'un Δ donnés, on en peut déduire quoi que ce soit. À cette fin, nous présentons le théorème suivant.

(2.9) *Meta-théorème*

Il existe des ensembles d'énoncés Γ et Δ tels que Γ n'est pas pertinent pour Δ (c'est-à-dire, la définition de la pertinence n'est pas triviale).

Démonstration: Il suffit de formuler des hypothèses pour Γ , Δ et φ , comme par exemple: Soit $\Gamma = \{ \text{La lune est faite de fromage vert} \}$; Soit $\Delta = \{ 2 + 2 = 4 \}$; Soit $\varphi = \text{Pierre est professeur de mathématique}$. Dans ce cas, selon la définition (2.7) de la pertinence, $\Gamma \cup \Delta \not\vdash_{\perp} \varphi$.

Pour éviter les problèmes exposés, en particulier les trivialités (du point de vue logique), nous devons définir une logique où ne vaut pas le théorème de la déduction classique et l'axiome $A \rightarrow (B \rightarrow A)$ non plus. Par conséquent, notre logique ne peut pas être la logique classique. Dans ce cas, nous devons partir d'une logique de la pertinence où les paradoxes de

l'implication matérielle (classique) ne sont pas valides.

Nous allons proposer par la suite, notre système pertinent P lequel a été obtenu en prenant pour base le système R de Anderson et Belnap (1975) car leur implication pertinente capte bien l'idée d'une relation de pertinence entre l'antécédent et le conséquent. L'axiomatique de R est la suivante:

(1)	$A \rightarrow A$	Idempotence
(2)	$(A \rightarrow B) \rightarrow ((C \rightarrow A) \rightarrow (C \rightarrow B))$	Transitivité
(3)	$(A \rightarrow (A \rightarrow B)) \rightarrow (A \rightarrow B)$	Contraction
(4)	$(A \rightarrow (B \rightarrow C)) \rightarrow (B \rightarrow (A \rightarrow C))$	Permutation
	La règle d'inférence: $A, A \rightarrow B \vdash B$	Modus Ponens
(5)	$A \wedge B \rightarrow A, A \wedge B \rightarrow B$	Élimination de la conjonction
(6)	$((A \rightarrow B) \wedge (A \rightarrow C)) \rightarrow (A \rightarrow (B \wedge C))$	Introduction de la conjonction
(7)	$A \rightarrow (A \vee B), B \rightarrow (A \vee B)$	Disjonction/Introduction
(8)	$((A \rightarrow C) \wedge (B \rightarrow C)) \rightarrow ((A \vee B) \rightarrow C)$	Disjonction/Élimination
(9)	$A \wedge (B \vee C) \rightarrow (A \wedge B) \vee C$	Distribution
	La règle d'inférence: $A, B \vdash A \wedge B$	Conjonction
(10)	$(A \rightarrow \sim A) \rightarrow \sim A$	Réduction
(11)	$(A \rightarrow \sim B) \rightarrow (B \rightarrow \sim A)$	Contraposition
(12)	$\sim \sim A \rightarrow A$	Double Négation

Les formules suivantes sont des méta-théorèmes de R :

- | | | |
|-----|--|------------------------------|
| (1) | $\vdash_R (A \rightarrow \sim B) \rightarrow ((A \rightarrow B) \rightarrow \sim A)$ | Faible réduction à l'absurde |
| (2) | $\vdash_R A \rightarrow \sim \sim A$ | |

$$(3) \quad \vdash_R A \vee \sim A$$

Comme nous allons le voir maintenant, notre système pertinent P a été construit à partir de l'axiomatisation du système R par l'exclusion de tous les axiomes de négation. Pour éviter des schémas d'axiomes de négation, nous introduirons tout de suite après une définition pour le symbole ' $\sim$ ' (négation).

3- Langage et axiomatique pour le système P

Le langage L de notre système formel P a les symboles primitifs suivants:

- a) un ensemble énumérable infini de variables propositionnelles ' p_i ';
- b) des opérateurs logiques propositionnels: $\wedge$ (conjonction), $\vee$ (disjonction), $\rightarrow$ (conditionnel);
- c) une constante propositionnelle: $\perp$ (faux);
- d) des symboles auxiliaires: $(,)$ (parenthèses).

La définition de formule bien formée ainsi que les conventions d'usage de parenthèses sont bien connues. L'ensemble des formules de L pour le système P sera nommé par $\text{FORM}_L(P)$.

Nous utiliserons les lettres latines majuscules, avec ou sans indices, comme des variables métalinguistiques pour formules, et pour ensembles de formules, des lettres grecques majuscules.

Les postulats (schémas d'axiomes (S.A) et règles primitives) de $\mathcal{P}$ sont les suivants:

- | | | |
|------|---|-----------------------------------|
| (1) | $A \rightarrow A$ | Idempotence |
| (2) | $(A \rightarrow B) \rightarrow ((B \rightarrow C) \rightarrow (A \rightarrow C))$ | Transitivité |
| (3) | $(A \rightarrow (A \rightarrow B)) \rightarrow (A \rightarrow B)$ | Contraction |
| | et sa variante: | |
| (3') | $(A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$ | Auto-Distribution |
| (4) | $(A \rightarrow (B \rightarrow C)) \rightarrow (B \rightarrow (A \rightarrow C))$ | Permutation |
| (5) | $(A \wedge B) \rightarrow A, (A \wedge B) \rightarrow B$ | Élimination de la
conjonction |
| (6) | $((A \rightarrow B) \wedge (A \rightarrow C)) \rightarrow (A \rightarrow (B \wedge C))$ | Introduction de la
conjonction |
| (7) | $A \rightarrow (A \vee B), B \rightarrow (A \vee B)$ | Introduction de la
disjonction |
| (8) | $((A \rightarrow C) \wedge (B \rightarrow C)) \rightarrow ((A \vee B) \rightarrow C)$ | Élimination de la
disjonction |
| (9) | $(A \wedge (B \vee C)) \rightarrow ((A \wedge B) \vee C)$ | Distribution |
| (10) | $A, A \rightarrow B \vdash B$ | <i>Modus Ponens</i> -
'M.P.' |
| (11) | $A, B \vdash A \wedge B$ | Conjonction |

Pour le symbole ' $\sim$ ' (négation) nous allons introduire la définition suivante:

$$(3.1) \sim A =_{\text{df}} A \rightarrow \perp$$

Le symbole ' $\vdash$ ' et les notions de démonstration, de déduction et de théorème se différencient de ceux de la logique classique seulement dans la mesure où nous ajoutons un critère de pertinence pour la notion de déduction (voir la définition de déduction pertinente plus loin).

Nous allons maintenant présenter quelques métathéorèmes de P . Les schémas d'axiomes (10) et (11) de R , tout comme quelques-uns de leurs théorèmes, présentés ci-dessous, peuvent être dérivés de schémas d'axiomes de P .

(3.2) Pour toutes les formules $A, B \in \text{FORM}_{L(P)}$:

- a) $\vdash_P (A \rightarrow \sim A) \rightarrow \sim A$ (S.A. 10) de R
- b) $\vdash_P (A \rightarrow \sim B) \rightarrow (B \rightarrow \sim A)$ (S.A. 11) de R
- c) $\vdash_P (A \rightarrow \sim B) \rightarrow ((A \rightarrow B) \rightarrow \sim A)$ (faible réduction à l'absurde) de R
- d) $\vdash_P (A \rightarrow B) \rightarrow (\sim B \rightarrow \sim A)$ (contraposition faible)

Démonstration:

- 1) $(A \rightarrow B) \rightarrow ((B \rightarrow \perp) \rightarrow (A \rightarrow \perp))$ (S.A.2)
- 2) $(A \rightarrow B) \rightarrow (\sim B \rightarrow \sim A)$ 1 / déf. (2.1)
- 3) $(A \rightarrow (A \rightarrow \perp)) \rightarrow (A \rightarrow \perp)$ (S. A. 3)

- 4) $(A \rightarrow \sim A) \rightarrow \sim A$ 3 / déf. (2.1)
- 5) $(A \rightarrow (B \rightarrow \perp)) \rightarrow (B \rightarrow (A \rightarrow \perp))$ (S.A.4)
- 6) $(A \rightarrow \sim B) \rightarrow (B \rightarrow \sim A)$ 5 / déf. (2.1)
- 7) $(A \rightarrow (B \rightarrow \perp)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow \perp))$ (S. A.3')
- 8) $(A \rightarrow \sim B) \rightarrow ((A \rightarrow B) \rightarrow \sim A)$ 7 / déf. (2.1)

(3.3) Pour toute formule $A \in \text{FORM}_{L(P)} \vdash_P A \rightarrow \sim \sim A$

Démonstration:

- 1) $(\sim B \rightarrow \sim B) \rightarrow (B \rightarrow \sim \sim B)$ Métathéorème (2.2.b)
- 2) $\sim B \rightarrow \sim B$ (S. A.1)
- 3) $B \rightarrow \sim \sim B$ 1, 2 / M. P.

(3.4) *Définition: Déduction Pertinente*

Une déduction d'une formule B à partir de formules $A_1, \dots, A_n$ est *pertinente* pour une hypothèse donnée A_i ssi A_i est effectivement utilisée dans la déduction de B , c'est-à-dire, s'il existe une chaîne d'inférences connectant A_i avec la formule finale B (Church, 1951). Dans ce cas, A_i est *nécessaire* dans la déduction de B . En symboles,

$$\{A_1, \dots, A_n\} - \{A_i\} \not\vdash_P B.$$

(3.5) *Théorème de la Déduction Pertinente* (Moh, Church)

S'il existe une déduction dans P de B à partir de $A_1, \dots, A_n, A$, laquelle est pertinente par rapport à A , alors il existe une déduction dans P de $A \rightarrow B$ à partir de $A_1, \dots, A_n$. Qui plus est, la nouvelle déduction sera "aussi pertinente" que l'ancienne, c'est-à-dire, quelconque A_i utilisée dans la déduction donnée sera utilisée aussi dans la nouvelle déduction. En abrégé, Si $A_1, \dots, A_n, A \vdash_P B$, alors $A_1, \dots, A_n \vdash_P A \rightarrow B$.

Démonstration:

Soit $B_1, \dots, B_k$ la déduction donnée avec une analyse particulière de la façon dont chaque pas a été justifié. Par induction, l'on peut montrer pour chaque B_i que si A a été utilisée pour l'obtention de B_i , alors il existe une déduction de $A \rightarrow B_i$ à partir de $A_1, \dots, A_n$ et une déduction de B_i à partir des mêmes hypothèses.

Cas 1: B_i a été justifiée comme hypothèse. Alors, ou B_i est A ou elle est l'une des A_j . Mais $A \rightarrow A$ est un axiome de P (et donc déductible de $A_1, \dots, A_n$), ce qui rend compte de la première alternative. Et clairement pour

la deuxième alternative B_i est déductible de $A_1, \dots, A_n$ (étant l'une de celles-ci).

Cas 2: B_i a été justifiée comme un axiome. Alors, A n'a pas été utilisée dans l'obtention de B_i , et certainement B_i est déductible (étant un axiome).

(Notez que A n'a pas été nécessaire dans l'obtention de B_i ; différemment de la logique classique, nous avons que $A \not\vdash_P B_i$).

Cas 3: B_i a été justifiée comme venant des étapes précédentes $B_j \rightarrow B_i$ et B_j par *modus ponens*. Il existe quatre sous-cas qui dépendent de la question de savoir si ou non A a été utilisée dans l'obtention de prémisses.

Sous-cas 3.1: A a été utilisée dans l'obtention de $B_j \rightarrow B_i$ et B_j . Alors, par l'hypothèse d'induction $A_1, \dots, A_n \vdash_P A \rightarrow (B_j \rightarrow B_i)$ et $A_1, \dots, A_n \vdash_P A \rightarrow B_j$. Donc $A \rightarrow B$ peut être obtenue en utilisant l'axiome (3') : $(A \rightarrow (B_j \rightarrow B_i)) \rightarrow ((A \rightarrow B_j) \rightarrow (A \rightarrow B_i))$ (*auto-distribution*).

Sous-cas 3.2: A a été utilisée dans l'obtention de $B_j \rightarrow B_i$ mais pas

B_j . En utilisant l'axiome (4): $(A \rightarrow (B_j \rightarrow B_i)) \rightarrow (B_j \rightarrow (A \rightarrow B_i))$ (Permutation) nous pouvons obtenir $A \rightarrow B_i$ à partir de $A \rightarrow (B_j \rightarrow B_i)$ et B_j .

Sous-cas 3.3: A n'a pas été utilisée dans l'obtention de $B_j \rightarrow B_i$ mais A a été utilisée dans l'obtention de B_j . En utilisant l'axiome (2): $(A \rightarrow B_j) \rightarrow ((B_j \rightarrow B_i) \rightarrow (A \rightarrow B_i))$ (*transitivité*), nous obtenons $A \rightarrow B_i$ à partir de $B_j \rightarrow B_i$ et de $A \rightarrow B_j$.

Sous-cas 3.4: A n'a pas été utilisée dans l'obtention ni de $B_j \rightarrow B_i$ ni de B_j . Donc B_i vient de $B_j \rightarrow B_i$ et de B_j en utilisant la règle *modus ponens*.

Remarque: Notre logique $\mathbf{P}$, à l'instar d'autres logiques de la pertinence, ne permet pas d'inférer $\Gamma, B \vdash_{\mathbf{P}} A$, à partir de $\Gamma \vdash_{\mathbf{P}} A$, c'est-à-dire, dans $\mathbf{P}$, ne vaut pas en général la monotonie.

À la suite de ces constatations nous introduisons la définition suivante:

(3.6): *Implication Pertinente*

A implique pertinemment B ssi B se déduit de A pertinemment.

En abréviation, $A \supset B =_{\text{df}} A \vdash_P B$.

Remarque: Les théorèmes classiques suivants ne sont pas des théorèmes de P :

- (1) $\vdash_P (A \wedge \sim A) \rightarrow B$ (Duns Scot)
- (2) $\vdash_P A \rightarrow (\sim A \rightarrow B)$ (Duns Scot)
- (3) $\vdash_P \sim A \rightarrow (A \rightarrow B)$ (Duns Scot)
- (4) $\vdash_P A \rightarrow (\sim A \rightarrow \sim B)$ (Duns Scot)
- (5) $\vdash_P \sim A \rightarrow (A \rightarrow \sim B)$ (Duns Scot)
- (6) $\vdash_P A \rightarrow (B \rightarrow A)$
- (7) $\vdash_P A \rightarrow (\sim B \rightarrow A)$
- (8) $\vdash_P (A \rightarrow (B \rightarrow C)) \rightarrow ((A \wedge B) \rightarrow C)$ (Importation)
- (9) $\vdash_P ((A \wedge B) \rightarrow C) \rightarrow (A \rightarrow (B \rightarrow C))$ (Exportation)
- (10) $\vdash_P (\sim A \rightarrow \sim B) \rightarrow (B \rightarrow A)$ (Contraposition forte)
- (11) $\vdash_P \sim \sim A \rightarrow A$ et, conséquemment, $\vdash_P A \vee \sim A$
- (12) $\vdash_P (A \rightarrow B) \vee A$
- (13) $\vdash_P (A \rightarrow B) \vee (B \rightarrow A)$
- (14) $\vdash_P (\sim A \rightarrow B) \rightarrow (\sim B \rightarrow A)$

Démonstration:

En prenant quelques exemples du langage naturel, nous voyons que les théorèmes de la logique classique (LC) ci-dessous ne sont pas des

théorèmes de **P** :

$$(1) \quad \vdash_{\text{LC}} (A \wedge \sim A) \rightarrow B$$

Par le théorème de la déduction classique, $A \wedge \sim A \vdash_{\text{LC}} B$.

Alors, est-ce que $\vdash_P (A \wedge \sim A) \rightarrow B$, c'est-à-dire, $A \wedge \sim A \vdash_P B$ et, par conséquent, $A \wedge \sim A$ implique pertinemment B ?

Soit $A \wedge \sim A$ l'hypothèse suivante: Pierre est vieux et Pierre n'est pas vieux.

Soit B l'hypothèse suivante: Pierre est âgé de 95 ans.

Selon notre définition d'*Implication Pertinente* (3.6), $A \wedge \sim A$ implique pertinemment B ssi B se déduit de A pertinemment; en abréviation, $A \supset B =_{\text{df}} A \vdash_P B$.

Supposons maintenant que $A \wedge \sim A$ implique B pertinemment. Si tel est le cas, selon notre définition de *Déduction Pertinente* (3.4), $A \wedge \sim A$ est effectivement utilisée dans la déduction de B , c'est-à-dire, il existe une chaîne d'inférences connectant $A \wedge \sim A$ avec la formule finale B . Dans ce cas, $A \wedge \sim A$ est nécessaire dans la déduction de B . En abréviation, $\{ \Gamma \} - \{ A \wedge \sim A \} \vdash \not\vdash_P B$.

Mais, comme nous pouvons le voir, cela ne se vérifie pas puisqu'il

n'y a pas de règles d'inférence dans la logique $\mathbf{P}$ qui permettent de déduire B à partir de $A \wedge \sim A$, ou en d'autres mots, il n'existe aucune chaîne d'inférences connectant $A \wedge \sim A$ avec la formule B . Dans ce cas, $A \wedge \sim A$ n'est pas du tout nécessaire dans la déduction de B . Ainsi, $A \wedge \sim A$ ne contribue en rien à la déduction de B . Donc, $A \wedge \sim A \not\vdash_P B$. Selon la définition (3.6), $A \wedge \sim A$ n'implique pas pertinemment B , et ainsi, $\not\vdash_P A \wedge \sim A \rightarrow B$.

Voyons maintenant pour les cas suivants:

$$(2) \quad \vdash_{\text{LC}} A \rightarrow (\sim A \rightarrow B)$$

$$(3) \quad \vdash_{\text{LC}} \sim A \rightarrow (A \rightarrow B)$$

$$(4) \quad \vdash_{\text{LC}} A \rightarrow (\sim A \rightarrow \sim B)$$

$$(5) \quad \vdash_{\text{LC}} \sim A \rightarrow (A \rightarrow \sim B)$$

Soient A , $\sim A$ et B les hypothèses suivantes:

$$A = 2 + 2 = 4$$

$$\sim A = 2 + 2 \neq 4$$

$$B = 2 \text{ est un nombre pair}$$

$$\sim B = 2 \text{ est un nombre impair.}$$

Par le théorème de la déduction classique, nous avons en (2) que

$A, \sim A \vdash_{\text{LC}} B$, et en (3) que $\sim A, A \vdash_{\text{LC}} B$. Mais, selon notre définition de *Déduction Pertinente* (3.4), les deux hypothèses $A, \sim A$ n'impliquent pas B pertinemment, c'est-à-dire, $A, \sim A \not\vdash_P B$, puisqu'elles ne sont pas effectivement utilisées dans la déduction de B . En d'autres termes, il n'existe pas de chaîne d'inférences connectant $A, \sim A$ avec la formule finale B . Alors, soit en (2) soit en (3), les hypothèses $A, \sim A$ ne sont pas nécessaires dans la déduction de B . Donc, selon notre définition d'*Implication Pertinente* (3.6), $A, \sim A$ n'impliquent pas pertinemment B , et ainsi, $\not\vdash_P A \rightarrow (\sim A \rightarrow B)$ et $\not\vdash_P \sim A \rightarrow (A \rightarrow B)$.

De même pour les théorèmes (4) et (5), nous avons que (4) $\not\vdash_P A \rightarrow (\sim A \rightarrow \sim B)$ et (5) $\not\vdash_P \sim A \rightarrow (A \rightarrow \sim B)$.

Notez que en (2) nous avons que l'affirmation d'une thèse A et de sa négation $\sim A$ implique une hypothèse quelconque B . De même, pour ses variantes (3), (4) et (5).

Du point de vue de la logique de la pertinence, en particulier de son sous-calcul P , ce qui est intéressant de capturer ce sont les aspects de validité des propositions du type “Si A alors B ” seulement quand A est pertinente pour B .

Pour les théorèmes (6) $\vdash_{\text{LC}} A \rightarrow (B \rightarrow A)$ et sa variante (7)

$\vdash_{\text{LC}} A \rightarrow (\sim B \rightarrow A)$, prenons les contre-exemples suivants pour montrer qu'ils ne sont pas théorèmes de **P** :

$$(6) \quad \vdash_{\text{LC}} A \rightarrow (B \rightarrow A)$$

Soit A, B les hypothèses suivantes:

$A =$ Je pratique des sports en soirée.

$B =$ Je dors en soirée.

Pour le théorème de la déduction classique, nous avons en (6) que $A, B \vdash_{\text{LC}} A$. Mais, selon notre définition d'*Implication Pertinente* (3.6), A, B n'impliquent pas pertinemment A , c'est-à-dire, $\not\vdash_P A \rightarrow (B \rightarrow A)$.

Pour la variante du théorème (6), nous avons:

$$(7) \quad \vdash_{\text{LC}} A \rightarrow (\sim B \rightarrow A)$$

Soient A et $\sim B$ les hypothèses suivantes:

$A =$ Je suis en Europe.

$\sim B =$ Je ne suis pas à Paris

Ainsi, $\not\vdash_P A \rightarrow (\sim B \rightarrow A)$

Voyons maintenant le théorème classique (8):

$$(8) \quad \vdash_{\text{LC}} (A \rightarrow (B \rightarrow C)) \rightarrow ((A \wedge B) \rightarrow C)$$

Soient A, B et C les hypothèses suivantes:

$A =$ Pierre est âgé de 95 ans.

$B =$ Pierre habite dans un asile de personnes âgées.

$C =$ Pierre est âgé.

Par le théorème de la déduction classique, nous avons en (8) que $A \rightarrow (B \rightarrow C), A \wedge B \vdash_{\text{LC}} C$.

Mais, selon notre définition d'*Implication Pertinente* (3.6), $A \rightarrow (B \rightarrow C)$ et $A \wedge B \not\vdash_{\text{P}} C$.

De même:

$$(9) \quad \vdash_{\text{LC}} ((A \wedge B) \rightarrow C) \rightarrow (A \rightarrow (B \rightarrow C))$$

Soient A, B et C les hypothèses suivantes:

$A =$ Je vais en France

$B =$ Je vais à Québec

$C =$ Je peux communiquer en français.

Ainsi, $\not\vdash_{\text{P}} ((A \wedge B) \rightarrow C) \rightarrow (A \rightarrow (B \rightarrow C))$

De même pour:

$$(10) \vdash_{LC} (\sim A \rightarrow \sim B) \rightarrow (B \rightarrow A)$$

Soient A, B les hypothèses suivantes:

$A =$ J'ai un permis de conduire ma voiture.

$B =$ Je suis libre d'infraction de la circulation.

De la même manière :

$$(11) \vdash_{LC} \sim\sim A \rightarrow A \text{ et conséquemment, } \vdash_{LC} A \vee \sim A$$

Soient A et $\sim A$ les hypothèses suivantes:

$A =$ Jean est au-dessus de 100 kilos.

$\sim A =$ Jean n'est pas au-dessus de 100 kilos.

Supposons que Jean pèse juste 100 kilos. Dans ce cas,

$$\sim\sim A \not\vdash_P A.$$

Et, de la même manière pour les deux cas suivants:

$$(12) \vdash_{LC} (A \rightarrow B) \vee A$$

$$(13) \vdash_{LC} (A \rightarrow B) \vee (B \rightarrow A)$$

Soient A, B les hypothèses suivantes:

$A =$ La lune est faite de fromage vert.

$$B = 2 + 2 = 4.$$

Alors, pour ce qui est, par exemple de (12), évidemment il nous semble bien contre-intuitif que: *Si la lune est faite de fromage vert alors $2 + 2 = 4$, ou alors la lune est faite de fromage vert.* Dans ce cas, $\vdash \neg_P (A \rightarrow B) \vee A$ et $\vdash \neg_P (A \rightarrow B) \vee (B \rightarrow A)$.

Et de même pour (13).

$$(14) \vdash_{LC} (\sim A \rightarrow B) \rightarrow (\sim B \rightarrow A)$$

Soient A, B les hypothèses suivantes:

$A =$ J'ai un permis de conduire ma voiture.

$B =$ J'ai fait une infraction au code de la route.

Alors, $\vdash \neg_P (\sim A \rightarrow B) \rightarrow (\sim B \rightarrow A)$

Tout bien considéré, il serait intéressant d'introduire maintenant une sémantique pour notre système P .

Comme nous l'avons dit précédemment, notre système P a été construit en prenant pour base le système R de Anderson et Belnap (1975) parce que la signification qu'ils ont donné au connecteur " $\rightarrow$ " (si ... alors)

capture bien l'idée d'une relation de pertinence entre l'antécédent et le conséquent de cette implication. Ainsi nous avons pensé faire une sémantique, laquelle serait une adaptation très simple de la sémantique de R construite à partir des concepts de base de la sémantique de Kripke (la sémantique des mondes possibles). Mais malheureusement le système axiomatique de Anderson et Belnap, qui définit l'implication correspondante, était privé d'une sémantique satisfaisante. Pour cette raison, au lieu d'en proposer une, nous indiquons simplement qu'un progrès décisif a été réalisé par G. Priest et R. Sylvan (1992) et qu'il serait enrichissant de l'ajouter à notre logique P (avec quelques adaptations mineures) parce que leur sémantique utilise des outils conceptuels de la sémantique des mondes possibles, comme *relation binaire d'accessibilité* (pour capturer les différents sens de l'opérateur unaire "nécessairement"), et la *relation d'accessibilité ternaire* (pour capturer les différents sens du *conditionnel pertinent* qui est un connecteur binaire) - voir leur article "Simplified Semantics for Basic Relevant Logic", 1992.

CHAPITRE IV

QUELQUES APPLICATIONS DE LA LOGIQUE DE LA PERTINENCE P À LA THÉORIE DE LA PERTINENCE DE SPERBER ET WILSON

Introduction

Nous avons proposé quelques définitions techniques de la pertinence dans le chapitre antérieur, qui visent à éclairer quelque peu et corriger la lacune rencontrée dans la théorie de la pertinence de Sperber et Wilson en ce qui concerne la définition d'implication non-triviale, ainsi que le critère d'évaluation du degré de pertinence d'une information.

Dans ce chapitre, nous allons discuter premièrement le problème de l'implication non-triviale car il s'agit d'une notion centrale dans la théorie de la pertinence, puisqu'elle est utilisée dans la définition même d'implication contextuelle; les implications contextuelles constituent une partie de la classe des effets contextuels, et ceux-ci sont déterminants pour établir la pertinence

d'une hypothèse par rapport à un contexte fixé.

1. Le problème de la réitération *ad infinitum* dans la définition d'implication triviale chez Sperber et Wilson.

À titre de rappel, revenons aux définitions d'implication logique non triviale et d'implication contextuelle chez Sperber et Wilson (1989):

Implication logique non triviale

Un ensemble d'hypothèses P implique logiquement et non trivialement une hypothèse Q si et seulement si, lorsque P est l'ensemble des thèses initiales d'une dérivation qui ne fait appel qu'à des règles d'élimination, Q appartient à l'ensemble de thèses finales (p. 152).

Implication contextuelle

Un ensemble d'hypothèses P implique contextuellement une hypothèse Q dans le contexte C si et seulement si:

- (a) l'union de P et de C implique non-trivialement Q ,
- (b) P n'implique pas non-trivialement Q , et

(c) C n'implique pas non-trivialement Q (p. 166).

Conformément à ce qu'on a déjà vu, dans la terminologie de Sperber et Wilson, les implications contextuelles, ainsi que les renforcements et effacements de prémisses, constituent la classe des effets contextuels et ce sont eux qui déterminent la pertinence d'une hypothèse: "une hypothèse est *pertinente* dans un contexte si et seulement si elle a un effet contextuel dans ce contexte" (p. 187).

En somme, dans la théorie de la pertinence de Sperber et Wilson, une information est *pertinente* si et seulement si, en s'intégrant à un système de croyances, elle y entraîne des modifications comme: (a) augmenter ou diminuer la force de certaines croyances du système qu'elle confirme, affaiblit ou élimine; (b) en conjonction avec certaines croyances du système, elle peut avoir des implications contextuelles. Ici, l'idée de base est que la pertinence d'une hypothèse est fonction de deux facteurs: plus une information a d'effets contextuels, plus elle est pertinente; plus il y a d'efforts requis pour traiter cette information, moins elle est pertinente. Dans ce cas, il nous semble qu'il faudra tout d'abord bien expliquer la théorie de l'implication non-triviale de Sperber et Wilson pour préciser la notion d'implication contextuelle et conséquemment, celle d'effet contextuel afin que leur propre théorie de la pertinence soit

débarrassée de cette lacune, soit le manque de précision dans la définition de l'implication non-triviale. Remarquons l'analyse critique qu'a faite Dominicy (1991) du problème de l'implication non-triviale. Comme lui-même le remarque, Sperber et Wilson se contentent d'opposer aux règles d'élimination, les règles d'introduction. Mais il peut y avoir d'autres règles avec un statut différent:

Ils semblent admettre que les deux classes ainsi définies sont complémentaires. Mais on débusque facilement des règles qui tombent dans une catégorie intermédiaire, telles " P . Donc, P " ou " P ou Q . Donc Q ou P ". Comme il est précisé aux p. 160-161 que toute hypothèse s'implique non-trivialement, nous pourrions croire que l'implication non-triviale exige, non pas le recours aux seules règles d'élimination, mais plutôt le non-recours aux règles d'introduction. Cependant, cette position de repli ne permet plus d'exclure l'application *ad infinitum* d'une règle récursive (" P . Donc, P . Donc, P , ..." ou " P ou Q . Donc, Q ou P . Donc, P ou Q , ..."), ce que voulait interdire la restriction aux règles d'élimination. En ce qui concerne l'itération d'hypothèse (" P . Donc, P ") et la commutativité (" P ou Q . Donc, Q ou P "), une conjecture séduisante, et sans doute testable, consiste à les traiter comme des stratégies de preuve effective, à ne pas confondre avec des règles de déduction. Il reste, malgré tout, qu'en se limitant aux règles d'élimination, on se condamne à encoder immédiatement des règles telles que le modus ponens conjonctif ou le modus ponens disjonctif, que les systèmes déductifs classiques obtiennent par dérivation (p. 154). Dans le cas de la contraposition, l'encodage direct ("Si P , alors Q . Donc, si non- Q , alors non- P ") produirait par contre une règle d'introduction; il faudra donc passer par l'équivalence traditionnelle (et fort gênante) entre "Si P , alors Q " et "Non- P ou Q ", ce qui ouvre de nouveau la porte à des récursions *ad infinitum* ("Si P , alors Q . Donc, non- P ou Q . Donc, si P , alors Q , ..."). Ce dernier exemple montre, en sus, que des règles d'élimination symétriques (des définitions formelles) créent les mêmes pathologies que les règles d'introduction. (Dominicy, 1991, p. 89)

Or, ce que Dominicy a montré, est que telle quelle, la définition

d'implication non-triviale chez Sperber et Wilson ne réussit pas à se libérer de l'application *ad infinitum* d'une règle récursive (à une hypothèse P quelconque) telle quelle: $P \vdash P \vdash P \vdash P \dots$

De notre part, comme nous l'avons déjà montré au chapitre III, nous avons proposé une définition de la pertinence, la définition (2.7) dans laquelle notre préoccupation majeure était d'éviter justement les trivialités. Cela nous a amené à y introduire une restriction (c'est-à-dire, une exigence de plus) dans le but d'empêcher l'application *ad infinitum* d'une règle récursive à un certain énoncé P . Selon notre définition,

(2.7) *La pertinence*

Un ensemble d'énoncés Γ du locuteur est *pertinent* pour un ensemble Δ d'énoncés exprimant des hypothèses de l'interlocuteur dans un contexte si et seulement si $\Gamma \cup \Delta$ est consistant et il existe une hypothèse φ telle que φ est une conséquence pragmatique non triviale de $\Gamma \cup \Delta$; et si pour chaque α et pour chaque β tel que $\Gamma \vdash_L \alpha$ et $\Delta \vdash_L \beta$, alors $\varphi \vdash_L \alpha \wedge \beta$.

Ainsi, tout de suite après cette définition, nous avons fait, dans le

chapitre mentionné, quelques commentaires pour bien expliquer l’objectif de nos restrictions, en particulier de celle qui permet d’éviter les implications triviales:

“ ... $\Gamma \cup \Delta \vdash_L \varphi$, et si pour chaque α et pour chaque β tel que $\Gamma \vdash_L \alpha$ et $\Delta \vdash_L \beta$, alors $\varphi \not\vdash_L \alpha \wedge \beta$ ”.

En d’autres mots, notre définition de la pertinence prévoit que $\Gamma \cup \Delta \vdash_L \varphi$; et si $\Gamma \vdash_L \alpha$ et $\Delta \vdash_L \beta$, alors $\varphi \not\vdash_L \alpha \wedge \beta$. Cela nous permet d’empêcher la réitération *ad infinitum*, vu que dans notre définition, nous exigeons que $\varphi \not\vdash_L \varphi$ (car de même que Sperber et Wilson nous nous intéressons à la production de conséquences “nouvelles”). Nous bloquons de même, la possibilité que $\varphi \vdash_L \varphi \wedge \varphi \dots$, c’est-à-dire, $\varphi \not\vdash_L \varphi \wedge \varphi \dots$, et du fait que notre définition ne fait appel à aucune règle d’élimination, il s’ensuit que $\not\vdash_L \varphi$. Ainsi, nous éliminons comme hypothèses pertinentes les lois logiques comme, par exemple, les théorèmes logiques où Γ et Δ sont vides (contrairement à ce qui survient chez Sperber et Wilson, car ils n’admettent que les règles d’élimination). Dans ces circonstances, notre définition de la pertinence a une portée plus grande que celle de Sperber et Wilson.

Qui plus est, notre définition de la pertinence ne permet pas non

plus l'introduction de la conjonction, et du fait que la conjonction est l'introduction la plus forte, nous pouvons conclure que telle quelle, notre définition n'accepte aucune règle d'introduction qui permet de tomber dans une réitération infinie. En effet, comme nous pouvons le voir, parmi les règles d'introduction de la logique classique, nous avons la règle d'*introduction* de *et*, la règle d'*introduction* de *si ... alors*, la règle d'*introduction* de *ou* et la règle d'*introduction* de *non*, et toutes ces règles sont bloquées dans notre définition de la pertinence, comme nous pouvons le voir par la suite, en introduisant la Proposition suivante :

Étant donné φ , α , β comme dans les conditions de la définition de la pertinence (2.7), il se suit que φ ne peut pas être obtenu à partir de α et de β par une règle d'introduction.

Démonstration : Examinons d'abord le cas de l'*introduction* de *et*.

Si φ a été obtenu par cette règle, φ serait $\alpha \wedge \beta$, mais cela est interdit par la définition de la pertinence, puisque $\alpha \wedge \beta \vdash_L \alpha \wedge \beta$. (Si $\varphi \neq \alpha \wedge \beta$, alors $\Gamma \vdash_L \varphi$ ou $\Delta \vdash_L \varphi$, ce que est exclu par la définition); pour ce qui est de la règle d'*introduction* de *ou* : si tel est le cas, φ serait $\alpha \vee \gamma$ ou $\beta \vee \gamma$. Donc, φ serait une conséquence de Γ tout seul ou de Δ tout seul, ce que est interdit par la définition de la pertinence. Quant à l'*introduction* de "*si ...*

alors”, dans ce cas, φ serait $\gamma \rightarrow \alpha$ ou $\gamma \rightarrow \beta$, ce que n’est pas possible, tout comme dans le cas antérieur. Relativement à l’introduction de *non*, φ serait $\sim \sim \alpha$ ou $\sim \sim \beta$, ce que ne serait pas possible, tout comme dans les deux cas antérieurs.

Tout bien considéré, du fait que telle quelle, la définition d’implication non-triviale de Sperber et Wilson ne réussit pas à se libérer de la réitération *ad infinitum* (comme le montre Dominicy, elle n’élimine pas les trivialités), nous pouvons conclure que notre définition est plus précise, plus rigoureuse, plus riche que celle de Sperber et Wilson, vu qu’elle rend compte effectivement de ce problème-là.

Notons que notre définition de la $A \vdash_{\perp} A \vee B$, aucune règle d’élimination, etc. Par contre notre logique **P** a pour axiome $\vdash_P A \rightarrow (A \vee B)$, mais n’a pas $A \vdash A \vee B$ comme règle de **P**, c’est-à-dire, $A \not\vdash_P A \vee B$. Il s’agit donc de deux choses bien distinctes et qu’il faut préciser: c’est une chose que de traiter de ce qui concerne la théorie elle-même (leurs définitions techniques), c’en est une autre que de traiter de ce qui concerne la logique qui correspond à cette théorie.

L’objectif principal de notre thèse est d’offrir une bonne définition de la pertinence. Dans notre théorie de la pertinence nous avons fixé un

ensemble Γ (concernant les énoncés du locuteur) plus un ensemble Δ (concernant les énoncés contextuels de l'interlocuteur), dans le but d'élaborer une définition adéquate de la pertinence. Plus précisément, notre définition de la relation " Γ est pertinent pour Δ " appartient à la théorie elle-même et non pas à la logique $\mathbf{P}$. Dans notre théorie nous avons $\Gamma \cup \Delta$ tandis que dans la logique $\mathbf{P}$, $\Gamma = \emptyset$ et $\Delta = \emptyset$ (Γ et Δ sont vides). Par exemple, notre définition de pertinence prévoit que $A \not\vdash_{\mathbf{L}} A$ tandis que dans la logique $\mathbf{P}$ nous avons l'axiome suivant: $\vdash_{\mathbf{P}} A \rightarrow A$. Considérons l'exemple suivant où LC signifie la logique classique. Soit Γ l'ensemble formé par la seule hypothèse: *Marie est partie en vacances*. Soit $\Gamma \cup \Delta \vdash_{\text{LC}} \text{Marie est partie en vacances ou } 2 + 2 = 5$ (appelons cette disjonction de notre conséquence ϕ) vu que dans la logique classique $A \vdash_{\text{LC}} A \vee B$. Cependant, dans la définition de pertinence $\Gamma \cup \Delta \not\vdash_{\mathbf{L}} \text{Marie est partie en vacances ou } 2 + 2 = 5$ car si tel était le cas, ϕ serait triviale.

Bien entendu, une règle est plus forte qu'un axiome dans le sens suivant: dans le cas d'un axiome comme $\vdash_{\mathbf{L}} A \rightarrow A \vee B$, nous avons que si l'antécédent A de cette implication est faux (en supposant une logique bivalente), toute la formule est vraie par la définition même du connectif vérifonctionnel " $\rightarrow$ ". Qui plus est, puisqu'un tel axiome est déductible de

l'ensemble $\emptyset$ (vide) d'hypothèses, et du fait que l'ensemble $\emptyset$ est sous-ensemble de n'importe quel ensemble, en particulier de Γ , nous avons que $\Gamma \vdash_{\mathcal{L}} A \rightarrow (A \vee B)$ où dans Γ nous pouvons avoir des énoncés également faux pour lesquelles l'axiome $\vdash_{\mathcal{L}} A \rightarrow (A \vee B)$ se maintient encore vrai. En d'autres mots, si $\vdash_{\mathcal{L}} A \rightarrow (A \vee B)$ est démontrable à partir d'un ensemble vide de prémisses, alors $\Gamma \vdash_{\mathcal{L}} A \rightarrow (A \vee B)$ est démontrable à partir de n'importe quel ensemble de prémisses. Cependant une règle du type $A \vdash_{\mathcal{L}} A \vee B$, par exemple, ne fonctionne que si l'hypothèse A à gauche de " $\vdash$ " est toujours vraie - et cela vaut (par définition de la déductibilité, ou de ce qu'est une règle d'inférence) pour n'importe quelle règle d'inférence, en particulier pour ce qui est des règles d'inférence de notre logique **P**. Dans $A, B \vdash_{\mathbf{P}} A \wedge B$ (règle de conjonction), il faut que les hypothèses A et B à gauche de " $\vdash$ " soient toujours vraies, et de même pour $A, A \rightarrow B \vdash_{\mathbf{P}} B$ (*modus ponens*). Enfin, les règles d'inférence sont correctes, c'est-à-dire, elles propagent la validité: ce qui doit être transmis des prémisses à la conclusion, c'est la *vérité*. Ainsi, si des prémisses comme $A \rightarrow B$ et A sont vraies, la conclusion B l'est aussi.

2- **La définition de degrés de pertinence chez Sperber et Wilson versus la définition de degrés de pertinence dans la Logique *P*.**

En prenant pour base les définitions du cadre théorique de Sperber et Wilson, nous discuterons maintenant leur définition de la pertinence présentée dans la section 3, Chapitre I.

Comme nous l'avons vu dans le Chapitre I, il y a deux contributions majeures que font Sperber et Wilson dans leur théorie de la pertinence visant à expliquer le processus d'inférence de l'esprit humain: le *Principe de Pertinence* selon lequel "tout acte de communication ostensive communique la présomption de sa propre pertinence optimale", et la notion de *contexte* qui tient compte non seulement des informations contenues dans les énonciations prononcées préalablement, mais de tout un ensemble de prémisses et présomptions à propos du monde ambiant et que l'interlocuteur utilise pour interpréter les énonciations qui contextuellement entraînent, par exemple, des implicatures.

La pertinence est caractérisée en termes d'effets contextuels. Sperber et Wilson soulignent que si une hypothèse n'a pas d'effets contextuels dans un contexte donné, elle n'est pas pertinente dans ce contexte - puisque le fait d'avoir des effets contextuels dans un contexte est une condition

nécessaire, mais non suffisante de la pertinence d'une énonciation dans un contexte. Pour eux, les intuitions de pertinence qu'il est essentiel d'expliquer portent surtout sur des degrés de pertinence plutôt que sur la présence ou l'absence de pertinence. Ainsi dans leur cadre théorique, la pertinence est traitée comme une notion comparative d'effet (directement proportionnelle) et d'effort (inversement proportionnelle) et qui est mesurée de la façon suivante: plus l'effet produit par une hypothèse dans un contexte donné est grand, plus elle est pertinente; moins l'effort de traitement de cette hypothèse dans un contexte est grand, plus elle est pertinente.

Voyons dans Sperber et Wilson (1989, p. 192) l'exemple suivant qui illustre cette notion comparative de pertinence. Soit dit en passant, il s'agit d'un exemple qui produit un seul type d'effet contextuel, soit des implications contextuelles.

Considérez un contexte constitué des hypothèses (1 a-c):

(1) a) Un homme et une femme qui vont se marier doivent consulter un médecin sur les maladies héréditaires qu'ils risquent de transmettre à leurs enfants.

b) Un homme et une femme atteints de thalassémie doivent éviter d'avoir des enfants ensemble.

c) Suzanne est atteinte de thalassémie.

Considérez les effets que les hypothèses (2) et (3) (dont Sperber et Wilson stipulent qu'elles ont la même force) auraient dans ce contexte:

(2) Suzanne, qui est atteinte de thalassémie, va épouser Pierre.

(3) Pierre, qui est atteint de thalassémie, va épouser Suzanne.

Comme le remarquent Sperber et Wilson, les hypothèses (2) et (3) ont chacune certains effets contextuels dans le contexte (1a-c) et sont donc pertinentes selon leur définition de la pertinence (Chapitre I de notre thèse):

La pertinence

Condition comparative 1: une hypothèse est d'autant plus pertinente dans un contexte donné que ses effets contextuels y sont plus importants.

Condition comparative 2: une hypothèse est d'autant plus pertinente dans un contexte donné que l'effort nécessaire pour l'y traiter est moindre. (Sperber et Wilson, 1989, p. 191)

En particulier, (2) et (3) entraînent chacune l'implication contextuelle (4).

(4) Suzanne et Pierre devraient consulter un médecin sur les maladies héréditaires qu'ils risquent de transmettre à leurs enfants.

La présence d'un effet contextuel peut expliquer, soutiennent Sperber et Wilson, notre intuition selon laquelle les deux hypothèses (2) et (3) sont pertinentes dans ce contexte.

Mais, d'après eux, notre intuition suggère en outre que, dans ce contexte, (3) est plus pertinente que (2). Ils affirment que c'est la définition de la pertinence qui permet d'expliquer ce second aspect de notre intuition. Dans le contexte (1a-c), (3) possède en effet une implication contextuelle de plus que (2), à savoir (5).

(5) Suzanne et Pierre devraient éviter d'avoir des enfants.

Quant à l'effort de traitement des hypothèses (2) et (3) dans le contexte (1a-c), Sperber et Wilson affirment qu'un supplément d'effort est nécessaire pour traiter (3) laquelle entraîne l'implication contextuelle (5). Vu que (2) n'entraîne pas cette implication, un tel supplément d'effort n'est pas nécessaire.

Chaque effet contextuel nécessite ainsi un supplément d'effort. Si les bénéfices tirés d'un effet contextuel ne dépassaient jamais le coût du supplément d'effort nécessaire pour réaliser cet effet, alors on ne pourrait jamais atteindre un degré positif de pertinence. Il ne vaudrait pas la peine de penser.

Pourtant, sauf lorsqu'ils sont totalement épuisés, les humains pensent. La conclusion empirique qui s'impose est que l'effort de traitement nécessaire

pour inscrire une implication contextuelle ou pour augmenter ou diminuer la force d'une hypothèse n'est pas coûteux au point d'annuler la contribution de ce processus à la pertinence. En outre, puisque ce supplément d'effort est toujours proportionné aux effets qui le rendent nécessaire, on peut simplement l'ignorer dans l'évaluation de la pertinence. L'esprit ne se soucie sans doute que des efforts qui peuvent être évités. Nous ferons de même et ne tiendrons compte que de l'effort de traitement qui a pour résultat un effet contextuel; nous ignorerons l'effort supplémentaire qui résulte lui-même du fait qu'un effet contextuel a été produit. (Sperber et Wilson, 1989, p. 193)

C'est pourquoi Sperber et Wilson suggèrent que les hypothèses (2) et (3) demandent le même effort de traitement au moment de les traiter dans le même contexte. Étant donné que (3) a plus d'effets contextuels que (2) dans le contexte (1 a-c), leur définition de pertinence prédit que (3) doit être plus pertinente.

Nous voulons maintenant attirer l'attention sur le fait suivant. Bien que Sperber et Wilson stipulent que les deux hypothèses (2) et (3) ont la même force (elles ont la même structure conceptuelle, et autorisent l'application des mêmes règles d'inférence du dispositif déductif), nous ne pouvons pas négliger le fait que par rapport au contexte (1a-c), l'hypothèse (2) présente une seule information nouvelle, à savoir, "Suzanne va épouser Pierre" tandis que l'hypothèse (3) a deux informations nouvelles, à savoir, "Pierre est atteint de thalassémie" et "Pierre va épouser Suzanne".

Cela étant dit, nous sommes tentés de dire que le fait que (3) présente une information de plus que (2), bien entendu, une information (ou

hypothèse) traitable¹⁹ dans le contexte (1a-c), permet d'expliquer la présence d'un effet contextuel de plus dans ce contexte, à savoir, l'implication contextuelle (5).

Il nous semble donc que plus informative est une hypothèse traitable dans un contexte donné, plus elle y entraîne d'effets contextuels. Autrement dit, plus une hypothèse contient d'informations nouvelles relativement au contexte où elle est traitée, plus elle entraînera d'effets contextuels dans ce même contexte donné.

Ce point de vue cependant soulève un problème en ce qui concerne la définition d'évaluation de pertinence de Sperber et Wilson (1989). D'après eux, une telle définition permet d'expliquer leur intuition selon laquelle "toutes choses étant égales d'ailleurs, une hypothèse ayant plus d'effets contextuels est plus pertinente; et, toutes choses étant égales d'ailleurs, une hypothèse demandant moins d'effort de traitement est plus pertinente" (p. 191). Pouvons-nous conclure qu'une hypothèse contenant plus d'informations a plus d'effets contextuels et, par conséquent, selon la définition de degrés de pertinence de Sperber et Wilson, qu'elle est plus pertinente? Remarquons que, par un

¹⁹ *Information Traitable* dans le sens utilisé par Sperber et Wilson, c'est-à-dire, dans le sens où l'information ("... pas nécessairement nouvelle pour l'organisme, elle peut être simplement de l'information qui est traitée à nouveau." p. 167), ajoutée à un ensemble d'hypothèses contextuelles, entraîne des implications contextuelles qui améliorent la représentation du monde de l'individu, modifient ou affaiblissent la force de l'information.

enchaînement d'inférences, cela revient à dire qu'une hypothèse plus informative est plus pertinente! Nous reviendrons sur ce sujet prochainement dans la troisième section de ce chapitre.

Comme nous pouvons le voir, l'exemple précédent est un cas typique de dimension littérale du sens. Avant de donner suite à l'analyse que nous présenterons bientôt, nous voulons développer quelques exemples liés à la non-littéralité, en particulier pour les implications conversationnelles. Soulignons que ces deux dimensions du sens, c'est-à-dire, la communication littérale et la communication non littérale sont respectivement dénommées par Sperber et Wilson (1989) *explicitation* et *implication*. Ainsi, ils définissent la communication explicite:

(6) L'*explicite*

Une hypothèse communiquée par un énoncé *U* est *explicite* si et seulement si elle résulte du développement d'une forme logique codée par *U*. (p. 271)

En d'autres mots, une hypothèse est une *explicitation* quand elle est explicitement communiquée. La forme logique de l'énoncé correspond, selon Sperber et Wilson, à certaines tâches qui doivent être accomplies par l'auditeur telles que: attribuer à l'énoncé une forme propositionnelle unique, ce

Cela revient à dire qu'il faut que l'information ajoutée soit à son tour pertinente.

qui pourra être fait, par exemple, par une désambiguïsation de la phrase énoncée (laquelle permet de sélectionner l'une des représentations sémantiques), ou par l'assignation de référents aux anaphoriques et aux déictiques; reconnaître l'attitude propositionnelle du locuteur; en somme, l'auditeur doit enrichir et compléter une représentation sémantique sélectionnée, au moyen d'informations contextuelles afin d'obtenir la forme propositionnelle exprimée par l'énoncé et l'incorporer dans un schéma d'hypothèse extrait de sa mémoire encyclopédique - schéma qui exprime une attitude par rapport à la forme propositionnelle qu'il incorpore. Et toutes ces tâches ne sont accomplies que par des processus inférentiels.

Ainsi, du fait qu'une explicitation n'est que le résultat du développement d'une forme logique d'un énoncé, elle n'est pas conforme au sens littéral de cet énoncé. Sperber et Wilson maintiennent qu'une explicitation est une combinaison de propriétés conceptuelles soit linguistiquement codées ou contextuellement inférées. "Plus grande est la part des propriétés codées, plus l'explicitation est explicite; plus grande est la part des propriétés inférées, moins l'explicitation est explicite" (p. 271).

D'autre part, les *implicitations* n'ont rien à voir avec le développement de la forme logique de l'énoncé. Elles ne tiennent compte que des hypothèses qui sont nécessaires pour obtenir une interprétation cohérente

avec le principe de pertinence. Sperber et Wilson soutiennent que le processus de reconstitution des *implications* d'un énoncé correspond à la reconnaissance des raisons manifestes qui ont amené le locuteur à penser que son énoncé serait optimalement pertinent pour l'auditeur. "Toute hypothèse qui est communiquée, mais pas explicitement, est communiquée implicitement: c'est une *implication*" (p. 271). Pour eux, il n'existe que deux types d'implication, à savoir, les *prémisses implicites* et les *conclusions implicites*.

Considérons l'exemple suivant:

- (7) (a) Henri: Est-ce que tu aimerais manger du boeuf ?
 (b) Hélène: Je ne mange pas de viande rouge.

Selon Sperber et Wilson en prenant pour point de départ des formes propositionnelles et des attitudes propositionnelles exprimées par un énoncé ainsi que le contexte, il est possible d'inférer toutes les explicitations de l'énoncé: la forme propositionnelle, étant l'explicitation la plus importante, permet de déclencher la plupart des effets contextuels de l'énoncé, et conséquemment de montrer sa pertinence. Dans ces circonstances, l'énoncé (7b) d'Hélène est une assertion ordinaire dont la principale explicitation (celle qui intéresse Henri) est sa forme propositionnelle. Bien que le fait que cette

forme propositionnelle ne réponde pas directement (ou explicitement) à la question (7a), Henri retient l'hypothèse que Hélène veut lui rendre manifeste une réponse contextuellement impliquée par son énoncé (7b) - sinon Hélène ne pourrait pas penser que son énoncé puisse être pertinent pour Henri. Dès lors, pour interpréter l'énoncé (7b), Henri utilise des schémas d'hypothèses extraits de sa mémoire encyclopédique (rappelons que le savoir encyclopédique d'un individu est sa représentation globale du monde), lesquels sont mutuellement manifestes à lui et à Hélène. Henri utilise en particulier, le schéma d'hypothèse suivant:

(8) La viande de bœuf est rouge.

En y ajoutant un tel schéma, l'explicitation (ou forme propositionnelle) de l'énoncé (7b) devient:

(9) Hélène ne mange aucune viande rouge,

Henri peut immédiatement inférer l'implication contextuelle suivante:

(10) Hélène n'aimerait pas manger du bœuf.

Ainsi, (8) est une prémisses implicitee, (10) est une conclusion implicitee. Selon Sperber et Wilson, une implicitee est une hypothese appartenant au contexte ou contextuellement impliquee. C'est cette hypothese la que le locuteur escomptait probablement rendre manifeste à l'auditeur par son desir de rendre son énoncé manifestement pertinent. Alors que les prémisses implicitees renvoient aux hypotheses que doit adopter l'interlocuteur afin de parvenir à une interpretation cohérente avec le principe de pertinence, les conclusions implicitees sont deduites à partir des explicitations de l'énoncé et du contexte, et elles ont la propriété de ne pas être entièrement determinées. C'est à ces conclusions (ou à certaines d'entre elles) que le locuteur escompte que l'interlocuteur aboutira s'il veut que son énoncé soit manifestement pertinent. D'autres prémisses implicitees comme (11) et (12) associees à l'explicitation (9) produisent respectivement des conclusions implicitees comme (13) et (14).

(11) La viande de mouton est rouge.

(12) La viande de cheval est rouge.

(13) Hélène n'aimerait pas manger du mouton.

(14) Hélène n'aimerait pas manger du cheval.

Toutes ces considérations ont, en fin de compte, l'objectif d'introduire notre discussion à propos du critère d'évaluation du degré de pertinence de Sperber et Wilson. Cependant, nous pouvons construire quelques interprétations intuitives qui ne satisfont pas à ce critère et nous sommes dès lors amenés à traiter de degrés de pertinence par le biais d'un autre critère - celui qui a été présenté dans le chapitre précédent où nous avons proposé nos définitions techniques de la pertinence ainsi qu'une logique de la pertinence *P*.

Considérons l'échange suivant:

(15) Un médecin traite deux patientes qui viennent d'arriver à l'hôpital. Il doit leur donner une injection qu'elles ne peuvent pas recevoir si elles sont enceintes. Alors, il pose la question:

Vous êtes enceintes?

Les deux femmes: Non!

Le médecin: Vous en êtes sûres?

(a) Louise: J'ai subi une ligature des trompes.

(b) Claire: Mon mari a subi une vasectomie.

Pour comprendre les réponses de Louise et de Claire, le médecin devrait tout d'abord attribuer à chaque énoncé une forme propositionnelle unique. Selon Sperber et Wilson, tout énoncé qui communique sa forme

propositionnelle constitue une affirmation ordinaire. Puisque, (15a) et (15b) sont des assertions ordinaires, la principal explicitation de chacune qui intéresse le médecin est la forme propositionnelle. Voyons comment l'analyse de Sperber et Wilson expliquerait le processus d'interprétation des deux réponses.

Pour ce qui est de la réponse de Louise (15a), considérons que la forme propositionnelle communiquée par son énoncé est l'explicitation (16):

(16) Louise a subi une ligature des trompes.

Or, cette forme propositionnelle ne répond pas explicitement à la question du médecin. Dans ce cas, pour interpréter l'énoncé (15a), le médecin doit en outre faire appel à des connaissances dites encyclopédiques (c'est-à-dire, à des hypothèses contextuelles accessibles en mémoire) qui sont mutuellement manifestes à lui et à Louise. Le médecin peut accéder en particulier, au schéma d'hypothèse (17):

(17) Une femme qui a subi une ligature des trompes, n'est plus capable d'enfanter.

En ajoutant à un tel schéma l'explicitation (ou forme propositionnelle de 15a) (16), le médecin peut immédiatement inférer l'implication contextuelle (18) :

(18) Louise n'est plus capable d'enfanter.

Supposons maintenant qu'à partir de (18) et d'autres informations mutuellement manifestes, le médecin soit en mesure d'inférer l'implication contextuelle (19) dans la tentative d'arriver à une interprétation cohérente avec le principe de pertinence:

(19) Louise est sûre de ne pas être enceinte.

Supposons encore qu'il soit mutuellement manifeste que c'est l'implication contextuelle (19) qui rend l'énoncé de Louise suffisamment pertinent pour qu'il mérite d'être traité par le médecin. Dans ce cas, (19) fait manifestement partie de l'ensemble des hypothèses communiquées par Louise, soit celles qu'elle avait l'intention de communiquer.

Pour ce qui est de la réponse de Claire (15b), considérons que la forme propositionnelle communiquée par son énoncé est (20):

(20) Le mari de Claire a subi une vasectomie.

De même que l'exemple précédent, cette forme propositionnelle ne répond pas directement à la question du médecin. Dans ce cas, le médecin peut

utiliser le schéma d'hypothèse extrait de sa mémoire (21):

(21) Un homme qui a subi une vasectomie n'est plus capable de faire un enfant à aucune femme,

lequel, associé à l'explicitation (20), produit l'implication contextuelle (22):

(22) Le mari de Claire n'est plus capable de lui faire un enfant.

Or, la tâche du médecin est de décider quelles hypothèses rendues manifestes par l'énoncé de Claire sont telles qu'il est mutuellement manifeste que Claire a l'intention de les rendre manifestes. Supposons alors, qu'à partir de (22) et d'autres informations mutuellement manifestes, le médecin soit au point d'inférer l'implication contextuelle (23):

(23) Claire est sûre de ne pas être enceinte.

Supposons maintenant qu'il soit mutuellement manifeste que c'est juste cette implication contextuelle qui rend la réponse de Claire suffisamment pertinente pour mériter l'attention du médecin. Si tel est le cas, (23) fait manifestement partie de l'ensemble des hypothèses qui sont communiquées par

l'énoncé de Claire, hypothèses qu'elle avait l'intention de rendre mutuellement manifestes.

Maintenant, qu'en est-il de l'ensemble des hypothèses rendues manifestes par les énoncés de Louise et de Claire? Lesquelles parmi toutes ces hypothèses, Louise ou Claire a l'intention de rendre mutuellement manifestes? Et puisque "les intuitions de pertinence qu'il est essentiel d'expliquer ne portent pas tant sur la présence ou l'absence de pertinence que sur des degrés de pertinence" (Sperber et Wilson, 1989, p. 188), peut-on tirer encore d'autres implications contextuelles dans le contexte (15a-b) pour ensuite comparer le degré de pertinence des deux hypothèses (15a) et (15b)?

De même que dans la communication non verbale, étant donné l'ensemble H des hypothèses rendues manifestes par le communicateur, il n'est jamais possible d'identifier à coup sûr, parmi toutes ces hypothèses, le sous-ensemble I des hypothèses communiquées par le communicateur, c'est-à-dire, celles qu'il avait l'intention de communiquer, de même pour la communication verbale, étant donné l'ensemble H des hypothèses rendues manifestes par l'énoncé du locuteur, il n'est jamais possible d'identifier à coup sûr parmi toutes ces hypothèses, le sous-ensemble I des hypothèses communiquées par l'énonciation du locuteur. En bref, I ne peut être défini qu'en termes généraux. Comme nous venons de le voir, aux yeux de Sperber et

Wilson, la description du processus de communication commence par la description linguistique d'un énoncé, laquelle est toujours subordonnée à la grammaire et non pas aux intérêts de l'auditeur. Dans la communication linguistique on peut compter sur un degré d'explicite qui nous permet d'énumérer toutes les hypothèses explicitement véhiculées par un énoncé (ce qui n'est pas toujours possible par rapport à la communication verbale dont les hypothèses sont souvent implicitement véhiculées). En second lieu, la description linguistique fournit, pour chaque sens de la phrase énoncée, un éventail de représentations sémantiques qui à leur tour doivent être complétées et intégrées dans une hypothèse sur l'intention informative du locuteur. Et c'est là le point plus délicat. Qui plus est, chaque sens schématique différent des autres peut être complété de plusieurs façons aussi différentes les unes des autres.

En ce qui touche à la possibilité de trouver d'autres implications contextuelles produites par les énoncés (15a) et (15b), cela nous permettrait de comparer le degré de pertinence de la réponse de Louise et celui de la réponse de Claire. Alors, ainsi que Sperber et Wilson l'ont remarqué pour ce qui est des hypothèses précédentes (2) et (3) pour le contexte (1a-c), de même nous observons que les hypothèses (15a) et (15b) ont la même structure conceptuelle, les mêmes propriétés conceptuelles, les unes linguistiquement

codées, les autres contextuellement inférées, et autorisent, en raison de cela, l'application des mêmes règles d'inférence du dispositif déductif. L'une et l'autre ont chacune quelques effets contextuels dans le contexte (15) et, conséquemment, selon la définition de la pertinence de Sperber et Wilson, elles sont pertinentes. En particulier, (15a) et (15b) entraînent respectivement les implications contextuelles (19) et (23):

(19) Louise est sûre de ne pas être enceinte.

(23) Claire est sûre de ne pas être enceinte.

Dans cet ordre d'idées, étant donné que la pertinence est traitée comme une notion comparative d'effet et effort, qu'en est-il du degré de pertinence entre les réponses de Louise et de Claire? Pourrions-nous conclure que dans le contexte (15), (15a) est plus pertinente que (15b) ou, au contraire, que (15b) est plus pertinente que (15a)?

Notre intuition ne suggère pas, dans le contexte (15), que (15a) soit plus pertinente que (15b) ou vice-versa. À première vue, nous pourrions penser que, en termes de degré, (15a) serait (intuitivement) aussi pertinente que (15b). Cependant, l'analyse de Sperber et Wilson impose que dans ce contexte, (15b) soit plus pertinente que (15a). En effet, dans (15), (15b) produit, en particulier, une implication contextuelle de plus que (15a), soit:

(24) Claire est fidèle à son mari.

Telle conclusion implicite peut être inférée par le médecin à partir de (23) et d'autres informations contextuelles qui sont mutuellement manifestes à lui et à Claire. Évidemment, si Claire communique contextuellement qu'elle est sûre de ne pas être enceinte du seul fait que son mari a subi une vasectomie, elle rend dès lors mutuellement manifeste l'hypothèse qu'elle est fidèle à son mari. Ainsi, (24) fait manifestement partie de l'ensemble des hypothèses qui sont communiquées par son énoncé. Et quant à l'effort pour inscrire l'implication contextuelle (24), puisque ce supplément d'effort est proportionné à l'effet qui le rend nécessaire, alors selon Sperber et Wilson, on peut l'ignorer dans l'évaluation de la pertinence.

Mais, comment la simple production d'une implication contextuelle de plus, peut-elle rendre la réponse de Claire plus pertinente que la réponse de Louise? Pourrions-nous dire que le critère de degré de pertinence est toujours valable, en particulier dans notre exemple (15a-b)? Nous ne partageons pas cette idée. Tout ce que nous admettons est seulement que (15b) communique plus que (15a) du fait qu'il serait inconséquent de dire que dans (15) Claire ne communique pas (24). Or, c'est justement cette implication qui fonde la

réponse indirecte de Claire, quoique dans ce contexte-là, Claire communique plus faiblement (24) que (23). Pour des raisons liées au sujet du contexte (15) où le médecin pose la question de savoir si Claire est sûre de ne pas être enceinte, nous tendons à conclure que Claire communique plus fortement (23) que (24). Rappelons que d'après Sperber et Wilson, les implications d'une hypothèse correspondent à l'ensemble des hypothèses indispensables pour obtenir une interprétation cohérente avec le principe de pertinence. Dans ce cas, (24) est nécessaire pour l'interprétation de (23). Cependant, il ne nous semble pas que le seul fait que (15b) ait l'implication (24) en plus que (15a) soit une condition suffisante pour considérer (15b) plus pertinente. Au contraire, nous croyons que Louise et Claire, l'une et l'autre, ont fait de leur mieux pour présenter au médecin la "meilleure" réponse, réponse qui lui permettrait de chercher de façon la plus économique possible ce qu'il voulait savoir: si les deux femmes étaient sûres de ne pas être enceintes. Nous insistons encore sur ce point; ce qui est tout à fait conforme à notre intuition c'est juste le fait que la réponse de Louise semble aussi intuitivement pertinente que la réponse de Claire. En d'autres mots, l'une et l'autre réponse permettent d'inférer, disons dans le même temps de compréhension, les implications contextuelles (19) et (23) qui sont exactement ce que voulaient communiquer les réponses de Louise et de Claire.

Supposons d'ailleurs que, tout comme Louise, Claire ait subi une ligature des trompes et qu'elle ait répondu au médecin de la même façon. Sperber et Wilson diraient que dans ce cas, la réponse de Claire serait, dans ce contexte-là, aussi pertinente que celle de Louise du fait qu'on en obtiendrait exactement les mêmes effets contextuels. Maintenant, si l'on retient les deux hypothèses à la fois, à savoir, "Claire a subi une ligature des trompes" et "Le mari de Claire a subi une vasectomie", alors laquelle entre les deux, Claire devrait-elle choisir pour répondre au médecin de façon optimalement pertinente? Rappelons que selon la théorie de la pertinence de Sperber et Wilson, étant donné le principe de pertinence, Claire devrait choisir la réponse qui véhiculera la conclusion prétendue de façon la plus économique possible (du point de vue du coût de traitement).

Mais comment mesurer dans ce cas en particulier, l'effort de traitement entre les deux réponses? Comme nous venons de le voir, dans notre exemple (15), (15b) entraîne une implication contextuelle de plus que (15a), à savoir (24). Or, de même que dans le contexte (1a-c) - voir l'exemple tout au début de ce chapitre - l'hypothèse (3) est plus pertinente que (2) puisqu'elle permet d'inférer une implication contextuelle de plus que (2), et dont le supplément d'effort proportionné à cette implication a été ignoré par les auteurs dans l'évaluation de la pertinence, de même nous suggérons que les hypothèses

(15a) et (15b) demandent (pour les mêmes raisons déjà présentées) le même effort de traitement dans le contexte (15). Dans ce cas, de même que (15a) et (15b), l'une ou l'autre réponse de Claire véhiculerait la conclusion prétendue de façon (également) plus économique, c'est-à-dire, l'une ou l'autre réponse confirmerait la condition comparative 2 de la définition de pertinence. Désormais, il ne resterait à Claire qu'à se préoccuper de la condition comparative 1 de pertinence et finalement décider pour la réponse "Mon mari a subi une vasectomie" (puisque'elle a une implication contextuelle en plus), réponse que confirmerait aussi la deuxième partie de la présomption de pertinence optimale, c'est-à-dire, réponse qui serait la plus pertinente des deux réponses que Claire pouvait utiliser pour communiquer qu'elle était sûre de ne pas être enceinte. Dans le fond, c'est ça que le critère de degrés de pertinence tel que proposé par Sperber et Wilson nous impose. Toutefois, pour nous, le choix d'une telle réponse comme étant la plus pertinente pêche contre notre intuition. Peut-être faudrait-il expliquer un peu mieux ce que sont de tels jugements (intuitifs) comparatifs. Sperber et Wilson soutiennent que le supplément d'effort de traitement nécessaire soit pour inscrire une implication contextuelle soit pour augmenter ou diminuer la force d'une hypothèse n'est pas coûteux au point d'annuler la contribution de ce processus à la pertinence et que pour cette raison, on peut l'ignorer. Nous sommes d'accord que cela

paraît être très conforme à notre intuition. Par contre, leur affirmation que l'esprit ne se soucie que des efforts qui peuvent être évités ne colle pas très bien à notre exemple. Selon nous, considérer la réponse "Mon mari a subi une vasectomie" comme étant plus pertinente que la réponse "J'ai subi une ligature des trompes", c'est admettre d'une certaine façon qu'en choisissant la première réponse, Claire ne se soucie pas des efforts qui peuvent être évités. Or, si Claire a le choix entre les deux réponses, celle qui nous semble accéder plus immédiatement à l'implication contextuelle prétendue c'est la deuxième réponse "J'ai subi une ligature de trompes" car elle évite un supplément d'effort de traitement nécessaire pour accéder à l'implication en question. Mais, d'après nous, ce n'est pas suffisant pour évaluer la pertinence. Pour nous il serait essentiel d'expliquer qu'en choisissant la deuxième réponse, Claire non seulement envisagerait de raccourcir le temps de compréhension du médecin mais aussi elle éviterait de produire chez le destinataire d'autres implications contextuelles outre celles qu'elle veut lui communiquer.

Qu'est-ce donc, pour une hypothèse, que d'être plus pertinente qu'une autre? Dans une conversation qu'est-ce qui est le plus important: le résultat de l'inférence ou le nombre d'inférences?

Par *résultat de l'inférence*, nous comprenons simplement la reconnaissance par le destinataire de l'intention du communicateur de vouloir

produire seulement des effets contextuels qu'il veut lui communiquer par l'énoncé qu'il utilise. Si à son tour, le destinataire infère exactement les implications contextuelles prétendues par le communicateur, alors parfait! Évidemment, il pourrait y avoir d'autres effets contextuels produits par le(s) énoncé(s) du communicateur chez le destinataire et qui ne seraient pas de la connaissance du communicateur, ou en d'autres mots, qui ne lui seraient pas manifestes. Mais de notre point de vue, ces effets-là ne doivent pas compter dans l'évaluation de la pertinence; au contraire l'on va se heurter aux mêmes problèmes que nous venons de mentionner dans l'exemple précédent.

Comme Sperber et Wilson eux-mêmes le soutiennent, la définition de la pertinence serait insuffisante si elle était définie comme une propriété purement formelle, sans décrire ses rapports à la réalité psychologique. D'après eux, il faudrait d'abord tenir compte de la relation de pertinence entre une hypothèse et un contexte et ensuite des degrés de pertinence entre des hypothèses relativement à un contexte donné. Cela étant dit, il nous semble que ces deux conditions ne sont pas toujours satisfaites comme on peut bien le voir dans notre dernier exemple.

À notre avis, les intuitions de pertinence qu'il est essentiel d'expliquer doivent porter sur le résultat de l'inférence à la fin du processus d'interprétation d'une énonciation (et non sur les inférences que l'auditeur peut

faire et qui n'appartiennent pas au contenu que le locuteur voulait effectivement communiquer). Ainsi, conformément à notre définition (1.8) de degrés de pertinence (Chapitre III), tant la réponse (15a) de Louise que la réponse (15b) de Claire ont le degré 1, ou degré maximum, de pertinence pour le médecin, c'est-à-dire, elles permettent au médecin d'inférer exactement la conséquence pragmatique que Louise avait l'intention de lui communiquer ainsi que la conséquence pragmatique que Claire avait l'intention de produire chez lui. Évidemment dans une conversation, le locuteur rationnel doit faire des énonciations qui soient en même temps appropriées mais aussi pertinentes²⁰ par rapport à cette conversation afin de permettre à l'allocataire d'inférer précisément ce qu'il a l'intention de lui communiquer (dans la plupart des cas, de la façon la plus économique possible, mais pas toujours: en effet, il se peut que le communicateur, par son énonciation, signifie beaucoup en disant peu); le cadre théorique de Sperber et Wilson impose que dans le processus d'interprétation d'une énonciation, plus grand est l'effort du destinataire pour la comprendre, moins grande est la pertinence de cette énonciation. Cependant, nous observons que conformément au contexte, le locuteur a pour but de produire linguistiquement une conséquence pragmatique chez l'interlocuteur et

²⁰ Nous discuterons dans le chapitre VI la distinction entre la notion d'appropriété et la notion de pertinence.

ce de façon minimale et suffisante. Considérons le passage suivant, dans le fameux *Leviathan* - I, chapitre III, de Hobbes:

Ainsi, dans une conversation sur notre actuelle guerre civile, rien ne paraissait plus impertinent que de demander (comme c'est arrivé effectivement) quelle était la valeur d'une pièce de monnaie romaine? Néanmoins, à mes yeux, la pertinence était assez manifeste, car la pensée de la guerre conduit à la pensée du roi livré à ses ennemis; cette pensée conduit à son tour à celle du Christ livré aux Romains; et celle-ci, à son tour, fait penser aux trente pièces de monnaie qui furent le prix de la trahison: et de là suit facilement la question malicieuse. Et tout cela en un bref instant, car la pensée est très rapide. (Hobbes, 1974, p.20-21 - notre traduction)

Comme nous pouvons le voir dans cet exemple, l'interlocuteur a dû étendre le contenu de la question du locuteur en y ajoutant des éléments pour obtenir une interprétation qu'il croît être consistante avec l'intention informative du locuteur.

Mais, que diraient Sperber et Wilson au sujet du degré de pertinence de cette énonciation? Que, dans le cas présent, le locuteur devrait avoir choisi une énonciation alternative qui permettrait au destinataire d'interpréter sa demande de façon moins coûteuse? Nous ne partageons pas cette idée. Il faut remarquer que les conversations ne sont pas toujours orientées vers le coût de traitement des hypothèses du locuteur par rapport à son interlocuteur: nous soutenons que le but propre des conversations tourne autour du rendement, c'est-à-dire, du produit, du résultat des inférences que

les hypothèses du communicateur peuvent entraîner chez l'auditeur à partir de ses intentions informatives de communiquer de telles hypothèses. Un homme qui vient de connaître une dame, par exemple, et qui tombe amoureux d'elle, peut avoir l'intention de lui communiquer sa passion, mais pas de façon si immédiate. Alors, dans sa conversation, il peut lui donner des indices à ce propos de plusieurs façons, soit en l'invitant à dîner, soit à aller au cinéma voir un film romantique, etc.

Pour ce qui est de notre définition de la pertinence, nous concluons que la question (rhétorique - car elle n'est pas une véritable demande d'information) "quelle était la valeur d'une pièce de monnaie romaine?", posée dans le contexte de cette conversation-là, est pertinente car son énonciation entraîne des conséquences pragmatiques dans ce contexte. Qui plus est, dans notre approche, l'exemple ci-dessus résume tout ce que l'allocutaire peut tenir pour le degré maximum de pertinence (tel que nous l'avons défini dans le chapitre III) car il a réussi à inférer exactement ce que le locuteur a eu l'intention de lui communiquer.

3- Convergence entre la seconde clause de la présomption de pertinence optimale chez Sperber et Wilson et la maxime de quantité de Grice.

En retournant maintenant à l'approche de Sperber et Wilson, qu'en est-il de notre observation (dans la section précédente) selon laquelle, il nous semble qu'une hypothèse véhiculant plus d'information entraîne plus d'effets contextuels (et pour cette raison, selon Sperber et Wilson, elle serait plus pertinente) ?

Comme on l'a vu dans la section 1, Chapitre I, selon Grice les contributions des locuteurs dans le déroulement de la conversation sont gouvernées par le principe de coopération, auquel quatre types de maximes conversationnelles sont subordonnées.

Mais, Sperber et Wilson soutiennent qu'il serait commode de remplacer le principe de coopération et les maximes conversationnelles de Grice par un principe de pertinence plus explicite que ne le sont les principes conversationnels de Grice: du fait que l'esprit humain est orienté vers la pertinence, un acte de communication ostensive véhicule une garantie de sa propre pertinence. Chez Grice cette garantie peut être décrite en termes de présomption que les interlocuteurs respectent de telles normes

conversationnelles. Comme le soutiennent Sperber et Wilson, le degré de coopération nécessaire à la communication est plus grand chez Grice que pour eux. “Il est donc plus facile d’être optimalement pertinent que de respecter les maximes de Grice” (Sperber et Wilson, 1989, p. 243); par rapport à la maxime de quantité de Grice, on peut être optimalement pertinent sans être aussi informatif que le demandent les objectifs de l’échange verbal en cours: “si, par exemple, on garde pour soi des informations dont la connaissance serait pertinente aux interlocuteurs” (idem, p. 243). Dans ce cas, les interlocuteurs eux-mêmes ne peuvent pas toujours compter sur le degré de coopération décrit par Grice. Qui plus est, d’après eux, la présomption que les interlocuteurs respectent de telles normes conversationnelles ne sert qu’à inférer les implications véhiculées dans l’échange.

À la différence de l’approche de Grice, le modèle inférentiel fondé par Sperber et Wilson sur le principe de pertinence tient compte selon eux de tous les niveaux de l’interprétation, c’est-à-dire non pas seulement les implications de l’énoncé, mais aussi son explicitation (l’enrichissement de sa forme logique).

L’analyse de la conversation présentée par Grice ne propose, selon Sperber et Wilson, aucune explication de la communication explicite, et quant aux implications, il les explique comme des hypothèses que le destinataire doit

construire pour préserver l'idée que le locuteur respecte le principe de la coopération. Par contre, leur principe de pertinence vise à expliquer la communication ostensive dans sa totalité, qu'elle soit explicite ou implicite.

Comme nous venons de le remarquer, Sperber et Wilson soulignent qu'il est plus commode d'être optimalement pertinent que de respecter les maximes de conversation de Grice, en particulier, sa maxime de quantité²¹. Mais comme l'a bien montré Dominicy (1991), la seconde clause de la présomption de pertinence optimale, chez Sperber et Wilson, n'est qu'une version plus précise des deux maximes de quantité suivantes de Grice (1975): (1) que votre contribution soit aussi informative que nécessaire; (2) que votre contribution ne soit pas plus informative que nécessaire.

Revenons à la définition de *présomption de pertinence optimale*, proposée par Sperber et Wilson (1989):

- (a) L'ensemble d'hypothèses I que le communicateur veut rendre manifestes au destinataire est suffisamment pertinent pour que le stimulus ostensif mérite d'être traité par le destinataire.
- (b) Le stimulus ostensif est le plus pertinent de tous ceux que le communicateur pouvait utiliser pour communiquer I. (p. 237)

Alors, ce que remarque Dominicy, c'est que Sperber et Wilson

²¹ Ils soutiennent que leur principe de pertinence n'est pas, comme chez Grice, une maxime, mais une généralisation: les interlocuteurs ne "suivent" pas le principe de pertinence et ils ne

maintiennent, à cet égard, qu’“on peut être optimalement pertinent sans pour autant être ‘aussi informatif que le demandent les objectifs de l’échange en cours’ (première maxime de quantité de Grice)” (p. 243).

Par ailleurs, nous trouvons aux p. 252-253 un développement, illustré d’un exemple, qui tendrait à prouver que “la deuxième partie de la présomption de pertinence [optimale] (*sic*) entraîne la conséquence suivante: s’il y a plusieurs interprétations du stimulus qui confirment la première partie de la présomption de pertinence, le communicateur a dû vouloir communiquer l’interprétation qui viendra la première à l’esprit du destinataire”. L’exemple est celui de la phrase *Georges a un grand chat*, où le destinataire D n’envisagera que l’interprétation “Georges a un chat domestique”, parce qu’il présumera que si le communicateur C avait voulu l’informer de ce que Georges possède un “animal de l’espèce *Felis*” tel qu’un guépard ou un léopard, il aurait choisi un stimulus suffisamment pertinent pour épargner à D l’effort de construire deux interprétations successives (“chat domestique”, puis “félin”). On voit immédiatement que cette illustration nous place sur le terrain de la première maxime de quantité: en effet, *Georges a un guépard* implique (non-trivialement) *Georges a un chat* (au sens dit “minimal”), mais non réciproquement; la première phrase est donc plus informative, techniquement parlant, que la seconde. Il en irait de même si nous comparions des hypothèses des formes “P et Q” et “P ou Q” dans un contexte ne renfermant que “non-P”: l’hypothèse “P et Q” aurait pour effet contextuel l’effacement de “non-P”, et nous aboutirions à un contexte renfermant “P” et “Q”; l’hypothèse “P ou Q” aurait pour effet contextuel d’impliquer contextuellement “Q”, et nous aboutirions à un contexte renfermant “non-P” et “Q”. Dans le cas de “P ou Q”, D n’envisagera pas l’interprétation où P et Q seraient toutes deux vraies. En termes de pertinence, la raison en est, comme dans l’exemple précédent, que C est présumé épargner à D l’effort de construire deux interprétations successives (“ou exclusif”, puis “et”). Mais à chaque fois (pour “chat” et pour “ou”), on observe que toute hypothèse adjointe logiquement et non-trivialement au contexte initial le serait également si C avait choisi une stratégie plus informative du type “guépard” ou “et” (l’inverse n’étant pas vrai). Compte tenu du fait que C n’aurait pas, de la sorte, accru l’effort de traitement chez D, on peut penser qu’une mesure de pertinence qui soit fondée non seulement sur les effets contextuels, mais aussi sur l’extension des contextes, permet d’expliquer pourquoi l’interprétation forte (“chat domestique”, “ou

exclusif”) s’impose dès l’abord. (Dominicy, 1991, p 90)

Après toutes ces considérations, nous croyons que dans son analyse, Dominicy a raison d’affirmer que la seconde clause de la présomption de pertinence optimale est seulement une version plus précise des deux maximes de quantité de Grice. Dans ce cas, Sperber et Wilson n’ont pas pris en considération la place qu’occupe les maximes de Grice et son analyse inférentielle, dans leur cadre théorique. Faudrait-il que Sperber et Wilson révisent mieux les différences qu’il y a entre leur théorie de la pertinence et la conception de Grice (qui, selon eux, sont nombreuses). Par contre, si l’analyse de Dominicy est correcte, c’est-à-dire, si la seconde clause de la définition de présomption de pertinence optimale n’est qu’une version plus rigoureuse des maximes de quantité de Grice, alors que dire à propos de nos considérations dans la section 2 ci-dessus, où nous avons posé la question de savoir si la définition d’évaluation de pertinence (définition (27), chapitre I) de Sperber et Wilson peut être réduite à quelque chose du type: “une hypothèse plus informative est, en vertu des effets contextuels qu’elle entraîne, plus pertinente dans un contexte?”

On pourrait même penser qu’à ce stade, il se peut que la définition d’évaluation de pertinence telle quelle a été proposée par Sperber et Wilson

nous amène à confondre pertinence avec information. Si tel est le cas, alors des hypothèses plus pertinentes seraient, selon leur approche, des hypothèses véhiculant plus d'informations. Voyons: il nous semble que la seule façon d'obtenir des effets contextuels dans un contexte est par le biais d'informations traitables dans ce contexte. Nous croyons qu'il existe une relation optimale entre informations et effets contextuels qui se révèle de la façon suivante: pour chaque information (ou hypothèse) traitable dans un contexte donné, il y correspond (au moins) un effet contextuel. Par conséquent, plus grand est le nombre d'informations véhiculées par une hypothèse dans un contexte, plus elle entraînera un nombre considérable d'effets contextuels dans ce contexte.

De fait, nombreux sont les exemples qui nous permettent de vérifier que chaque fois qu'on ajoute plus d'information à une hypothèse (bien entendu information pertinente dans le contexte avec lequel cette hypothèse interagit), plus elle y entraînera d'effets contextuels. Et cela vaut aussi non seulement pour des hypothèses dont l'énoncé est une expression littérale de la pensée, mais aussi pour des hypothèses dont l'énoncé est une expression non littérale de la pensée. Prenons par exemple le contexte constitué des hypothèses exprimées littéralement dans l'exemple (1a-c) de Sperber et Wilson où ont été traitées les hypothèses (2) et (3). Comme nous pouvons le voir, l'hypothèse (2) a, par rapport au contexte (1a-c), une seule information

nouvelle à savoir “Suzanne va épouser Pierre”. Dans ce contexte elle entraîne donc, l’implication contextuelle (4). Quant à l’hypothèse (3), elle présente, par rapport au même contexte (1a-c), deux informations nouvelles soit “Pierre est atteint de thalassémie” et “Pierre va épouser Suzanne”. Pour cette raison, l’hypothèse (3) permet de déclencher en plus de l’implication contextuelle (4), l’implication contextuelle (5). C’est pourquoi Sperber et Wilson suggèrent que l’hypothèse (3) est plus pertinente que (2), elle entraîne dans le contexte (1a-c) une implication contextuelle de plus, soit (5). Contrairement à ce que pensent Sperber et Wilson, nous disons que dans ce cas (3) n’est que plus informative que (2) et qu’il ne s’agit pas de traiter les degrés de pertinence par la quantité d’informations contenue dans une hypothèse. Tout va dépendre de l’intention du communicateur de communiquer seulement les hypothèses qu’il a l’intention de communiquer. Par exemple, il se peut que dans l’exemple (1a-c), le communicateur avait l’intention de vouloir produire chez le destinataire seulement l’implication (4):

(4) Suzanne et Pierre devraient consulter un médecin sur les maladies héréditaires qu’ils risquent de transmettre à leurs enfants.

Dans ce cas, suffirait l’hypothèse (2) pour entraîner une telle implication dans le contexte (1a-c). Ainsi, si le communicateur n’avait aucune intention de

produire (5) chez le destinataire:

(5) Suzanne et Pierre devraient éviter d'avoir des enfants,

il ne serait pas nécessaire d'ajouter au contexte (1a-c) une nouvelle information, à savoir (3). Il se pourrait que le communicateur ait exactement l'intention de ne pas produire (5) en limitant les effets contextuels chez le destinataire à ce qu'il veut lui communiquer.

En ce qui touche à la dimension non littérale du sens, il suffit de noter que dans la communication non littérale comme par exemple l'indirection et les implicatures, les hypothèses elles-mêmes déjà véhiculent plus d'information dans leurs énoncés que celles de dimension littérale. Prenons notre exemple (7a-b) en haut. Comme nous pouvons le remarquer, la réponse indirecte d'Hélène véhicule plus d'information que si elle avait simplement communiqué littéralement et complètement sa pensée; et cela justifie les implications contextuelles (13) et (14) qui suivent de sa réponse. Sur ce point, Sperber et Wilson eux-mêmes l'avaient déjà remarqué: "Une réponse indirecte à une question véhicule plus d'information qu'une réponse directe; de manière générale, il découle du principe de pertinence que cette information supplémentaire doit posséder une pertinence propre" (Sperber et Wilson, 1989,

p. 293). Dans ce cas, si Hélène n'a pas répondu simplement (10) *qu'elle n'aimerait pas manger du bœuf* c'est parce qu'elle voulait en (7b) (*je ne mange pas de viande rouge*), communiquer plus que (10), sinon elle aurait répondu directement, par exemple *Je ne mange jamais du bœuf*. En vertu du fait que, par (7b), Hélène communique plus qu'une réponse directe (en effet, elle communique, par exemple, (13) et (14)), il nous semble contre-intuitif de dire que (7b) devient plus pertinent que l'énoncé *Je ne mange jamais du bœuf* dans le contexte de cette conversation-là.

À partir de toutes ces considérations, pouvons-nous interpréter chez Sperber et Wilson leur définition de degré de pertinence d'une hypothèse, de la façon suivante: *une hypothèse est d'autant plus pertinente dans un contexte donné qu'elle y est plus informative ?* (Remarquons que la deuxième clause de la définition de présomption de pertinence optimale fait appel à la notion d'évaluation de pertinence - définition (27), chapitre I: en effet, selon cette clause-là “ *le stimulus ostensif est le plus pertinent de tous ceux que le communicateur pouvait utiliser pour communiquer I* ”. C'est en ce sens que la définition de degrés de pertinence d'une hypothèse est ici présupposée). Mais si tel est le cas, nous avons cette fois-ci une certaine divergence entre telle interprétation et les deux maximes de quantité de Grice (1. que votre contribution soit aussi informative que nécessaire et 2. que votre contribution

ne soit pas plus informative que nécessaire). Or, la deuxième maxime de quantité de Grice prescrit que l'hypothèse (ou énonciation) du locuteur ne doit pas être plus informative que nécessaire tandis que dans l'approche de Sperber et Wilson l'évaluation de pertinence semble entraîner la conséquence suivante: pour être plus pertinente dans un contexte, une hypothèse doit être plus informative - afin de produire beaucoup d'effets contextuels - dans ce contexte. Donc, ce n'est qu'une divergence entre la définition d'évaluation de pertinence chez Sperber et Wilson et les maximes conversationnelles (de quantité) de Grice, et nous en concluons qu'il faudra, de façon urgente, débarrasser telle définition de cette divergence-là car si d'un côté notre analyse est correcte, d'un autre côté, nous sommes d'accord avec l'analyse de Dominicy selon laquelle il existe une convergence entre la seconde clause de la définition de pertinence optimale chez Sperber et Wilson (laquelle fait également appel à leur notion d'évaluation de pertinence) et les deux maximes de quantité de Grice. Cela revient à dire que de fait la seconde clause de la définition de présomption de pertinence optimale n'est qu'une version plus précise des maximes de quantité de Grice (dans ce cas, Sperber et Wilson n'ont pas pris en considération la place qu'occupe les maximes de Grice et son analyse inférentielle, dans leur cadre théorique; il faudrait que Sperber et Wilson révisent mieux les différences qu'il y a entre leur théorie de la pertinence et la

conception de Grice lesquelles, selon eux, sont nombreuses).

Alors, du fait que la définition d'évaluation de pertinence chez Sperber et Wilson entraîne, dans la seconde clause de la définition de présomption de pertinence optimale, à la fois une divergence et une convergence, nous proposons notre propre définition de *degrés de pertinence pour un interlocuteur* (définition (1.8), chapitre III) laquelle rend compte de cette lacune-là et dans laquelle on constate une certaine convergence avec les maximes de quantité de Grice.

CHAPITRE V

PRÉSENTATION DE LA THÉORIE DES ACTES DE DISCOURS DE SEARLE ET VANDERVEKEN

1- Introduction

Dans ce chapitre nous allons présenter quelques résultats importants pour notre travail obtenus par Searle et Vanderveken en Théorie des Actes de Discours (T.A.D.). Ces résultats concernent avant tout les actes de discours qui sont importants pour la théorie de la signification. De tels actes sont du type appelé par Austin *actes illocutoires*.

La sémantique formelle du succès et de la satisfaction est une extension conservatrice de la grammaire de Montague, laquelle s'est limitée seulement à l'analyse sémantique des énoncés déclaratifs. Cette théorie a enrichi la grammaire de Montague, en faisant une analyse adéquate de la forme logique des actes de discours et en procédant à une interprétation sémantique

des énoncés non déclaratifs et *performatifs*. Une telle sémantique caractérise tant les aspects véri-conditionnels que les aspects illocutoires de la signification des énoncés.

1- **Considérations historiques**

Le courant logique de la philosophie du langage, fondé par Frege, B. Russell (1905), le premier Wittgenstein (1961), Carnap (1956) et Tarski (1956), a étudié la forme logique et les éléments constitutifs des énoncés déclaratifs qui servent à exprimer nos connaissances, croyances et opinions concernant les états de choses dans le monde. Ultérieurement développé par Church (1951), Montague (1970), David Lewis (1972), Davidson (1984) et Kaplan (1975), le courant logique a permis de discuter de la forme à donner à une théorie adéquate de la vérité en analysant les conditions de vérité des propositions exprimées par les énoncés déclaratifs. Les énoncés déclaratifs expriment des propositions qui ont des conditions de vérité. Dans ce courant, le seul critère d'identité propositionnelle est celui de la stricte équivalence ou d'identité des conditions de vérité.

Mais comment expliquer la signification d'un énoncé en la

restreignant à ses conditions de vérité? Par exemple, comment expliquer la signification de l'énoncé "Je te promets de venir", lequel semble dépourvu de conditions de vérité?

Pour rendre compte de ce genre de problème, le courant de la philosophie du langage ordinaire fondé par G.E.Moore, le second Wittgenstein, Strawson et Austin a étudié principalement l'utilisation et les fonctions du langage dans le discours. Sous l'influence du second Wittgenstein, Austin (1962) amorce une nouvelle analyse de la signification, soit une analyse des actes de discours. Il a considéré que, par leurs énonciations, les locuteurs entendent d'abord et avant tout accomplir des actes de discours du type illocutoire tels que des assertions, excuses, promesses, offres et questions. Austin a découvert de tels actes en prêtant attention aux énoncés performatifs comme "je t'ordonne de partir", "je te baptise Marie", dont les énonciations littérales réussies constituent l'accomplissement d'actes illocutoires. L'idée principale d'Austin est que les actes illocutoires sont les unités premières de signification et de communication dans l'usage et la compréhension du langage en général. Malheureusement son analyse s'est limitée à des actes de discours isolés accomplis lors de l'énonciation d'un seul énoncé. Ce courant, qui a été développé par la suite par des philosophes contemporains comme Searle, Strawson et Grice, traite des conditions de succès et de satisfaction des actes

de langage accomplis par les locuteurs lorsqu'ils utilisent (littéralement ou non) des énoncés de n'importe quel type syntaxique.

Comme nous l'avons déjà dit dans l'introduction de notre thèse, Searle (1969) a élaboré une théorie générale des actes de discours où il fait ressortir l'idée de la force comme partie de la signification. Les actes illocutoires élémentaires complets sont composés d'une force illocutoire F en plus d'un contenu propositionnel P . Ainsi, les énoncés qui sont utilisés pour accomplir de tels actes illocutoires élémentaires sont composés à la fois d'un indicateur de force illocutoire (assertif, directif, exclamatif, force de promesse, etc.) et d'une clause exprimant un contenu propositionnel.

À son tour, Grice (1975) révolutionne la pragmatique en introduisant les implicatures (*implicatures*). Il fait voir que la signification de l'énonciation par le locuteur dans un contexte déterminé ne peut pas toujours être réduite à sa force et à son contenu propositionnel littéraux.

En remarquant que certains énoncés communiquent plus que ce que les mots constituants de la phrase ne signifient, Grice a montré qu'il peut exister une divergence entre ce qui est dit et ce qui est implicite (implicat). Alors que le contenu communiqué par la phrase correspond au contenu logique de l'énoncé, c'est-à-dire à ses conditions de vérité, le sens transmis (implicitement) par l'énonciation de l'énoncé échappe aux conditions de vérité.

Les implications conventionnelles sont déclenchées par des mots ou des expressions linguistiques tandis que les implications conversationnelles sont déclenchées par des principes généraux de nature rationnelle (elles dépendent du contexte, de la situation, d'informations d'arrière-plan dont les interlocuteurs sont mutuellement au courant et du sens conventionnel de l'énoncé). Au sens de Grice, ces implications conversationnelles mettent en cause le principe de coopération et l'utilisation des maximes conversationnelles.

Comme nous allons voir par la suite, les implications conventionnelles et les implications conversationnelles sont définies, selon lui, à partir des critères suivants: l'implication conversationnelle est, au contraire de l'implication conventionnelle, calculable, annulable, non détachable, non conventionnelle, dépendante de l'énonciation et indéterminée (ouverte). Remarquons que de tels critères ne constituent pas des conditions nécessaires ou suffisantes pour s'assurer de la présence d'une implication dans un certain contexte.

L'implication conversationnelle est calculable car elle est déterminée par une procédure (inférence) basée sur le principe de coopération et les maximes conversationnelles; l'implication conventionnelle au contraire, n'est déterminée que de façon automatique à partir du contenu de l'énoncé.

L'implication conversationnelle est annulable soit explicitement (par l'addition d'une proposition qui affirme ou sous-entend que le locuteur se met hors jeu relativement au Principe de Coopération), soit contextuellement (si la forme de l'énoncé qui produit l'implication est utilisée dans un contexte qui rend clair l'intention du locuteur de se mettre hors jeu). L'annulation d'une implication conversationnelle ne produit pas une contradiction, comme nous pouvons le voir dans l'exemple suivant:

(1) (a) Elle a eu des enfants et elle s'est marié.

Ainsi, (b) est une implication conversationnelle de (a):

(b) D'abord, elle a eu des enfants; ensuite elle s'est marié.

Mais, cette implication peut être annulée par l'addition de (c):

(c) Elle a eu des enfants et elle s'est marié, mais pas forcément dans cet ordre.

Comme on peut le remarquer, le "et" vérifonctionnel qui apparaît dans (a) n'a aucune connotation temporelle et il est commutatif comme le montre (c).

Selon Grice, l'implication conversationnelle déclenchée par

l'énonciation d'un énoncé, dans un contexte, est non détachable quand elle est déclenchée par la signification de l'énoncé dans ce contexte (et non pas par sa forme), si bien que toute une autre façon de dire la même chose produit l'implication en question. Ainsi quand une implication est non détachable on peut remplacer n'importe quelle expression utilisée par un synonyme. Voyons l'exemple suivant:

(2) Dans un contexte où Jean demande si Pierre est un bon philosophe et que quelqu'un lui répond:

(a) Pierre a une belle calligraphie.

Cela implique conversationnellement (b):

(b) Pierre est un très mauvais philosophe.

Quelconque énoncé (du type c-e , par exemple) qui a le même sens que (a) et qui est utilisé dans le même contexte, produit la même implication (b):

(c) Pierre a une calligraphie remarquable.

(d) La calligraphie de Pierre est très satisfaisante.

(e) Pierre a une calligraphie admirable, etc.

S'assurer qu'il s'agit bien d'une implication conversationnelle

requiert une connaissance du sens conventionnel des expressions linguistiques à partir desquelles les implications sont produites, mais ce qui est implicite conversationnellement ne fait pas partie de la définition de départ du poids conventionnel de l'expression.

L'implication conversationnelle est dépendante de l'énonciation dans le sens où elle n'est pas produite par ce qui est dit, mais par la manière dont on l'exprime. Ce qui est dit peut être vrai alors que ce qui est implicite peut être faux.

Une implication conversationnelle peut être indéterminée (ouverte). Prenons le cas des métaphores. Les implications qui peuvent être inférées d'une métaphore, comme "Pierre est un ange" peuvent être: "Pierre est délicat, affectueux, doué de bonnes qualités, parfait, etc.". Dans ce cas, si un tel inventaire n'est pas exhaustif, il n'extrait pas en totalité ce que veut dire le locuteur.

À la suite de Grice, Searle (1979) étudiera les actes de discours non littéraux tels que les actes de discours indirects, et les métaphores où ce que le locuteur signifie est différent de ce qu'il dit. Différemment des énonciations littérales (où la signification du locuteur est identique à celle de l'énoncé qu'il utilise), les actes de discours indirects, par exemple, sont des énonciations dans lesquelles ce que le locuteur veut dire dépasse ce qu'elles

signifient.

2- La sémantique des actes de discours

Comme nous l'avons mentionné auparavant, les actes de discours importants pour la théorie de la signification, appelés par Austin (1962) *actes illocutoires*, sont les actes accomplis au moyen de l'énonciation d'énoncés. Dans la T.A.D., les actes illocutoires élémentaires complets sont composés d'une force illocutoire F qui sert à indiquer la fonction réalisée lors de l'usage littéral de l'énoncé (assertif, directif, force de promesse, etc.), et d'un contenu propositionnel P auquel la force est appliquée. Ainsi, les actes illocutoires élémentaires complets sont de la forme $F(P)$. Quelques-uns sont de la forme $F(R)$ où R est seulement une référence, par exemple "Vive l'équipe brésilienne de soccer!" qui n'est pas une proposition complète; d'autres ont la forme $F(\emptyset)$, où il y a seulement une force, par exemple l'expression de douleur "Aie!" dépourvue de contenu propositionnel. Les forces illocutoires des énoncés élémentaires sont déterminées par des marqueurs, à savoir, les verbes performatifs, l'ordre des mots, le mode du verbe, les signes de ponctuation, l'intonation et autres traits syntaxiques des énoncés dont la

signification consiste à indiquer quelles forces illocutoires aurait l'énonciation littérale de ces énoncés. "Le sujet et le temps du verbe, par exemple, sont des traits constitutifs des clauses propres aux énoncés 'Est-ce que Paul le fait?', 'Jean est-il en train de le faire?' et 'Est-ce que Julie le fera?' " (Vanderveken, 1988, p. 22). Un énoncé ambigu produit dans un contexte d'emploi possible peut avoir plusieurs forces illocutoires. Selon cette théorie, les actes illocutoires qui ont la forme $F(P)$ sont accomplis dans les langues naturelles par des énoncés élémentaires $f(p)$ où f représente un marqueur de force illocutoire et p représente une clause.

Pour ce qui est des actes illocutoires complexes, comme par exemple, les actes conditionnels de la forme $P \rightarrow F(Q)$, les actes de dénégation illocutoire de la forme $\sim F(P)$ et les actes illocutoires liés par conjonction, de la forme $F_1(P_1) \wedge F_2(P_2)$, on remarque que leur forme logique n'est pas réductible à la forme logique des actes illocutoires élémentaires.

Ainsi, les actes illocutoires conditionnels sont des actes de discours complexes de la forme $P \rightarrow F(Q)$ dont le but est d'accomplir un acte illocutoire $F(Q)$ non pas catégoriquement, mais à la condition qu'une proposition P soit vraie. (Vanderveken, 1988, p. 21)

Dans l'approche de Searle et Vanderveken (1985), les actes

illocutoires ont des *conditions de succès* et de *satisfaction*. Du fait qu'ils sont pourvus d'intentionnalité, ils sont dirigés vers des états de choses que le locuteur représente avec l'intention d'effectuer une correspondance entre les mots et les choses. Ainsi les *conditions de succès* d'un acte illocutoire sont les conditions nécessaires au locuteur, dans le contexte de son énonciation, pour qui lui réussisse à accomplir cet acte dans ce contexte. Une promesse, par exemple, a pour une condition de succès, que le locuteur s'engage à accomplir une action future dans le monde. Vanderveken définit, en outre, les *conditions de satisfaction* des actes illocutoires. Selon lui, elles sont en général différentes de leurs conditions de succès. Elles sont "les conditions qui doivent être remplies dans le monde d'un contexte d'énonciation pour que cet acte soit satisfait dans ce contexte" (1988, 34). Il se peut que le monde ne corresponde pas au contenu propositionnel de l'acte illocutoire et, dans ce cas, il peut être accompli avec succès, mais dépourvu de satisfaction. Par exemple, en donnant un conseil à l'allocutaire, un locuteur manifeste un désir qu'il fasse quelque chose dans le monde, mais il se peut que l'allocutaire ne suive pas le conseil qui lui a été donné. Dans l'approche de Searle et Vanderveken les actes illocutoires complets - avec des conditions de succès et de satisfaction (et non pas des propositions isolées - avec des conditions de vérité) sont les unités premières de signification et de communication dans l'usage littéral du

langage.

Comme nous l'avons déjà remarqué, les actes illocutoires élémentaires complets sont composés d'une force illocutoire F servant à indiquer la fonction réalisée lors de l'usage littéral de l'énoncé et d'un contenu propositionnel P auquel la force est appliquée. Chaque force est divisée en sept composants:

(i) Le *but illocutoire*

Le but illocutoire est la composante la plus importante des forces illocutoires parce qu'il détermine, en général, les conditions de sincérité, les conditions sur le contenu propositionnel et les conditions préparatoires ainsi que la *direction d'ajustement* des énonciations qui ont cette force. Les énonciations peuvent avoir cinq buts illocutoires: 1) le but illocutoire *assertif*, qui est de représenter comment les choses sont dans le monde; 2) le but illocutoire *engageant* (commissif), qui consiste à engager le locuteur à accomplir une action future dans le monde; 3) le but illocutoire *directif*, qui consiste à essayer de faire en sorte, de façon linguistique, que l'allocutaire accomplisse une action future dans le monde; le but illocutoire *déclaratif*: rendre existant un nouvel état de choses représenté par le contenu

propositionnel d'une énonciation du seul fait de cette énonciation, et (5) le but illocutoire *expressif*, qui consiste à exprimer des états psychologiques du locuteur à propos d'états de choses dans le monde. Ces cinq buts illocutoires correspondent à quatre directions possibles d'ajustement entre le langage et le monde: la direction d'ajustement des mots aux choses (les actes illocutoires avec le but assertif ont cette direction d'ajustement; en faisant une assertion, prédiction, témoignage ou conjecture, le locuteur entend représenter comment les choses sont dans le monde); la direction d'ajustement des choses aux mots (c'est la direction d'ajustement qu'ont les actes illocutoires avec le but engageant et le but directif; en faisant une promesse, acceptation, recommandation, demande ou question, le monde sera transformé par une action future du locuteur ou de l'interlocuteur, afin qu'il corresponde au contenu propositionnel); la double direction d'ajustement (les actes illocutoires avec le but déclaratif ont cette direction d'ajustement; en cas de satisfaction d'une déclaration telle qu'une bénédiction, une ratification, un licenciement, le locuteur accomplit l'action qu'il dit accomplir et, dans ce cas, le monde est transformé par l'action du locuteur au moment de son énonciation); et la direction vide d'ajustement (les actes illocutoires dont le but est expressif ont la direction vide d'ajustement; il n'est pas question de succès ou d'échec d'ajustement: en remerciant, le locuteur exprime de la gratitude, en se vantant,

le locuteur exprime son orgueil).

(ii) *Le degré de puissance du but illocutoire*

Un but illocutoire peut être réussi avec plus ou moins de force. Par exemple, lors d'une supplication, le locuteur fait une forte tentative linguistique pour pousser l'allocutaire à faire quelque chose.

(iii) *Le mode d'atteinte*

C'est la façon et le moyen par lequel le but est atteint. Par exemple, pour donner un ordre à quelqu'un, il faut invoquer, en donnant l'ordre, une position d'autorité qui ne laisse aucune option de refus à l'allocutaire. En accomplissant un acte illocutoire comme une demande, le locuteur doit donner une option de refus à l'allocutaire.

(iv) *Les conditions sur le contenu propositionnel*

Plusieurs forces illocutoires imposent des conditions sur les contenus propositionnels d'actes de discours avec ces forces. C'est le cas, par exemple, du contenu propositionnel d'une prédiction, lequel doit représenter,

par rapport au moment de l'énonciation, un état de choses futur; la prédiction ne porte ni sur le passé ni sur le présent.

(v) *Les conditions préparatoires*

Le locuteur présuppose que certaines propositions sont vraies dans chaque contexte où il accomplit un acte de discours de la forme $F(P)$. Ces présuppositions du locuteur sont fonctions de la force F , du contenu propositionnel et du contexte. Par exemple, quand un locuteur donne un ordre à quelqu'un il présuppose que l'allocataire est capable de lui obéir. Ainsi les conditions préparatoires d'une force F déterminent quelles propositions un locuteur doit présupposer dans un contexte d'énonciation pour qu'il accomplisse des actes illocutoires ayant cette force.

(vi) *Les conditions de sincérité*

Lors d'une énonciation, le locuteur exprime des états mentaux relativement au contenu propositionnel de cette énonciation. Par exemple, lors d'un remerciement, le locuteur exprime de la gratitude; lors d'une assertion, une croyance, etc. Une condition de sincérité n'est qu'un ensemble de modes

psychologiques d'attitudes. Quiconque s'engage exprime une intention.

(vii) *Le degré de puissance des conditions de sincérité*

Lors de l'accomplissement d'un acte de discours, le locuteur exprime un état psychologique avec plus ou moins de force. Par exemple, un témoignage augmente le degré de puissance d'une assertion: le locuteur qui témoigne exprime une croyance plus forte.

Ces sept composantes de la force illocutoire que nous venons de mentionner déterminent les conditions qui doivent être remplies, dans un contexte d'énonciation, pour qu'un acte illocutoire soit accompli avec succès et soit sans défaut dans ce contexte. Vanderveken (1988) réduit à six les composantes de la force illocutoire car selon lui, le degré de puissance de but illocutoire est déterminé par le mode. Ainsi, chaque force illocutoire consiste en un but illocutoire, un mode d'accomplissement de ce but, des conditions sur le contenu propositionnel, des conditions préparatoires, des conditions de sincérité et un degré de puissance des conditions de sincérité.

Comme Searle et Vanderveken (1985) l'avaient déjà signalé, les forces illocutoires ont souvent des degrés de puissance identiques de but illocutoires et de conditions de sincérité. Par exemple, un locuteur qui supplie

exprime un désir relativement fort, mais cela correspond aussi à une forte tentative linguistique pour déterminer l'allocataire à réaliser quelque chose. Néanmoins, certaines forces illocutoires comme, par exemple, la force de commandement, augmentent le degré de puissance de leur but illocutoire, mais non pas nécessairement le degré de puissance de leur conditions de sincérité: il est possible pour un locuteur de donner un commandement mais sans désirer fortement d'être obéi.

Soulignons que dans une théorie sémantique des actes de discours, la signification du locuteur dans un contexte est identifiée à celle de l'énoncé qu'il utilise dans ce contexte. De ce point de vue, un *contexte possible d'énonciation* est formé par les éléments suivants: le(s) locuteur(s), (les) allocataire(s), le temps et le lieu de l'énonciation, le monde de l'énonciation et les énoncés que le locuteur a utilisés dans ce contexte. Ainsi, l'interprétation d'un énoncé dépend du contexte de son énonciation. La sémantique est conçue comme la théorie de la signification de l'énoncé. Elle rend compte seulement de la compréhension des actes de discours littéraux. Mais, comme nous l'avons vu précédemment, souvent les interlocuteurs font des usages non littéraux des énoncés et la caractéristique principale des actes de discours non littéraux, c'est que la signification du locuteur diffère de celle de l'énoncé qu'il utilise lors de son énonciation. Ce dont nous avons besoin c'est d'une

pragmatique pour analyser l'accomplissement et la compréhension des actes de discours non littéraux.

3- La pragmatique des actes de discours

La pragmatique, conçue comme une étude de la signification non littérale est aussi analysée dans la T.A.D. Les auteurs Searle et Vanderveken analysent les figures de signification telles que l'ironie (le locuteur accomplit un acte de discours littéral, mais son acte de discours principal, à savoir, celui qu'il a l'intention d'accomplir, est le contraire de ce qu'il dit), les métaphores (le locuteur accomplit un acte littéral avec l'intention de signifier autre chose que ce qu'il dit), et de même les actes de discours indirects (le locuteur entend accomplir, en plus de l'acte illocutoire littéral, un acte illocutoire principal par l'accomplissement du premier), et les *implicitations* conversationnelles (le locuteur accomplit un acte de discours littéral, mais il a l'intention de signifier davantage que ce qu'il dit). Pour Vanderveken, la pragmatique est l'étude de la signification du locuteur surtout quand celle-ci ne coïncide pas avec la signification de l'énoncé qu'il utilise. En d'autres mots, la pragmatique est la théorie du discours non littéral.

En termes généraux, la pragmatique traite des liaisons existantes entre les expressions linguistiques, leur signification et l'usage qu'on peut en faire en parlant; elle est conçue comme la théorie de la signification du locuteur. Sa fonction consiste à interpréter les énonciations non littérales et à expliquer la capacité qu'ont les interlocuteurs d'accomplir et de comprendre les actes de discours non littéraux.

À notre avis, la syntaxe, la sémantique et la pragmatique sont trois niveaux dans l'étude du langage. De façon générale, la pragmatique d'une langue présuppose la sémantique de cette langue (règles de bonne formation, une connaissance de la signification linguistique, ainsi que l'habileté à identifier les conditions de vérité des propositions, etc.). Sur ce point, la T.A.D. soutient que dans la mesure où la signification est dans l'usage (et cela renvoie à la sémantique); en bref, la pragmatique doit inclure la sémantique tout comme la sémantique doit inclure la syntaxe. Dans le programme de la T.A.D. il n'y a pas de distinction entre compétence/performance. Cela fait partie de la compétence que de savoir comment utiliser la langue en situation. Une théorie complète de la signification de l'énoncé détermine les conditions de vérité des propositions et établit une relation de conséquence logique entre de telles propositions. Pour cette raison, la pragmatique qui traite des divers usages des énoncés (qui ont déjà une signification) présuppose la sémantique pour

développer une théorie plus générale de la signification du locuteur qui soit capable d'interpréter de tels usages.

Selon nos conclusions précédentes, tant Grice (1975) que Searle (1979) ont montré que dans la conversation on ne comprend l'acte de discours principal non littéral d'une énonciation qu'en identifiant, premièrement, l'acte de discours littéral (c'est là que la pragmatique a besoin de la sémantique) et en comprenant que cet acte littéral ne peut pas être l'acte principal dans le contexte de l'énonciation si le locuteur respecte des principes ou règles pragmatiques.

4- L'analyse de la conversation

L'une des contributions les plus importantes dans les recherches pragmatiques contemporaines sur le problème de la communication est sans doute la contribution de Grice (1975). Comme nous l'avons déjà mentionné dans le chapitre I (section 1), son idée est que l'usage du langage est régi par le principe général de Coopération, à partir duquel quatre types de maximes conversationnelles sont dérivées: les maximes de quantité, de qualité, de relation et de manière.

Jusqu'à présent aucune théorie ne rend compte adéquatement de la signification non littérale. Mais, selon Vanderveken (1990), il est possible de formuler, en utilisant la logique illocutoire, des principes généraux pour traiter adéquatement le phénomène de la signification non littérale. Pour ce faire, il faut caractériser la nature des inférences qui sont faites lors de la compréhension des énonciations non littérales. Comme nous venons de le remarquer, dans *Expression and Meaning* (1979), Searle présente un modèle des actes de langage indirects, de l'ironie et des métaphores. Ce modèle comprend:

- (i) une théorie de la signification linguistique;
- (ii) une théorie des actes de discours (laquelle permet à l'auditeur d'identifier les conditions de félicité des actes illocutoires exprimés ou tentés);
- (iii) certains principes généraux de conversation coopérative (quelques-uns décrits par Grice), lesquels permettent à l'auditeur d'identifier, ce que le locuteur veut dire non littéralement;
- (iv) une analyse de certains faits de l'arrière-plan conversationnel de l'énonciation dont les interlocuteurs sont mutuellement au courant et
- (v) la capacité de l'auditeur à faire des inférences sur la base de l'hypothèse que le locuteur respecte les maximes conversationnelles et que les faits de l'arrière-plan en question existent.

L'influence de l'arrière-plan sur la détermination de la signification littérale ou non littérale est soulignée par Searle dans sa philosophie du langage. On considère l'arrière-plan fondamental en théorie des actes de discours en général relativement à la détermination de la signification littérale, et également pour la pragmatique des actes de discours. Selon Searle (1983), l'arrière-plan est constitué de régularités naturelles (en particulier biologiques), communes à tous les êtres humains normaux comme, par exemple, les capacités de percevoir, de marcher, etc., et de régularités sociales (culturelles), comprenant des choses comme ouvrir les portes, etc. Autrement dit, l'arrière-plan est constitué de capacités, de pratiques, d'hypothèses, présuppositions, enfin, de formes primitives du *savoir-faire* (savoir comment faire quelque chose) et du *savoir-que* (savoir comment sont les choses dans le monde - l'information acquise) qui sont ou des capacités fondamentales communes à tous les êtres humains (ce que Searle appelle l'*arrière-plan profond*), ou des capacités acquises socialement (que Searle appelle l'*arrière-plan local*). Ces capacités de l'arrière-plan permettent aux individus de comprendre la signification des locuteurs, qu'elle soit littérale ou non littérale.

Selon Searle et Vanderveken, les allocutaires recourent à l'arrière-plan pour interpréter les énonciations non littérales et trouver l'acte principal ou les *implications* conversationnelles que le locuteur entend exprimer non

littéralement.

En ce qui concerne les principes de la logique conversationnelle formulés par Grice, il est possible pour l'auditeur de déterminer quels actes de discours non littéraux le locuteur entend accomplir en faisant de telles inférences.

Rappelons les maximes conversationnelles de Grice (1975, 1979) lesquelles ont été présentées dans le chapitre I. Les maximes de quantité: "Que votre contribution contienne autant d'information qu'il est requis; que votre contribution ne contienne pas plus d'information qu'il n'est requis". Les maximes de qualité: "N'affirmez pas ce que vous croyez être faux; n'affirmez pas ce pour quoi vous manquez de preuves". La maxime de relation: "Soyez pertinent". Les maximes de manière: "Soyez clair; soyez bref; soyez ordonné".

Malheureusement, ces maximes de Grice tiennent compte seulement des assertions; elles ne formulent pas des principes généraux pour l'analyse de la signification non littérale. Selon Grice, de telles maximes sont dérivées de principes de rationalité. Asa Kasher (1976), qui s'inspire de Grice, a formulé des maximes de rationalité (instructions) pour les substituer aux principes et maximes de Grice. À toutes les étapes en cours de la réalisation de votre fin, agissez comme requis pour l'atteinte de votre fin, *ceteris paribus*. Selon lui, le principe gricéen impose aux locuteurs d'une conversation,

l'exigence de partager un même but et aussi de savoir que ce but est mutuel. Dans les maximes de Kasher, il n'y a pas de notions sémantiques, elles ne font appel à aucune compétence linguistique particulière; ce sont des maximes générales d'action. Les autres instructions générales sont formulées en respectant le principe général.

En utilisant la logique illocutoire, Vanderveken (1990) propose la réformulation des maximes de quantité et de qualité de Grice en les étendant à tous les types d'actes de discours. Pour Vanderveken, ces deux maximes sont les principales maximes que les locuteurs respectent dans le discours. Sa réformulation du principe de coopération se fonde sur le principe de rationalité selon lequel les locuteurs sont rationnels (minimalement cohérents) dans leur utilisation du langage. Comme tout acte illocutoire est un moyen pour atteindre certaines fins linguistiques, les locuteurs doivent accomplir dans chaque contexte les actes illocutoires qui conviennent pour parvenir à leurs fins linguistiques. Par exemple, dans une conversation, les locuteurs ne doivent pas accomplir des actes illocutoires incompatibles: ordonner à quelqu'un de faire quelque chose et, en même temps, lui interdire de le faire.

Vanderveken souligne que, d'un point de vue logique, pour qu'une énonciation soit de qualité parfaite, l'acte illocutoire qu'elle sert à accomplir doit être en même temps réussi, sans défaut et satisfait dans ce contexte. Par

exemple, un commandement a pour condition de satisfaction, l'obéissance du commandé. Pour qu'une énonciation soit de quantité satisfaisante dans un contexte d'énonciation, l'acte illocutoire qu'elle sert à accomplir doit être aussi fort qu'il le faut pour atteindre le but linguistique du locuteur dans ce contexte. Par exemple, si le but du locuteur est de supplier l'aide à son allocataire, alors une demande d'aide n'est pas assez forte pour atteindre son but.

Voyons maintenant les généralisations des maximes de quantité et de qualité faites par Vanderveken (1990) dans le but de traiter non seulement les énonciations assertives mais tous les types d'énonciations.

Pour la maxime de qualité: Faites en sorte que l'acte illocutoire principal que vous tentez d'accomplir soit réussi, sans défaut et satisfait dans le contexte de votre énonciation. (Un acte de discours manifestement raté, défectueux ou insatisfaisable n'est pas un bon moyen linguistique d'arriver à ses fins).

Pour la maxime de quantité: Faites en sorte que votre acte de discours principal soit aussi fort qu'il le faut (ni trop fort ni trop faible) pour atteindre vos buts linguistiques courants dans le contexte de votre énonciation.

En cas de signification littérale, il n'y a pas de faits dans l'arrière-plan manifestement incompatibles avec l'hypothèse que le locuteur respecte de telles maximes en accomplissant l'acte littéral. En cas de signification non

littérale, par contre, certains faits de l'arrière-plan conversationnel empêchent le locuteur de respecter les maximes si l'acte illocutoire littéral (exprimé par l'énoncé qu'il utilise) est l'acte illocutoire qu'il entend principalement accomplir. Vanderveken distingue, entre autres, les cas d'*exploitation* et d'*usage* d'une maxime. Selon lui, un locuteur exploite une maxime conversationnelle dans un contexte d'énonciation si et seulement si:

(...) certains faits de l'arrière-fond conversationnel que le locuteur présume être mutuellement connus par lui et l'allocutaire sont tels qu'il a l'intention que l'allocutaire reconnaisse qu'il est incapable de respecter cette maxime conversationnelle dans le contexte de l'énonciation si son acte de discours principal est l'acte illocutoire littéral et si ces faits existent; (...) le locuteur est capable de respecter cette maxime sans violer une autre maxime (à cause d'un conflit) et il veut poursuivre la conversation; (...) le locuteur entend également que l'allocutaire sache qu'ils sont tous deux au courant de ceci. (Vanderveken, 1992, p. 41)

Selon Vanderveken, un instituteur qui dit impérativement à son élève "Quitte la classe immédiatement, s'il te plaît!" exploite la maxime de qualité dans le contexte de son énonciation car il ne réalise pas le but directif avec le mode littéral d'accomplissement d'une demande. Puisque, dans ce contexte-là, il ne donne aucune option de refus à l'élève, celui-ci comprend que l'instituteur entend accomplir principalement un acte illocutoire directif par le mode opposé d'accomplissement. D'autre part, un instituteur manifestement fort impressionné par le travail de son élève exploite la maxime de quantité

pour faire une litote quand il lui dit “Ton travail n’est pas mal”. Dans ce contexte, l’élève comprend que l’instituteur veut indirectement faire une assertion plus forte que l’assertion littérale (à savoir, que son travail est très bon).

D’un autre côté,

un locuteur utilise une maxime conversationnelle dans un contexte d’énonciation si et seulement si certains faits de l’arrière-fond conversationnel qu’il présume être mutuellement connus par lui et par l’audience sont tels (...) qu’il entend que l’audience reconnaisse que, dans ces conditions, il ne respecte cette maxime en accomplissant son acte de discours principal que si un autre acte illocutoire non littéral secondaire est accompli, sans défaut et satisfait dans le contexte de son énonciation et (...) s’il entend également que chaque allocataire sache qu’il veut qu’il comprenne cela. (Vanderveken, 1992, p. 42)

D’après Vanderveken, en répondant à la question “Est-ce que Paul a voté communiste?” en affirmant qu’il est membre du P.C., le locuteur utilise la maxime de qualité quand il veut répondre positivement à la question (en invoquant le fait de l’arrière-plan conversationnel il est manifeste que les membres d’un parti votent pour celui-ci). D’un autre côté, en répondant à la question “Dans quel pays Paul se trouve-t-il?” par l’énonciation “Il est au Japon ou en Chine”, le locuteur utilise la sous-maxime de quantité “Soyez aussi informatif que possible!” quand il veut que l’allocataire comprenne qu’il n’est pas en mesure de faire aucune des deux assertions plus fortes “Il est au Japon”

ou “Il est en Chine”.

Ainsi, tous les actes de discours indirects sont des cas d’exploitation de la maxime de quantité; tous les actes de discours ironiques ou métaphoriques sont des cas d’exploitation de la maxime de qualité. D’autre part, toutes les implicatures conversationnelles sont des cas d’utilisation des maximes.

CHAPITRE VI

CONSIDÉRATIONS CRITIQUES ET APPLICATIONS DE LA LOGIQUE DE LA PERTINENCE *P* À LA THÉORIE DES ACTES DE DISCOURS DE SEARLE ET VANDERVEKEN

Introduction

Comme nous l'avons montré dans le chapitre antérieur, la T.A.D. de Searle et Vanderveken intègre des idées de Austin et de Grice. Ce programme comprend une logique, et cette logique est appelée par ses auteurs la *logique illocutoire*.

La logique illocutoire est une partie du programme de la T.A.D. qui permet de construire la classe de toutes les forces illocutoires possibles et de montrer comment justifier les inférences valides que l'on peut faire avec des énoncés de n'importe quel type (déclaratif ou non déclaratif) exprimant des

actes de n'importe quelle force illocutoire.

La T.A.D. à l'origine est atomiste (non dialogique)²²; elle n'analysait que des actes de discours isolés et non pas des séquences d'actes de discours dans la poursuite d'une conversation. Dans leur classification des actes de discours, Searle et Vanderveken signalent des actes de discours dont les conditions préparatoires renvoient à d'autres actes de discours, mais leur analyse reste à la base atomiste; elle ne tient compte que des actes de discours en tant qu'unités de la conversation. Donc, leur sémantique à l'origine rend compte seulement des actes de discours isolés accomplis dans un seul contexte d'énonciation.

Toute la théorie tient compte premièrement de la littéralité; la sémantique réduit la compétence linguistique à l'habileté des humains à accomplir et à comprendre des actes de discours littéraux: dans un contexte d'énonciation un locuteur signifie seulement ce que l'énoncé qu'il utilise signifie dans ce contexte. Par contre, selon ces auteurs, la T.A.D. permet de formuler des principes généraux pour expliquer en pragmatique l'habileté des humains à accomplir et à comprendre des actes de discours non littéraux ainsi que des implications conversationnelles.

²² Comme nous allons le voir plus loin, Vanderveken propose actuellement un nouveau programme pour analyser la structure des discours dans la conversation.

Voici quelques considérations critiques à la T.A.D. en ce qui concerne plus particulièrement la pragmatique conversationnelle.

1- **Considérations critiques à la pragmatique des actes de discours à la lumière de la logique de la pertinence *P***

La contribution de la T.A.D. pour un traitement formel de la pragmatique consiste à reconstruire formellement les procédures que le locuteur utilise pour donner des indications à l'interlocuteur, soit pour que l'interlocuteur puisse interpréter de telles indications, et faire des inférences afin de calculer la signification du locuteur. Cette théorie nous propose des modèles de procédures pour les métaphores, les actes de discours ironiques, les actes de discours indirects et les implications conversationnelles.

Comme Searle et Vanderveken le reconnaissent, la nécessité d'étudier les actes de discours en sémantique est justifiée par des raisons théoriques liées à la nature même du langage et de l'intentionnalité. La T.A.D. est la seule théorie qui permette de donner à tous les types d'énoncés (déclaratifs ou non) une signification. Elle est nécessaire pour rendre la sémantique formelle capable d'analyser les marqueurs de force et les verbes

performatifs. Nous sommes capables de faire, en parlant, non pas seulement des inférences théoriques valides mais aussi des inférences pratiques valides (ayant comme conclusion des énoncés non déclaratifs). Grâce à la logique des actes de discours, la sémantique formelle peut analyser de tels raisonnements pratiques. Comme nous l'avons souligné plus haut, la sémantique et la pragmatique ne peuvent pas être dissociées; au contraire, la pragmatique présuppose la sémantique. Pour ce qui est de la pragmatique, nous croyons que la T.A.D. est la seule qu'on puisse utiliser comme point de départ pour construire une théorie générale de la signification du locuteur.

Toutefois, comment comprenons-nous les actes de discours non littéraux ainsi que les implications conversationnelles? On a déjà vu que d'après Grice et Searle, nous raisonnons en utilisant des maximes conversationnelles (comme: dire le vrai et être sincère) et notre connaissance partagée de l'arrière-plan conversationnel. Selon la T.A.D., l'arrière-plan conversationnel est nécessaire pour analyser les actes de discours littéraux aussi bien que les non littéraux et les implications conversationnelles, en considérant des faits qui peuvent être "pertinents" pour comprendre la signification du locuteur. Mais quel types de faits sont *pertinents* pour comprendre la signification du locuteur? Telle quelle, la T.A.D. nous offre, à partir de la logique illocutoire, une reformulation du principe de coopération de

Grice en l'étendant à tous les types d'actes de discours; la reformulation ne retient que deux des quatre catégories de maximes proposées par Grice: la maxime de quantité et la maxime de qualité; elle ne reformule pas la maxime de pertinence de Grice laquelle selon nous est fondamentale dans le modèle qui explique comment font les locuteurs pour communiquer une signification autre ou plus riche que celle de l'énoncé qu'ils utilisent. Il nous semble que la maxime de pertinence est la plus importante des maximes de Grice, et que toutes les autres dépendent de cette maxime pour leur application; entre le principe de coopération et la maxime de pertinence il y a une grande ressemblance et l'application des maximes dépend de la coopération des interlocuteurs. Une conversation ne se réalise que si elle est coopérative. Nous croyons qu'il faudrait d'abord expliquer la maxime de pertinence et ensuite la reformuler en l'étendant à tous les actes de discours. Vu que la pertinence sert souvent à déclencher une recherche qui, présupposant la coopération entre les interlocuteurs, nous conduit à chercher une interprétation qui diffère du sens littéral de l'énoncé, nous pensons qu'il faudrait de façon urgente expliquer quel est le rôle que joue la pertinence aussi bien dans le discours littéral que non littéral.

De même que les maximes de Grice ne sont pas des maximes générales de la conversation (elles ne sont que des maximes appliquées à des

actes de discours isolés), la T.A.D. n'est pas à l'origine une théorie de la conversation²³. Elle était limitée à analyse des actes de discours isolés, accomplis dans un contexte d'énonciation. Dans ce sens, nous comprenons que le problème de la pertinence n'était pas spécifiquement une affaire de la T.A.D. vu que cette théorie ne tenait pas compte de séquences ordonnées d'énonciations résultant d'une interaction langagière entre un ou plusieurs locuteurs engagés dans une conversation (monologue ou dialogue). Le problème de la pertinence est donc une affaire relevant d'une théorie du discours ou de la conversation. Ainsi, il nous semble que cela suffit pour justifier que la T.A.D., telle qu'elle a été développée, n'a pas senti le besoin de tenir compte du problème de la pertinence des énoncés et des actes illocutoires relativement aux buts des interlocuteurs (linguistiques ou non) dans une conversation.

Toutefois, lorsque la T.A.D. nous propose les reformulations des maximes de qualité et de quantité de Grice (il s'agit là d'une généralisation mais c'est une généralisation seulement au niveau des forces illocutoires), nous ne comprenons pas la raison pour laquelle cette théorie évite une reformulation de la maxime de pertinence de Grice. Nous soutenons dès lors que la

²³ Comme nous le verrons bientôt, Vanderveken développe un nouveau programme pour rendre compte de séquences entières d'actes illocutoires dans la poursuite d'une

pertinence devrait jouer un rôle beaucoup plus important dans la T.A.D. et que pour cette raison, il faut également traiter formellement la maxime de pertinence de Grice.

Or, la logique illocutoire rend la pragmatique capable d'expliquer les deux maximes de quantité et de qualité de Grice. Mais en ce qui touche à la maxime de quantité, il nous semble que sa généralisation fait appel à la maxime de pertinence pour les raisons que nous présenterons par la suite.

Dans une conversation particulière, le locuteur a un objectif (non pas l'objectif ou but illocutoire, mais un objectif différent, plus général, qui peut être perlocutoire, par exemple). Lorsque le locuteur est engagé dans une conversation, il délibère à partir d'objectifs très précis afin d'utiliser le meilleur moyen linguistique (ou un moyen satisfaisant) pour atteindre ses objectifs (on ne peut pas caractériser la délibération du locuteur sans mentionner d'une façon ou d'une autre, que le locuteur utilise des moyens plus ou moins effectifs pour atteindre son objectif dans un contexte d'énonciation). Pour évaluer si l'analyse illocutoire respecte la maxime de quantité, le locuteur évalue si le moyen est maximal pour atteindre ses objectifs (et bien entendu, l'objectif principal du locuteur dans la plupart des cas n'est pas le but illocutoire; son objectif est plutôt perlocutoire, c'est-à-dire, réaliser quelque chose comme

conversation.

divertir, insulter, informer, etc.). Pour évaluer si le locuteur respecte la maxime de quantité et si l'acte de discours atteint son objectif (ce qui est important pour évaluer si l'acte de discours est littéral ou non), l'auditeur a besoin de connaître l'objectif du locuteur; pour identifier l'objectif du locuteur, l'auditeur cherche à voir, par exemple, si l'objectif est approprié en considérant la conversation antérieure (s'il y a lieu), ou la situation où il prend place (le contexte en question), c'est-à-dire, si l'acte de discours accompli satisfait les conditions d'appropriété contextuelle. Et, d'après nous, cela nous rapproche beaucoup de la maxime de pertinence (comme nous le verrons dans la suite).

Sur ce point, nous faisons usage des "conditions d'appropriété contextuelle" telles que définies par Verschueren (1979). D'après Verschueren, les *conditions d'appropriété contextuelle* ne sont des conditions que pour le locuteur tandis que pour l'auditeur, elles sont plutôt des éléments de signification qui déterminent la place qu'occupe un énoncé dans un contexte. Pour cette raison, le locuteur peut les utiliser pour transmettre une information indirectement, à l'exemple d'une femme qui en voulant communiquer indirectement à son destinataire qu'elle est mariée, énonce simplement la phrase "Mon mari est avocat"; la présupposition en question est qu'elle a un mari. Ainsi, pour Verschueren, les conditions d'appropriété sont, par exemple:

les présuppositions logiques ainsi que pragmatiques²⁴; la félicité des actes de discours déterminée au moyen de conditions et de règles; les implications conversationnelles, par exemple, quand le locuteur affirme “Je n’ai plus d’essence” et, à son tour, l’auditeur répond “Il y a un garage de l’autre côté de l’église”, l’implication conversationnelle que le garage est ouvert et qu’on y a de l’essence, est indispensable pour l’appropriété de la réponse; le principe que le locuteur respecte les maximes conversationnelles de Grice, etc. Selon Verschueren, la conversation, les écrits, tant littéraires que scientifiques, ou même les articles ou annonces de journaux, sont des *types de communication* correspondant à des chaînes de conditions différentes. Par exemple, les maximes conversationnelles de Grice sont les “conditions d’appropriété associées à un certain type de communication, c’est-à-dire la forme standard de la conversation. On ne peut pas appliquer ces maximes aux autres types de communication sans les modifier plus ou moins” (p. 280). Il faudrait généraliser de telles maximes (comme nous l’avons déjà montré, dans sa pragmatique illocutoire Vanderveken a formulé une généralisation des maximes de qualité et de quantité).

²⁴ Verschueren cite là-dessus la distinction entre la présupposition logique (ou sémantique) et la présupposition pragmatique faite par Keenan (1971): une phrase S présuppose logiquement une autre phrase S’ si la vérité de S’ est une condition nécessaire pour la vérité et la fausseté de S. D’autre part, l’énonciation d’une phrase présuppose pragmatiquement que son contexte (linguistique et extra-linguistique) est approprié (p.277).

Comme nous l'avons dit plus haut, l'objectif conversationnel et sa réalisation dépendent partiellement d'une discussion, d'une conversation et en général, on peut dire que l'objectif du locuteur est de faire une remarque pertinente parce que l'objectif n'a pas beaucoup de sens s'il n'est pas *lié* à des actes de discours antérieurs, ou dans le cas des situations, aux hypothèses formées lors de la traduction des traits saillants du contexte²⁵. Le caractère d'appropriété est une propriété de certains moyens linguistiques utilisés pour atteindre un but communicationnel (il se peut qu'un énoncé soit approprié dans un contexte, mais non pas approprié dans un autre contexte). Sous ce rapport, notre définition de la pertinence, présentée au chapitre III, dans laquelle un contexte²⁶ particulier est fixé, pourrait être considérée comme étant un *explicatum* (au sens de Carnap) pour expliquer une notion intuitive d'efficacité d'un moyen linguistique utilisé pour réaliser une certaine fin conversationnelle; en ce sens, le caractère d'appropriété du moyen exact linguistique choisi se mesure à partir de l'objectif du locuteur. Et cela renvoie à la rationalité, c'est-à-dire à des relations moyens-fins où les fins sont discursives. De même, nous pouvons utiliser notre définition de la pertinence pour expliquer le caractère d'appropriété de la relation entre un énoncé et un contexte, et cela renvoie aussi

²⁵ Rappelons notre postulat (1.1), chapitre III, selon lequel tous les traits saillants du contexte peuvent être traduits par des hypothèses.

à la rationalité ainsi qu'à des relations moyens-fins car il s'agit d'une interaction moyen-fin. Supposons qu'un interlocuteur veuille atteindre telle fin; donc il doit utiliser le bon moyen, vu ses opinions, ses croyances, ses désirs, ses intentions, ses capacités, ses moyens disponibles; et parfois, les moyens ne sont pas linguistiques (supposons que quelqu'un veuille bouger la table et qu'il existe dans ce contexte, des contraintes comme: pour bouger la table, il faut plusieurs personnes, sinon elle va casser). Dans ce cas, c'est la théorie de la décision rationnelle qui intervient ici, quels moyens choisir dans tel contexte où les gens ont tels désirs, telles capacités, telles intentions - vu leurs opinions sur la relation moyen-fin. Comme le disait Descartes,

Et je m'étais ici particulièrement arrêté à faire voir que, s'il y avait de telles machines, qui eussent les organes et la figure d'un singe, ou de quelque autre animal sans raison, nous n'aurions aucun moyen pour reconnaître qu'elles ne seraient pas en tout de même nature que ces animaux; au lieu que, s'il y en avait qui eussent la ressemblance de nos corps, et imitassent autant nos actions que moralement il serait possible, nous aurions toujours deux moyens très certains, pour reconnaître qu'elles ne seraient point pour cela de vrais hommes. Dont le premier est que jamais elles ne pourraient user des paroles, ni d'autres signes en les composant comme nous faisons pour déclarer aux autres nos pensées. Car on peut bien concevoir qu'une machine soit tellement faite qu'elle profère des paroles, et même qu'elle en profère quelques-unes à propos des actions corporelles qui causeront quelque changement en ses organes: comme, si on la touche en quelque endroit, qu'elle demande ce qu'on lui veut dire; si en un autre qu'elle crie qu'on lui fait mal, et choses semblables; *mais non pas qu'elle les arrange diversement, pour répondre au sens de tout ce qui se dira en sa présence, ainsi que les hommes les plus hébétés peuvent faire*²⁷. Et le second est que, bien qu'elles fissent plusieurs

²⁶ Voir notre définition intuitive de contexte, chapitre III.

²⁷ C'est nous qui soulignons. Il s'agit de la marque de la rationalité pour Descartes.

choses aussi bien, ou peut-être mieux qu'aucun de nous, elles manqueraient infailliblement en quelques autres, par lesquelles on découvrirait qu'elles n'agiraient pas par connaissance, mais seulement par la disposition de leurs organes. Car, au lieu que la raison est un instrument universel, qui peut servir en toutes sortes de rencontres, ces organes ont besoin de quelque particulière disposition pour chaque action particulière; d'où vient qu'il est moralement impossible qu'il y en ait assez de divers en une machine pour la faire agir en toutes les occurrences de la vie, de même façon que notre raison nous fait agir. (Descartes, 1966, p. 120-121)

Puisque l'objectif du locuteur et sa réalisation dépendent du contexte, nous soulignons que les actes de discours que le locuteur accomplit pour atteindre une fin communicationnelle dans un contexte fixé, doivent satisfaire les conditions d'appropriété par rapport à ce contexte. Remarquons, qu'un énoncé du locuteur peut être approprié pour un ensemble Δ d'énoncés exprimant des hypothèses de l'interlocuteur dans un contexte, mais non pas pertinent (et vice-versa). Supposons par exemple, que dans un contexte où les gens sont en train de voter pour les élections présidentielles, quelqu'un fait l'énonciation "*les élections présidentielles ont lieu au niveau du sol*"; alors, comme on peut l'observer, il s'agit d'un énoncé approprié (par la condition thématique) dans ce contexte mais qui n'entraîne aucune conséquence pragmatique dans ce contexte. Dans ce cas, par notre définition de pertinence, cet énoncé n'est pas pertinent dans ce contexte-là (sauf s'il était adressé, par exemple, à des martiens qui étaient présentes au moment de son énonciation dans le contexte donné). C'est dans cet ordre d'idées, que notre définition de la

pertinence pourrait être considérée dans l'explication de la notion intuitive d'efficacité d'un moyen linguistique pour atteindre un but conversationnel.

Sous ce rapport, dans son nouveau programme Vanderveken (1997a) définit ainsi les quatre buts discursifs de base que les interlocuteurs peuvent avoir l'intention commune de réaliser quand ils sont engagés dans une discussion: le *but descriptif* sert à décrire ce qui se passe, par exemple, les exposés, les débats théoriques, les reportages, les entrevues, etc. Les discours avec le but descriptif ont la direction d'ajustement des mots aux choses; le *but discursif délibératif* sert à délibérer sur les actions futures des interlocuteurs dans le monde, par exemple, les délibérations, les négociations, les sermons, les instructions, les règlements, etc. Ce but est en même temps engageant et directif; les délibérations peuvent engager des locuteurs ou tenter d'engager des interlocuteurs à accomplir des actions communes dans le monde; elles ont la direction d'ajustement des choses aux mots; le *but discursif interne déclaratif* est le but par lequel les locuteurs entendent transformer le monde par des déclarations communes. Dans ce genre de discours il faut que les locuteurs aient l'autorité de faire ensemble une série d'actions par le seul fait de dire qu'ils le font: par exemple, des assemblées constituantes ont le pouvoir de promulguer une nouvelle constitution et les assemblées législatives ont le pouvoir de légiférer. Ainsi les discours avec le but discursif interne déclaratif

ont la double direction d'ajustement. Et finalement, le *but discursif expressif* sert à exprimer ou manifester les états d'âme et les attitudes des locuteurs, par exemple, les prières, les lamentations publiques, les protestations verbales, etc. Ce but discursif expressif a la direction vide d'ajustement.

Revenons à l'analyse que nous avons commencé plus haut à propos de la généralisation de la maxime de quantité dans la T.A.D. Nous avons dit que pour évaluer si l'acte de discours accompli par le locuteur est de quantité parfaite, l'auditeur a besoin de connaître son objectif (qui peut être, comme nous l'avons déjà remarqué, un objectif plus général comme, par exemple, l'objectif perlocutoire d'"informer"). Mais d'après nous, seulement cela et les intentions du locuteur ne suffisent pas à l'auditeur car la conversation peut contraindre les intentions et les objectifs (à comprendre) du locuteur et dans ce cas, il nous semble que le locuteur a besoin de la maxime de pertinence. Nous croyons donc que la réformulation de la maxime de pertinence est nécessaire: il s'agit même d'une question d'ordre d'application des maximes de Grice, telles que reformulées par Vanderveken.

La réformulation de la maxime de quantité que nous offre Vanderveken est sans doute une généralisation très importante. Il offre une bonne description formelle de cette maxime. Néanmoins, si on veut l'appliquer (au niveau formel) pour rendre compte de cas plus spécifiques de discours

particuliers dans une conversation, nous devons tenir compte du fait que l'auditeur peut identifier l'objectif du locuteur et évaluer si l'énoncé qu'il utilise est pertinent pour lui dans le contexte donné (en rapport à l'objectif). Mais pour traiter de cas spécifiques nous nous trouvons devant une grande difficulté, à savoir, le problème de l'identification de l'objectif linguistique. Et, conséquemment, le problème de l'identification de l'objectif du locuteur affecte, chez l'auditeur, l'évaluation de la pertinence de son énoncé par rapport à l'objectif ou but conversationnel.

Quel est l'objectif de la maxime de quantité ainsi généralisée? L'objectif de la maxime de quantité est seulement d'évaluer si l'acte de discours est suffisamment fort pour atteindre l'objectif. Mais du point de vue de l'auditeur, on a besoin de connaître l'objectif pour pouvoir l'évaluer. Nous avons donc deux choses à considérer: l'objectif et les actes de discours. Quant à ces derniers, un acte illocutoire est de quantité parfaite si il est suffisamment fort pour atteindre l'objectif linguistique. Quant à l'objectif plus général, comme nous le savons, il n'existe aucune procédure effective qui permette à l'auditeur d'identifier le véritable objectif du locuteur. Évidemment il n'est pas toujours possible de déterminer avec précision les objectifs du locuteur en rapport à un contexte particulier, car le locuteur peut bien dissimuler l'objectif à atteindre et accomplir des actes de discours dans l'intention de rendre

manifeste à l'auditeur un autre objectif, différent de celui qu'il a l'intention d'atteindre. Si l'auditeur n'est pas en mesure d'identifier l'objectif du locuteur, alors il n'est pas possible pour lui d'évaluer, dans cette situation spécifique, si le locuteur respecte la maxime de quantité, c'est-à-dire, si l'acte de discours accompli par l'énoncé qu'il utilise est suffisant pour atteindre son objectif, afin de décider si cet acte est littéral ou non, enfin si l'acte de discours est ou non pertinent dans le contexte de son énonciation. Mais, en règle générale, lorsqu'il n'est pas en notre pouvoir de discerner les véritables objectifs du locuteur, nous devons compter sur les objectifs linguistiques lesquels sont destinés à être communiqués. Commençons par dire que indépendamment de l'objectif plus général que le locuteur peut vouloir dissimuler au destinataire par l'accomplissement d'actes de discours dans un contexte fixé, quand le locuteur accomplit un acte de discours dans ce contexte, il doit avoir l'intention de communiquer son intention²⁸ d'accomplir cet acte à l'auditeur et conséquemment de rendre manifeste **au moins un** autre objectif perlocutoire (quoiqu'il s'agisse maintenant d'un objectif "stratégique") au destinataire pour que le destinataire puisse évaluer si le locuteur respecte ou non la maxime de quantité. Lorsqu'un tel objectif du locuteur est rendu manifeste au destinataire,

²⁸ L'intentionnalité commune est constitutive de la relation d'interlocution entre sujets parlants et elle est déjà présente dans les actes de discours eux-mêmes.

alors il devient une hypothèse formée dans l'ensemble des hypothèses sur le contexte et qui est désormais manifeste à l'un et l'autre des interlocuteurs. Mais, si tel objectif n'est pas manifeste à l'auditeur, alors, il n'est pas en mesure d'évaluer si la maxime de pertinence a été respectée, et par conséquent, si l'acte de discours du locuteur est ou non de quantité parfaite. Rappelons notre définition (2.8), chapitre III, de degrés de pertinence pour un interlocuteur:

Soient $\varphi_1, \dots, \varphi_k$ un ensemble d'énoncés représentant les conséquences pragmatiques nouvelles que le locuteur veut produire chez l'interlocuteur dans un contexte où il utilise l'ensemble d'énoncés Γ et où l'ensemble d'énoncés Δ exprime les hypothèses de l'interlocuteur dans ce contexte. Considérons que l'interlocuteur n'infère aucune de ces conséquences pragmatiques nouvelles $\varphi_1, \dots, \varphi_k$ (dans ce cas, l'ensemble de ses hypothèses exprimées par des énoncés Δ dans le contexte C est différent de l'ensemble des hypothèses que le locuteur fait sur le contexte au moment de son énonciation). C'est ainsi que nous définissons le *degré zéro* (ou degré minimum) de pertinence pour un interlocuteur.

Alors, un tel objectif perlocutoire peut très bien être l'objectif de dissimuler l'objectif perlocutoire principal, soit pour confondre le destinataire en faisant des énonciations ambiguës, soit pour atteindre quelque autre objectif; enfin pour n'importe quel objectif perlocutoire que le locuteur a l'intention de réaliser dans un contexte C , il va utiliser le meilleur moyen linguistique pour

l'atteindre; or, comme nous pouvons le vérifier dans le quotidien, dans plusieurs situations, le locuteur réussit (et réussit bien) à dissimuler son objectif principal et atteindre un autre objectif. Alors, c'est sur l'objectif du locuteur de vouloir communiquer son intention d'accomplir un certain acte de discours à l'interlocuteur que celui-ci s'appuie pour évaluer si le locuteur respecte la maxime de pertinence, et conséquemment, la maxime de quantité par rapport à cet objectif-là. C'est pourquoi seulement les objectifs linguistiques sont à prendre en considération.

Évidemment il existe des situations où l'auditeur perçoit l'attitude dissimulatrice de la part du locuteur et peut inférer exactement son vrai but (dans ce cas, son ensemble d'hypothèses faites sur le contexte est plus ample que celui que le locuteur en fait, c'est-à-dire, l'ensemble d'hypothèses contextuelles de l'auditeur a au moins une hypothèse en plus que l'ensemble d'hypothèses contextuelles du locuteur, à savoir celle que le locuteur suppose que l'auditeur ne connaît pas et sur laquelle le locuteur s'appuie pour essayer de le tromper). Mais c'est une autre discussion. Ce que nous voulons discuter maintenant est le besoin qu'on a de préciser la relation de pertinence entre des actes de discours accomplis par des énoncés dans un contexte par rapport aux objectifs des locuteurs lorsqu'ils font des énonciations. Car nous croyons que cela va nous permettre de déterminer si dans la conversation les interlocuteurs

respectent la maxime de quantité, et selon nous, cela n'est possible que si l'on ne renonce pas à la maxime de pertinence.

Pour résumer, l'explication que nous donnons au problème de l'identification du but conversationnel du locuteur lorsqu'il ne veut pas rendre manifeste son objectif perlocutoire principal, est qu'en accomplissant un acte de discours dans un contexte d'énonciation, le locuteur devra rendre manifeste à l'auditeur au moins un autre objectif, à savoir l'objectif linguistique sur lequel l'auditeur puisse s'appuyer pour interpréter son acte de discours. Et quel que soit l'objectif perlocutoire (principal ou secondaire) que le locuteur veut transmettre au destinataire, nous soutenons que le locuteur ne peut pas réaliser un objectif perlocutoire sans réaliser une remarque pertinente (faute de quoi la conversation échouerait). Alors, la réalisation des objectifs perlocutoires dépend de l'objectif du locuteur de les réaliser par des remarques appropriées, pertinentes, par rapport au but communicationnel qu'il veut atteindre chez le destinataire dans le contexte d'une énonciation. Dans ce cas, l'explication de la maxime de quantité doit faire appel à la maxime de pertinence, en d'autres mots, la maxime de quantité est essentiellement liée à la maxime de pertinence: la maxime de quantité est la forme logique de l'acte illocutoire; la forme logique détermine les conditions de félicité, et les conditions de félicité sont liés aux objectifs linguistiques. Intrinsèquement, l'objectif linguistique et la

forme logique de l'acte illocutoire sont liés; en particulier, l'objectif linguistique du locuteur et les conditions de félicité de l'acte illocutoire sont liés (les actes illocutoires avec les mêmes conditions de félicité sont identiques car ils remplissent les mêmes objectifs linguistiques). Supposons par exemple, qu'un locuteur veut inviter indirectement quelqu'un pour sortir le soir en lui disant: "Est-tu libre ce soir?". Dans ce cas, son objectif est d'inviter l'autre à sortir, pas seulement de lui demander s'il est libre ou non.

Tous les objectifs ne sont pas atteints en accomplissant l'acte littéral; dans l'arrière-plan il y a des faits, comme par exemple: le locuteur aimerait bien sortir avec lui, le locuteur voudrait lui demander de sortir avec lui et le locuteur ne l'a pas fait explicitement; et tous ces faits, toutes ces conditions de succès non littérales sont remplies dans l'arrière plan; alors, l'auditeur va comprendre que, en plus de poser la question, le locuteur veut l'inviter à sortir le soir.

Les autres objectifs du locuteur sont dans l'arrière plan. Et alors, ce qu'il fait, il les ajoute à l'acte illocutoire littéral et il obtient l'acte non littéral à dériver.

Nous croyons que la maxime de quantité et la maxime de pertinence sont liées, elles sont liées intrinsèquement parce que la maxime de pertinence dit:

Veillez accomplir un acte qui ait les bonnes conditions de félicité que vous voulez avoir: aucune superflue et toutes celles que vous voulez avoir; c'est-à-dire, veuillez accomplir un acte pertinent pour atteindre tous vos (et seulement vos) objectifs linguistiques.

Supposons que le locuteur utilise un énoncé qui exprime un acte moins fort. Il y a toute une série de conditions non littérales de félicité, d'objectifs qui sont linguistiques, mais qui ne sont pas exprimés, et qu'il faut ajouter à l'acte littéral pour obtenir l'acte indirect. C'est alors un cas d'exploitation de la maxime de quantité.

Pour ce qui est de l'utilisation de la maxime de quantité, supposons qu'un locuteur pose la question suivante: "Où est Jean?" Si au lieu de répondre: "Il est à Paris", "Il est à Bruxelles", l'interlocuteur répond: "Il est à Paris ou à Bruxelles", il se pose la question de savoir pourquoi l'interlocuteur n'a pas répondu exactement dans quelle ville Jean est? S'il veut seulement dire que "Jean est à Paris ou à Bruxelles", il veut alors affirmer la disjonction et non pas un disjoint. Dans ce cas, s'il respecte la maxime de quantité, alors cette assertion littérale est exactement le bon acte (ni trop fort ni trop faible) pour atteindre ses objectifs. Dans ce cas, il utilise la maxime de quantité. Si les gens attendaient que le locuteur fasse un acte plus fort (mais il ne l'a pas fait), il pourrait répondre: "J'ai affirmé la disjonction parce que je sais qu'il est dans

l'une de ces deux villes mais pas exactement laquelle". Ainsi, les deux assertions "Jean est à Paris" et "Jean est à Bruxelles" seraient des assertions trop fortes dans ce contexte, elles seraient défectueuses, elles pourraient ne pas être satisfaites.

Toutes ces considérations ont pour but de dire que la maxime de quantité est liée à des objectifs linguistiques (illocutoires). Quand on parle de conditions de félicité et d'objectifs, c'est la même chose.

Quand deux actes illocutoires servent à atteindre les mêmes objectifs, ils ont la même forme logique et les mêmes conditions de félicité; quand un acte illocutoire a plus de conditions de félicité qu'un autre, il sert à atteindre plus d'objectifs; quand un acte illocutoire a moins de conditions de félicité qu'un autre, il sert à atteindre moins d'objectifs. "Veuillez accomplir un acte qu'il sert à atteindre tous les objectifs et rien qu'eux!", c'est la même chose que "Veuillez faire un acte qui soit ni trop fort ni trop faible pour atteindre les objectifs!".

Donc, la notion même de la maxime de quantité est liée à la notion d'objectif parce que le nombre des conditions de félicité d'un acte de discours, correspond au nombre d'objectifs de cet acte. Quand on a les mêmes conditions de félicité, on a les mêmes objectifs. Dans le cas de la maxime de quantité, on parle seulement des objectifs linguistiques (elle ne considère pas

les objectifs perlocutoires).

Bien entendu le problème que nous venons de soulever au sujet de la généralisation des maximes de quantité et de qualité n'est pas un problème de la T.A.D. elle-même. Il s'agit d'un problème lié à l'extension de cette théorie lorsqu'elle veut rendre compte de la littéralité et de la non-littéralité. Ainsi, les reformulations des maximes de qualité et de quantité vont au-delà de la pragmatique formelle.

Pour ce qui est de la reformulation de la maxime de qualité, on n'y trouve pas ce genre de problème, puisque la maxime de qualité prescrit d'exécuter un acte de discours réussi, sans défaut et satisfait (un acte illocutoire est de *qualité parfaite* si et seulement s'il est accompli avec succès, sans défaut et satisfait).

Ce serait bien si on pouvait enrichir la T.A.D. de façon à expliquer comment se créent les liaisons entre les actes de discours, c'est-à-dire comment créer le lien entre des actes de discours accomplis par des locuteurs lorsqu'ils sont engagés dans une conversation. Autrement, de quelle autre façon la T.A.D. pourrait développer la pragmatique et la formaliser? Vanderveken lui-même manifeste récemment sa préoccupation à ce sujet:

Pourrait-on enrichir la théorie actuelle des actes de discours et la logique illocutoire de façon à analyser rigoureusement la structure logique des

séquences ordonnées d'énonciations résultant d'une interaction langagière entre plusieurs locuteurs? En particulier, pourrait-on formuler une **classification raisonnée des types de discours** qui sont pourvus d'un but discursif propre? Quelle est la **structure logique** et quelles sont les **conditions de succès** qui doivent être remplis pour que des locuteurs réussissent à tenir de tels discours? (Vanderveken, 1997, p 59)

Dès lors, Vanderveken nous offre une analyse philosophique des composantes et des conditions essentielles de succès des discours qui ont un but linguistique intrinsèque. Il formule des principes logico-philosophiques de son analyse de la structure du discours et présente en particulier, une typologie logique des discours et une définition de leurs conditions de succès. Malheureusement, le sujet de la pertinence pragmatique n'y est jamais abordé. Or, la compétence conversationnelle a pour noyau la pertinence - qui doit expliquer ou justifier les jugements de pertinence. Nous soutenons donc, que pour généraliser la théorie des actes illocutoires à l'analyse de discours et formuler une logique du discours, il faut entre autres choses: enrichir la logique illocutoire actuelle en y intégrant un traitement formel de la pertinence. Un tel traitement formel pourrait prendre comme point de départ les définitions techniques de la pertinence que nous avons élaborées dans le chapitre III.

2- **La Logique propositionnelle dans la sémantique formelle de la théorie des actes de discours**

Nous avons dit plus haut que la T.A.D. est la seule théorie qu'on puisse utiliser comme point de départ dans la construction d'une théorie générale de la signification du locuteur.

Les unités de base de la conversation sont les actes illocutoires, littéraux ou non, que les locuteurs accomplissent au moyen de leur énonciations. La pertinence qu'un locuteur exprime en accomplissant un acte illocutoire proféré dans le contexte d'une énonciation dépend du but propre aux interlocuteurs. La question de la pertinence n'est pas celle de savoir comment choisir des buts illocutoires en accord avec certains buts conversationnels. Les actes illocutoires que nous utilisons ont non seulement un contenu mais aussi une force et la pertinence n'est pas seulement une question de contenu, elle a aussi rapport à la force. Bien sûr, la pertinence concerne l'aspect illocutoire et l'aspect propositionnel. Nous pouvons parler pertinemment seulement au niveau propositionnel en faisant des énonciations telles que des exclamations, des questions, etc., mais la force de ces actes pourrait ne pas être appropriée dans le contexte d'une conversation: dans un interrogatoire (enquête) policier, par exemple, l'accusé pourrait répondre à toutes les questions en énonçant

d'autres questions ou en se limitant à des exclamations.

Dans le programme de la T.A.D. il existe une sémantique formelle, c'est-à-dire, une théorie de la signification des énoncés à côté d'une logique illocutoire, laquelle rend compte des aspects illocutoires de la signification, et une logique propositionnelle qui rend compte des aspects vériconditionnels. Comme nous l'avons déjà remarqué plus haut, chaque acte illocutoire est de la forme $F(P)$. La logique illocutoire tient compte des aspects illocutoires de la signification F ; la théorie des propositions tient compte des aspects vériconditionnels P , et la sémantique générale, des aspects vériconditionnels et illocutoires $F(P)$. Mais cette théorie présuppose le développement d'une théorie des propositions, ou plus précisément, d'une théorie des propositions faisant partie de la sémantique générale.

La logique propositionnelle a été proposée par Vanderveken pour prendre en considération le fait que les propositions ont des conditions de vérité et qu'elles sont des sens d'énoncés aussi bien que des contenus d'actes illocutoires et d'attitudes propositionnelles.

D'un point de vue logique, les propositions sont porteurs de valeurs de vérité tandis que les énoncés sont des entités linguistiques formées de sons ou de caractères et dont les valeurs de vérité dépendent de leur sens dans un contexte, c'est-à-dire du fait que la proposition exprimée soit vraie ou

fausse. Par exemple, un énoncé déclaratif est vrai selon que son sens est une proposition vraie, et un acte illocutoire assertif est vrai ou faux selon la valeur de vérité de son contenu propositionnel. D'un point de vue philosophique, la valeur de vérité d'une proposition est le vrai quand l'état de choses qu'elle représente existe dans le monde. Et d'un point de vue cognitif, comprendre une proposition exige que l'on comprenne quel état de choses doit exister dans le monde pour qu'elle soit vraie (même si on ne sait pas si cet état de choses est réalisé). Ainsi, comprendre une proposition c'est appréhender sa forme logique.

Cette forme logique peut être décrite de la façon suivante. Chaque proposition a des *constituants propositionnels* qui sont des sens (des propriétés, des relations ou d'autres attributs) logiquement liés en termes de prédication de façon à représenter un état de choses. Ainsi, par exemple, la proposition qui est le sens de l'énoncé (1) "Le président français était le mois dernier en Irlande" dans le contexte de cette énonciation, a deux constituants propositionnels, qui sont le sens de la description définie "le président français" et la propriété exprimée par le prédicat "était le mois dernier en Irlande" dans ce contexte. Ces constituants propositionnels sont logiquement liés par la relation de *prédication*, parce que la propriété exprimée par le prédicat est prédiquée à la dénotation de la description définie dans cette proposition. Ainsi, cette proposition identifie un état de choses possible et elle a des conditions de vérité. Dans cette optique, un locuteur comprend une proposition si et seulement s'il comprend quels attributs certains objets doivent avoir dans le monde pour que cette proposition soit vraie. (Vanderveken, 1988, p. 84)

En bref, les propositions sont des contenus d'actes illocutoires: elles sont ce qui est exprimé par les indicateurs du contenu propositionnel. Les

sens des expressions référentielles, les propriétés, les relations ou autres attributs, les notions modales, comme la notion de nécessité, ou les modalités temporelles, le temps des verbes, etc. sont des constituants des propositions, c'est-à-dire, des indicateurs du contenu propositionnel, et ils font une contribution à la détermination des conditions de vérité du contenu propositionnel. En raison de cela, nous croyons qu'une logique propositionnelle générale doit prendre en considération des notions modales, temporelles, etc.

Comme nous venons de le remarquer, la logique illocutoire analyse la forme logique des actes de discours de la forme $F (P)$. La logique illocutoire traite des aspects illocutoires de la signification (les aspects liés à la force F) et la logique propositionnelle, des aspects liés aux contenus propositionnels P . Au sujet de la logique propositionnelle, nous avons déjà souligné dans le chapitre III, que l'identité des conditions de vérité (ou la stricte équivalence) n'est pas un critère suffisant d'identité propositionnelle en sémantique. Il faut que le locuteur comprenne les constituants propositionnels de la proposition et la façon dont ils sont liés par les prédications faites dans les propositions atomiques qui forment la proposition; en outre, il faut qu'il comprenne aussi quelles conditions possibles de vérité des propositions atomiques rendent la proposition vraie en chaque contexte. Qu'est-ce que cette logique fait de nouveau sur l'implication entre propositions?

D'un point de vue logique, chaque proposition en *implique strictement* beaucoup d'autres: elle ne peut être vraie dans un monde possible sans que ces autres le soient également. Ainsi la proposition que Paris est la capitale de la France implique strictement toutes les propositions vraies de l'arithmétique. Cependant, ces implications strictes n'engendrent pas d'engagements illocutoires (ou psychologiques). Par exemple, l'on peut affirmer que Paris est la capitale de la France sans pour autant affirmer que $2 + 2 = 4$. (Vanderveken, 1992, p. 23)

Il existe une compatibilité restreinte de l'implication forte avec l'accomplissement des buts illocutoires (pour lesquels la direction d'ajustement est non vide). En logique illocutoire, une proposition *P implique fortement* une autre proposition *Q* si et seulement si 1) toutes les propositions atomiques de *Q* sont également des propositions atomiques de *P* et 2) toutes les conditions possibles de vérité des propositions atomiques qui rendent *P* vraie dans un contexte rendent également vraie *Q* dans ce même contexte (voir Vanderveken, 1999d).

Comme le souligne Vanderveken, en vertu de leur compétence linguistique, les locuteurs réalisent cognitivement l'implication forte: quand une proposition *P* implique fortement une autre *Q*, il n'est pas possible d'appréhender *P* sans comprendre que sa vérité implique celle de *Q*. En d'autres mots, on ne peut pas avoir *P* à l'esprit sans avoir *Q* à l'esprit et sans savoir que si *P* alors *Q*. Alors, il s'agit d'une métalogue en fait, et tout cela

renvoie à l'implication forte qui est décidable.

Maintenant, nous posons la question suivante: notre logique de la pertinence **P** est-elle compatible avec la logique propositionnelle intégrée dans la sémantique formelle de la T.A.D.?

Notre réponse est qu'il s'agit de deux choses différentes. Comme nous pouvons le vérifier dans Vanderveken (1991, p. 25), le système formel proposé par lui-même est une axiomatisation correcte et complète²⁹ des lois de sa logique propositionnelle où le symbole ' $\rightarrow$ ' exprime l'implication matérielle, ' $\Box$ ' la nécessité universelle, ' $>$ ' l'inclusion de propositions atomiques et où l'implication stricte ' $\multimap$ ' et l'implication forte ' $\vdash$ ' sont définies par abréviation. Parmi les 17 schémas d'axiomes du Calcul Propositionnel présentés dans son système, on y trouve: les schémas d'axiomes de la logique des connecteurs de vérité (propositionnels) de Church (1956), c'est-à-dire, les axiomes de la logique classique; les schémas d'axiomes de la logique modale **S5**, de C.I. Lewis (1918), et des schémas d'axiomes naturels pour la relation d'identité, etc. De cette façon, la logique propositionnelle de Vanderveken suit simplement la logique classique. Il a élaboré la logique illocutoire en y intégrant une logique propositionnelle qui a comme logique

²⁹ Correcte et complète au sens que nous avons défini au chapitre II: un système axiomatique *S* est correct quand tous ses théorèmes sont valides; *S* est complet quand toute

sous-jacente la logique classique. Mais, pour ce qui est de notre logique de la pertinence P , il s'agit d'un système pertinent, ou plus précisément, d'un sous-calcul de la logique de la pertinence qui a été construit à partir du système R de Anderson et Belnap (1975), (voir le chapitre III). Nous avons montré que les logiciens de la pertinence donnent au connectif ' $\rightarrow$ ', une autre signification que celle de la logique classique si bien que des formules comme $A \rightarrow (B \rightarrow A)$ et $\sim A \rightarrow (A \rightarrow B)$ ne sont plus valides. La forme et la valeur de vérité d'une proposition sont les seules propriétés que intéressent la logique classique tandis que les logiciens de la pertinence veulent capturer la relation de pertinence entre A et B qui existe dans de nombreux usages de "Si A alors B ". Ainsi, une implication est pertinente quand l'antécédent suffit pertinemment au conséquent.

Il n'y a donc pas de divergence entre la logique classique (la logique sous-jacente à la logique propositionnelle de Vanderveken) et la logique de la pertinence, en particulier notre logique P . Nous pouvons voir qu'il ne s'agit que de deux significations différentes données au connecteur ' $\rightarrow$ ' selon l'objectif de chacune de ses logiques.

Notre logique de la pertinence P (aussi intuitionniste et paraconsistante) pourrait-elle être intégrée dans la logique propositionnelle de

la logique illocutoire, c'est-à-dire, dans la logique des contenus propositionnels des actes de discours?

Notre réponse est oui. La logique illocutoire analyse la forme logique des actes de discours de la forme $F (P)$ pour traiter des aspects illocutoires de la signification (les aspects liés au F). La logique propositionnelle doit rendre compte des aspects de signification liés aux contenus propositionnels P . Vanderveken formule toute l'axiomatique de sa logique en y intégrant des axiomes pour l'implication vérifonctionnelle, l'implication stricte (C. I. Lewis), et l'implication forte. Dans une sémantique plus générale du discours, il faut ajouter un traitement pour les propositions modales, temporelles et d'action - voir ses articles (Vanderveken, 1997a, 1999b et 1999e) où il a formulé un système plus général. Cette nouvelle formulation de la logique propositionnelle prédicative de Vanderveken est une grande contribution à la sémantique actuelle du discours. Nous croyons qu'il serait enrichissant d'ajouter la logique de la pertinence P et celle du conditionnel contraire au fait (D. Lewis) à la nouvelle logique propositionnelle des actes de discours pour rendre compte des autres significations que l'on peut donner au "Si ... alors" du langage naturel.

Comme nous venons de le remarquer dans la section 2 de ce chapitre, pour généraliser la théorie des actes illocutoires à l'analyse du

discours, Vanderveken a besoin d'y intégrer un traitement formel de la pertinence. Cela permettra d'expliquer mieux sa généralisation de la maxime de quantité qui, d'après nous, fait appel à la maxime de pertinence. C'est, comme nous l'avons déjà dit, cet ordre d'application qui permet à l'auditeur d'identifier si l'objectif du locuteur est approprié, si son acte de discours est pertinent et ainsi évaluer si la maxime de quantité a été respectée.

Alors, si Vanderveken veut que sa logique illocutoire rende compte du problème de la *pertinence*, il doit ajouter à la logique propositionnelle, une logique de la pertinence quelconque (pas nécessairement notre sous-calcul P). Étant donné que l'implication matérielle classique et des axiomes de $S5$ (modale) ne servent pas pour discuter de la pertinence, nous croyons que si Vanderveken prenait la logique de la pertinence comme une logique sous-jacente à sa logique propositionnelle dans le but de l'intégrer à la sémantique formelle de la T.A.D., alors il pourrait rendre compte de la *pertinence* dans une analyse générale du discours. Vanderveken pourrait prendre pour base propositionnelle une logique de la pertinence en plus d'assumer $S5$ dans sa logique actuelle; donc, il pourrait compter sur les axiomes du système pertinent d'Anderson et Belnap (1975) qui tiennent compte aussi des modalités.

À ce propos, nous croyons, comme nous l'avons déjà dit, que notre logique de la pertinence P , qui est intuitionniste et paraconsistante, pourra être

utilisée comme une partie de la logique propositionnelle des actes de discours. Puisque la discussion de la T.A.D. porte aussi sur des problèmes d'ordre pragmatique et si une telle théorie prend en considération la nécessité d'y traiter de la *pertinence*, un outil servant à l'analyse de jugements de pertinence serait effectivement une logique de la pertinence. Nous insistons sur ce point: à l'exemple des logiques intuitionnistes, la logique de la pertinence est la plus proche du langage commun que toute autre logique dont nous disposons.

Néanmoins, il y aurait peut-être quelques difficultés concernant, par exemple, la notion de *co-entailment* $\models P \leftarrow \rightarrow (P \wedge (Q \vee \sim Q))$ dans la logique de la pertinence n'est pas la même que l'identité propositionnelle $\models P = (P \wedge (Q \vee \sim Q))$, c'est-à-dire, $\leftarrow \rightarrow \neq =$. Mais, nous croyons que Vanderveken pourrait arranger la logique de la pertinence **P** de telle façon que *co-entailment* soit plus proche de l'identité propositionnelle.

3. Applications de la logique de la pertinence *P* aux forces illocutoires des actes de discours

Du fait que nous avons effectivement l'habileté de faire des jugements de pertinence, comment décider de ce qui est pertinent ou pas? Comment les gens procèdent-ils pour reconnaître les faits pertinents de l'arrière-plan conversationnel?

La logique de la pertinence *P* que nous avons proposée dans le Chapitre III, peut-elle contribuer à mesurer, évaluer la pertinence d'actes illocutoires tels que les assertions, les interrogations, les exclamations, etc.?

Selon Vanderveken (1988), tout comme toute théorie adéquate des propositions doit procéder à l'analyse de leurs conditions de vérité, de même toute théorie adéquate des actes de discours de la forme F (P) doit procéder à l'analyse de leurs conditions de succès et de satisfaction.

En prenant pour base la T.A.D., nous pouvons considérer une hypothèse comme étant sémantiquement pertinente dans le contexte de son énonciation, quand il est possible de déterminer les conditions de succès et de satisfaction de l'acte illocutoire principal à partir de cette hypothèse.

D'une manière générale, une théorie adéquate de la pertinence devrait avoir pour but d'expliquer comment se crée le lien entre des

énonciations, c'est-à-dire, comment se crée le lien (ou relation) entre des actes de discours dans une conversation. Comme nous l'avons souligné auparavant, d'une façon générale, les agents conversationnels doivent avoir une intentionnalité conjointe de coopérer en faisant des énonciations appropriées qui soient surtout pertinentes. Dans ce cas, la relation de pertinence amène les interlocuteurs à un équilibre de coordination, basé sur la compréhension mutuelle. Dans notre approche une hypothèse est pertinente dans la mesure où elle fait une contribution pour atteindre certaines fins. Supposons, par exemple, que je sois dans la rue avec une cigarette à la main et je demande à quelqu'un s'il a des allumettes et qu'on me réponde coopérativement: "non, je ne fume pas". Dans ce cas, il y a, évidemment, une relation de pertinence entre la réponse de l'interlocuteur et ma demande.

Ainsi, pour définir un ensemble d'hypothèses comme étant sémantiquement ou pragmatiquement pertinent, il nous a fallu en premier lieu, nous rendre à l'évidence (que la réalité nous présente empiriquement) que la conversation ne devient intelligible que par le biais d'une coordination d'actions entre des interlocuteurs utilisant non pas seulement leur habileté linguistique (la reconnaissance de la signification linguistique d'un énoncé par exemple) mais aussi d'autres recours, lesquels sont essentiels pour identifier leur intention communicative. Comme nous l'avons déjà remarqué dans les

chapitres I et le chapitre précédent, l'un de ces recours est, selon la T.A.D., le principe que la conversation est régie par des principes généraux de coopération (cf. Grice 1975) lesquels permettent aux interlocuteurs de reconnaître l'existence d'un (des) but (s) réciproques propres tels que communiquer leurs croyances, désirs, etc. C'est pourquoi, ils essaient de collaborer mutuellement pour atteindre de tels buts. La coopération (fonctionnant ainsi comme une présupposition) joue un rôle essentiel dans l'activité conversationnelle car elle fonctionne comme un repère pour bien orienter les interlocuteurs et leur permettre de créer les conditions adéquates à la compréhension mutuelle. D'autres recours entrent aussi dans la coordination d'actions: les régularités naturelles (biologiques) et sociales (culturelles); les informations contextuelles; la capacité de l'auditeur à faire des inférences, etc. La capacité de l'auditeur à faire des inférences est fondamentale pour la réussite de la communication lorsque la reconnaissance de la signification linguistique d'une énonciation n'est pas suffisante pour la compréhension.

Comme nous pouvons le remarquer dans le quotidien, dans la plupart des cas, les conversations ne sont pas, au contraire des actes illocutoires, pourvues d'un objectif linguistique qui leur soit propre (mais leurs parties le sont). Dans le cours d'une conversation, les locuteurs peuvent changer le sujet et accomplir un acte de discours non approprié sans violer des

règles constitutives ou des conditions de succès de la conversation qui se poursuit quand même.

À partir des définitions techniques de la pertinence et de la logique de la pertinence *P*, que nous avons élaborée dans le chapitre III, nous allons analyser par la suite un *jugement de pertinence* pour expliquer comment les gens procèdent afin de décider ce qui est pertinent ou pas, c'est-à-dire, comment les interlocuteurs reconnaissent des informations (ou des indices) qui peuvent être pertinents pour comprendre la signification du locuteur, qu'elle soit littérale ou non littérale.

En prenant pour base la T.A.D., voyons, à titre d'exemple, comment nos définitions de la pertinence, proposées dans le chapitre III, permettent d'expliquer des *jugements de pertinence* que nous faisons intuitivement dans le quotidien. Considérons le contexte formé par les hypothèses suivantes:

(1) Dans une banque deux individus attendent, au milieu d'une longue file, leur tour à la caisse. Monsieur *X* se retourne vers Monsieur *Y* et utilise l'énoncé suivant (en montrant son intention de sortir de la file):

Monsieur, j'ai besoin d'une minute pour chercher ma carte de crédit que j'ai oubliée à la réception.

Monsieur *Y*: D'accord!

Monsieur *X* part très rapidement vers la réceptionniste.

Reconstruisons les étapes nécessaires, à savoir, l'ensemble d'énoncés Δ représentant les hypothèses que l'interlocuteur fait dans ce contexte³⁰ pour découvrir comment fait *Y* pour comprendre que *X* accomplit un acte de discours littéral, mais que son intention est d'accomplir principalement un autre acte illocutoire non littéral, c'est-à-dire, de signifier plus que ce qu'il dit. En d'autres mots, en prenant comme base le modèle proposé par Searle (1979) des actes de discours indirects, de l'ironie et des métaphores (dont on a parlé dans la section 4, chapitre antérieur) et en suivant notre définition de la pertinence, voyons quelle conséquence pragmatique l'information de *X* peut entraîner pour Monsieur *Y*. Considérons les hypothèses 1 à 7 suivantes:

h_1 : *X* dit: Monsieur, j'ai besoin d'une minute pour chercher ma carte de crédit que j'ai oubliée à la réception.

Monsieur *Y*: D'accord! (faits de conversation).

h_2 : *Y* suppose que *X* est coopératif et que, pour cette raison, son information devra lui fournir des indices suffisants pour traiter

l'hypothèse de X dans le contexte C afin de lui faire inférer une conséquence pragmatique et comprendre ce qu'il veut dire (principe de la coopération conversationnelle de Grice).

h_3 : Y sait donc que le but illocutoire primaire de l'énonciation de X diverge de son but illocutoire littéral, c'est-à-dire, Y infère que X veut probablement dire plus qu'il ne dit (inférence de h_2 et théorie des actes de discours).

h_4 : Il y a beaucoup de monde à la banque et quelqu'un qui veut se diriger à la caisse doit attendre son tour dans la file (information factuelle d'arrière-plan en commun - pratique ou régularité sociale).

h_5 : Quelqu'un qui quitte sa place, va la perdre (inférence à partir de h_4).

h_6 : Une manière pour X ne pas perdre sa place dans la file est de demander à une personne qu'elle la lui garde et lui serve de témoin du fait que X occupait cette place-là quand X y retournera (information factuelle d'arrière-plan en commun).

h_7 : Y comprend donc que le but illocutoire primaire de X est probablement de lui adresser une demande: pouvez-vous garder ma place pendant que je vais à la réception? (inférence de h_3 , h_5 et h_6).

³⁰ Rapellons que tous les traits saillants du contexte peuvent être traduits par des énoncés.

Cette demande nous la considérons comme notre hypothèse φ .

L'énoncé utilisé par X dans le contexte C a été accepté par Y comme pertinent dans la mesure où Y a pu en dériver une conséquence pragmatique qui a contribué à atteindre certaines fins dans la conversation où elle a pris place (le contexte).

Nous soulignons que dans notre travail, notre préoccupation n'est pas d'expliquer par quel mécanisme les allocutaires peuvent identifier les faits de l'arrière-plan conversationnel qui sont pertinents pour comprendre les énonciations accomplies par des actes de discours dans un contexte donné. Nous avons limité nos investigations à la pertinence tout en proposant dans le Chapitre III, une notion théorique adéquate de la pertinence qui puisse être utilisée comme un point de départ dans l'explication de la production et de la compréhension du discours oral et écrit. Nous avons proposé en outre, une logique de la pertinence ***P***, laquelle pourra, avec quelques limitations, traiter formellement cette question. Or, notre logique de la pertinence ***P*** sert à analyser la notion de conséquence (implication) pragmatique et à mieux définir la notion de la pertinence dans la pragmatique générale. Mais, d'après nous, une bonne définition de la pertinence devrait rendre compte de tous les types d'illocutions, c'est-à-dire, elle devrait contribuer à rendre compte de la

capacité des gens de juger la pertinence d'une question, d'une exclamation, d'une promesse, etc. Cela étant dit, pourrions-nous étendre la logique de la pertinence *P* à tous les types de forces illocutoires?

Certainement, oui. Pour étendre la logique de la pertinence *P* aux forces illocutoires, nous limitons notre travail à la construction de quelques exemples au niveau des forces illocutoires de chaque catégorie (assertive, engageante, directive, déclarative et expressive). Car ce sont les forces illocutoires *les plus simples* dans l'usage du langage et cela suffit pour montrer que si la logique basique des propositions dans la sémantique formelle de la T.A.D. était une logique de la pertinence, alors les forces illocutoires pourraient être analysées aussi du point de vue de la pertinence, comme par exemple, si une certaine question est pertinente dans le contexte de son énonciation.

Selon la T.A.D., il y a cinq et seulement cinq forces illocutoires primitives³¹ dans l'usage du langage. Premièrement, la force illocutoire primitive d'assertion a le but assertif de représenter comment les choses sont dans le monde (par exemple, les prédictions, témoignages, conjectures, assertions, objections, etc.). Cette force, nommée en français par le verbe performatif "affirmer", est réalisée syntaxiquement dans le type des énoncées

³¹ C'est à partir de la force illocutoire primitive, en ajoutant ou en retirant certaines conditions sur chaque composante, que l'on obtient toutes les autres forces de la catégorie.

déclaratifs. La force illocutoire primitive d'engagement a le but engageant d'engager le locuteur à accomplir une action future dans le monde (par exemple, les promesses, vœux, recommandations, supplications, menaces). Cette force, nommée par le verbe performatif "s'engager", n'est pas réalisée syntaxiquement dans un mode verbal en français. La force illocutoire directive primitive a le but directif d'essayer de faire en sorte, de façon linguistique, que l'auditeur accomplisse une action future dans le monde (par exemple, les menaces, promesses, vœux, recommandations). Cette force est réalisée syntaxiquement en français dans le type des énoncés impératifs. La force illocutoire primitive de déclaration a le but illocutoire déclaratif de rendre existant un nouvel état de choses représenté par le contenu propositionnel d'une énonciation par cette énonciation même (par exemple, les actes de congédier, d'endosser, de définir, de capituler, d'excommunier). Cette force, nommée par le verbe performatif "déclarer" est réalisée syntaxiquement dans les marqueurs complexes des énoncés performatifs. La force illocutoire primitive expressive a le but expressif d'exprimer des états psychologiques du locuteur à propos d'états de choses dans le monde (par exemple, les excuses, remerciements, condoléances, félicitations). Cette force est réalisée syntaxiquement en français dans le type des énoncés exclamatifs.

Maintenant nous allons vérifier si notre définition de la pertinence

rend compte de ces forces illocutoires. Pour simplifier notre analyse, construisons un seul exemple typique de pertinence dans lequel nous présentons toutes les forces illocutoires:

Un voleur donne l'assaut à un magasin dans lequel il y a très peu de gens. Il vient de mettre un couteau sous la gorge d'un otage pendant qu'il demande de l'argent au propriétaire du magasin. Il est visiblement saoul et ne réussit pas à distinguer les choses très clairement. En commençant à ouvrir la caisse, le propriétaire lui demande:

- Tu vois le contenu de la caisse? Pas besoin de voler. La moitié est à toi. Prends-le vite et pars tranquillement.

À quoi le voleur répond:

-Non. Je veux tout l'argent qui est dans la caisse.

Le propriétaire:

-Dommage! Tu viens de perdre ta chance de ne pas voler.

Et, tout à coup, le propriétaire voit une perceuse portative qui est à côté de la caisse, la saisit dans ses mains, et comme s'il s'agissait d'une arme, la pointe vers le voleur en lui disant:

- Laisse l'otage tout de suite sinon je tire.

Silence... Les gens regardent le voleur en silence. Le voleur crie au propriétaire:

- Ce n'est pas une arme ce que tu as dans les mains! C'est une perceuse.

Le propriétaire:

- Quel dommage! J'imagine que tu es saoul au point d'avoir la vision troublée.

- Non, je suis sûr que c'est une perceuse.

- Ah oui? Alors, fais quelque chose à l'otage et je déchargerai aussitôt mon arme à feu sur toi.

À ce moment, un client entre dans le magasin et est surpris par la scène: le voleur, l'otage, les gens (en silence) et le propriétaire avec la perceuse à la main. Le voleur lui demande très rapidement:

- Dis-moi vite ce que le Monsieur a dans les mains.

Le client confus reste en silence. Le voleur insiste:
 - Allons, dépêche-toi, et dis-moi: qu'est-ce que le Monsieur de la caisse a dans les mains?
 Le client regarde les gens, le propriétaire, et très calmement répond au voleur:
 - C'est une arme à feu.
 Et finalement, le voleur laisse tomber par terre son couteau.

Le contexte décrit ci-dessus exemplifie la relation de pertinence entre ensembles d'énoncés exprimant des actes de discours dans un même contexte. C'est-à-dire, chaque acte illocutoire est traité dans le contexte produisant ainsi des conséquences pragmatiques. Il faut souligner ici, la relation étroite entre les objectifs des interlocuteurs et le contexte intentionnel mutuellement manifeste aux participants de cette conversation - ce qui met en évidence la nécessité du rapport entre ce qui est approprié et ce qui est pertinent (quand les individus sont engagés dans un discours intelligent).

Voyons maintenant si notre définition de la pertinence est adéquate pour expliquer les jugements de pertinence relativement à chaque force illocutoire dans l'exemple ci-dessus.

D'abord, pour la force illocutoire et déclarative. Nous concluons que la force illocutoire **déclarative** exprimée par l'énoncé utilisé par le propriétaire "*tu vois le contenu de la caisse? Pas besoin de le voler. **La moitié est à toi.***", dans le contexte de cette conversation, est pertinente car, selon notre

définition de la pertinence, elle entraîne des conséquences pragmatiques dans ce contexte. Par cette énonciation, le propriétaire a pour but de produire chez le voleur, la conséquence pragmatique de lui rendre manifeste son intention de lui donner la moitié de l'argent dans la caisse. À quoi le voleur répond: "*Non. Je veux tout l'argent qui est dans la caisse.*" (en confirmant le degré 1 de pertinence pour son interlocuteur). Alors, l'un et l'autre des énoncés sont, dans le contexte fixé, pertinents.

D'un autre côté, quel serait le degré de pertinence de cette force déclarative si on la comparait à d'autres énonciations alternatives? Par exemple, supposons que le but du propriétaire étant celui de déclarer au voleur, dans le contexte fixé, son intention de lui donner la moitié de l'argent de la caisse, il dit maintenant: *Tu vois le contenu de la caisse? La moitié n'est plus à moi.* Dans ce cas, à cause de l'ambiguïté sémantique, le voleur pourrait ne pas être en mesure d'inférer que le propriétaire voulait lui rendre manifeste son intention de lui donner la moitié de l'argent de la caisse. Évidemment le voleur pourrait penser que le locuteur avait seulement l'intention de l'informer que la moitié de l'argent de la caisse appartenait déjà à une autre personne. Et, dans ce cas, puisque l'ensemble des hypothèses que le voleur fait sur le contexte serait différent de l'ensemble des hypothèses que le locuteur fait sur le contexte en question, le degré de pertinence de l'énoncé utilisé par le locuteur (le

propriétaire) serait le degré zéro (ou degré minimum) de pertinence pour l'interlocuteur.

Deuxième exemple, relatif à la force illocutoire expressive. Comme nous pouvons le voir, la force illocutoire primitive **expressive** dans l'énonciation du propriétaire "***Domage!** Tu as perdu ta chance de ne pas voler.*" est pertinente car elle entraîne, dans le contexte de cette énonciation, la conséquence pragmatique que le propriétaire veut exprimer sa tristesse à propos de l'attitude du voleur. Dans ce cas, le degré de pertinence de l'énoncé du propriétaire, pour son interlocuteur, est le degré 1. Supposons maintenant, qu'au lieu de cette énonciation, le propriétaire fasse l'énonciation suivante, dans le but d'exprimer de la tristesse: "***Je proteste!** Tu as perdu ta chance de ne pas voler.*" Alors, puisqu'une protestation exprime d'une façon formelle une désapprobation de la part du locuteur, il se pourrait que le voleur n'infère pas la conséquence pragmatique voulue par le propriétaire, à savoir, qu'il voulait lui exprimer de la tristesse (et non pas de la désapprobation) par rapport à son comportement. Dans ce cas, le degré de pertinence de la force illocutoire primitive expressive manifestée par cette énonciation, serait le degré 0 (zéro) de pertinence pour l'interlocuteur.

En ce qui concerne la force illocutoire assertive de l'énonciation du voleur. Nous concluons que la force illocutoire **assertive** exprimée par

l'énonciation du voleur "*Ce n'est pas une arme ce que tu as dans les mains! C'est une perceuse.*", dans le contexte de cette conversation, est pertinente car, selon notre définition de pertinence, elle entraîne une conséquence pragmatique dans ce contexte. Par cette énonciation, le voleur a pour but de produire, chez le propriétaire, la conséquence pragmatique de lui rendre manifeste son intention de l'informer de ses doutes quant à ce que le propriétaire prétend tenir dans ses mains. À quoi le propriétaire répond: "*Quel dommage! J'imagine que tu es saoul au point d'avoir la vision troublée.*" (en confirmant le degré 1 de pertinence de l'énoncé du voleur pour le propriétaire).

Mais, comment serait le degré de pertinence de cette force assertive dans le cas où elle serait réalisée par une énonciation alternative? Par exemple, supposons que le but du voleur était d'informer le propriétaire de ses doutes à propos de ce qu'il prétendait tenir dans ses mains, c'est-à-dire, une perceuse, en énonçant malicieusement *Tu ne peux pas tirer sur moi...* Dans ce cas, le propriétaire ne serait pas en mesure d'inférer la conséquence pragmatique désirée par le voleur, à savoir, que ce dernier était sûr qu'il s'agissait d'une perceuse dans les mains du propriétaire. Et de là ne suivrait aucunement la réponse du propriétaire "*J'imagine que tu es saoul au point d'avoir la vision troublée*", car l'ensemble d'hypothèses qu'il fait sur le contexte serait différent de l'ensemble des hypothèses que le voleur fait sur le

contexte; donc, à ce moment là, le degré de pertinence de l'énoncé du voleur pour son interlocuteur serait le degré 0 de pertinence.

Pour ce qui est de la force illocutoire **engageante**, nous déduisons que cette force exprimée dans le contexte de l'énonciation du propriétaire "*Ah oui? Alors, fais quelque chose à l'otage et je déchargerai aussitôt mon arme à feu sur toi*", est pertinente, car elle y entraîne des conséquences pragmatiques. En effet, en faisant cette énonciation, le locuteur a le but de manifester à l'allocutaire son intention de lui confirmer (c'est un engagement) qu'il a une arme à feu dans ses mains. Et cela confirme le degré 1 de pertinence de son énoncé pour l'allocutaire dans le contexte de cette conversation. D'autre part, en comparant le degré de pertinence de cette force engageante avec une autre énonciation possible, par exemple, en supposant maintenant que le propriétaire énonce: "*Alors, fais quelque chose à l'otage et je te ferai une perforation*", cela permettrait à l'allocutaire d'inférer une autre chose différente de ce que le propriétaire avait l'intention de lui communiquer; en effet, l'allocutaire pourrait inférer que le propriétaire lui confirme exactement qu'il avait une perceuse dans ses mains. Et, puisque l'allocutoire ne serait pas en mesure d'inférer la conséquence pragmatique voulue par le propriétaire, nous concluons que le degré de pertinence pour l'allocutoire, serait le degré 0 de pertinence.

Finalement, la force illocutoire directive. Selon notre définition de pertinence, la force illocutoire **directive** manifestée par l'énonciation du voleur “... *qu'est-ce que le monsieur de la caisse a dans les mains?*” est pertinente dans le contexte fixé, parce qu'elle entraîne la conséquence pragmatique que le voleur a pour but de faire en sorte que l'allocutaire accomplisse un acte de discours futur déterminé par le contenu propositionnel de la question, c'est-à-dire, que l'allocutaire réponde à sa question. Et, du fait que l'allocutaire a répondu pertinemment à cette question, son degré de pertinence pour lui est le degré 1 de pertinence. Mais si, dans ce contexte, le voleur posait la question à l'allocutaire, par une autre énonciation, comme par exemple, *Quelle sorte d'arme a le monsieur dans les mains?*, il se pourrait que l'allocutaire ne soit pas capable de répondre à la question de la façon anticipée par le voleur. Si tel était le cas, le degré de pertinence de cette question, pour l'allocutaire, serait le degré 0 de pertinence.

Étant donné l'application de notre définition de la pertinence aux exemples précédents, nous concluons que cette définition de la pertinence, rend compte de la pertinence d'énoncés exprimant des actes illocutoires avec des forces illocutoires déclarative, expressive, assertive, engageante et directive.

CONCLUSION

L'une des caractéristiques fondamentales de la théorie pragmatique est de tenter d'expliquer la question de la signification du locuteur qu'elle soit littérale ou non dans la communication humaine.

Plusieurs approches pragmatiques qualifiées d'inférentielles ont eu pour base les travaux du philosophe du langage H. P. Grice qui a montré comment expliquer les implicatures conversationnelles. Comme on le sait, dans la conception de Grice, il faut recourir à des principes de coopération et des maximes conversationnelles. Grice a proposé de rendre aussi simple que possible la description sémantique (en particulier de ne pas multiplier sans nécessité les significations linguistiques) pour favoriser la description pragmatique. Cependant ses maximes ne tiennent compte que des énonciations assertives; elles ne formulent pas des principes généraux pour l'analyse de la signification non littérale. Comme Vanderveken (1990a) l'a montré, on peut formuler à partir de la logique illocutoire, des principes généraux pour traiter adéquatement du phénomène de la signification non littérale. Dans sa pragmatique des actes de discours il y a une généralisation des maximes de

quantité et de qualité de Grice, mais il manque une généralisation de la maxime de pertinence.

Comme nous l'avons toujours remarqué dans notre thèse, la pertinence est inhérente à la communication humaine. Si on n'admet pas la pertinence comme une présupposition essentielle dans la description du processus de compréhension verbale, alors comment expliquer le besoin qu'ont les interlocuteurs de chercher les données qui peuvent leur permettre de juger ou d'évaluer la signification d'une énonciation? Par exemple, quel critère les interlocuteurs d'une langue emploient-ils pour "justifier" leur propos? Nous croyons qu'une théorie de la pertinence pourra apporter des réponses à ce genre de questions.

Comment expliquer formellement la maxime de pertinence de Grice? Il faut auparavant faire un travail d'élaboration et de précision de la notion de pertinence. Nous croyons que la définition que nous avons donnée de la pertinence est un point de départ à tout traitement formel de cette maxime dans une analyse pragmatique générale du discours.

Bien entendu, la portée du domaine pratique où s'insère la pertinence n'est pas restreinte à l'échange des actes de parole eux-mêmes entre des interlocuteurs désireux de collaborer dans la communication. Son domaine d'application est beaucoup plus ample. Il suffit d'observer dans le quotidien

comment les gens font usage de cette notion dans les pratiques sociales communes, comme par exemple, l'accomplissement d'une action commune exigeant de la coordination. Lorsque des interlocuteurs sont impliqués dans une négociation quelconque, on considère que la pertinence est à l'œuvre dans toutes les décisions qu'ils prennent pour tenter de parvenir à un accord. Dans la pratique scientifique, par exemple, c'est la théorie qui détermine les tests et les facteurs expérimentaux pertinents. Du point de vue des actes de discours proférés dans la conversation, l'inaccomplissement d'un acte illocutoire peut être dans certains cas pertinent. Pour mieux situer notre thèse, c'est dans les recherches pragmatiques sur la communication que nous avons précisé la notion de la pertinence.

Cette thèse trace le parcours de la problématique de la pertinence dans les recherches pragmatiques qui intègrent des idées de Grice en particulier, la théorie de Searle et Vanderveken et celle de Sperber et Wilson. Nous avons en outre formulé une nouvelle définition de la pertinence qui capture adéquatement la notion intuitive de pertinence. Les énonciations non pertinentes selon notre définition sont en fait inutiles et donc non pertinentes au sens intuitif. Cette nouvelle définition est utile pour analyser les jugements de pertinence que font les interlocuteurs dans la poursuite d'une conversation.

Sperber et Wilson ont formulé une théorie de la pertinence pour

rendre compte du même problème de la pertinence pragmatique dans la communication. Cependant ils n'ont pas, à notre avis, donné une définition adéquate de la pertinence pragmatique. Ils ont préféré souligner les divergences entre leur approche et les autres approches, celle de Grice et celle de la T.A.D.

Par contre, nous jugeons tous les travaux importants et nous les prenons en considération dans notre définition technique de la pertinence.

Comme nous l'avons remarqué au début de la thèse, cette idée de souligner les convergences entre les théories déjà existantes (plutôt que de chercher les divergences) nous a été inspirée par Verschueren (1979) dont l'article "À la recherche d'une pragmatique unifiée" souligne l'importance d'unifier des théories et des concepts proposés par les linguistes et philosophes.

La préoccupation de Verschueren est d'unifier les théories de la présupposition, des actes de discours et des implicatures conversationnelles, qui ont été conduites isolément. Ainsi Verschueren observe qu'il existe des relations (1) entre les maximes conversationnelles et les règles ou conditions formulées par Searle pour les actes de discours, (2) entre les implicatures conversationnelles et la force illocutoire, et (3) entre les actes de discours et la présupposition. Il énonce la relation (1) entre les maximes conversationnelles et les règles pour les actes de discours comme suit:

La règle selon laquelle une assertion que *p* ne devrait pas être prononcée à moins que le locuteur n'ait quelque raison de croire à la vérité de *p* (ce qui est une des conditions préparatoires de l'assertion) est manifestement une application spécifique de la maxime de qualité "Ne dites aucune chose qui ne soit fondée". En plus, les maximes de qualité sont liées aux conditions de sincérité. Et la condition selon laquelle le locuteur doit croire que l'auditeur est capable de faire ce qui lui est ordonné (ce qui est une condition préparatoire de l'ordre) est une application spécifique de la maxime de pertinence. (Verschueren, 1979, p. 275)

Et en ce qui concerne la relation (2) entre les implications conversationnelles et la force illocutoire:

Considérez la notion d'un *acte de langage indirect*. Une phrase comme "Il fait froid ici" a la forme d'une assertion, mais dans certaines circonstances elle peut véhiculer une demande de fermer la porte (...) seulement parce qu'autrement la maxime de pertinence [quantité] (*sic*) serait violée (dans certaines circonstances, il n'est pas vraiment nécessaire de dire quelque chose que chacun peut sentir par lui-même): l'interprétation alternative est construite (au moyen d'une implication conversationnelle) pour restaurer la balance de coopération qui a l'air d'être troublée. Ainsi l'implication conversationnelle est le processus par lequel la force primaire de beaucoup d'actes de langage indirects est obtenue. (Verschueren, 1979, p. 275-276)

Finalement, pour ce qui est de (3), les actes de discours et la *présupposition* (dont les définitions se diversifient dans la littérature contemporaine):

Considérons la définition suivante de la présupposition, proposée par C. J. Fillmore (1971: 380): les phrases d'une langue naturelle sont employées pour poser des questions, pour ordonner, pour asserter des propositions, pour exprimer des sentiments, etc.; on peut identifier les présuppositions d'une phrase et les conditions qui doivent être satisfaites pour qu'on puisse l'employer dans une de ces fonctions. (...) La même définition a été adoptée par D. Cooper (1974) et, sous une forme modifiée, par P. Harder et C. Kock (1976). (Verschueren, 1979, p. 276)

Nous aimerions remarquer que nous pouvons aller plus loin et étendre cette analyse de Verschueren en ajoutant un quatrième item: (4) Nous soutenons qu'il existe, en philosophie du langage, une convergence en ce qui concerne le concept de *pertinence* et son rôle dans la pragmatique, plus spécifiquement dans les travaux de Grice (1957, 1975), de Keenan (1971), et celui sur la pertinence, de Sperber et Wilson (1986). Nous observons que, en fait, une certaine convergence a lieu dans la définition de la présupposition pragmatique formulée par Keenan: l'énonciation d'une phrase présuppose pragmatiquement que son contexte (linguistique et extralinguistique) est approprié, et le principe de pertinence formulé par Sperber et Wilson: tout acte de communication ostensive communique la présomption de sa propre pertinence optimale.

Alors, pourquoi les étudier isolement?

Selon Verschueren (1979), la pragmatique est l'étude des conditions d'appropriété contextuelle; son idée est que pour unifier les théories de la présupposition, des actes de discours et des implications conversationnelles, la notion unifiant serait celle d'appropriété contextuelle.

Je crois que la pragmatique n'est pas une discipline linguistique de même

niveau que la phonologie, la syntaxe et la sémantique, mais plutôt un point de vue sous lequel on peut envisager la conduite des entités sémantiques, syntaxiques et même phonologiques. (Verschueren, 1979, p. 277)

Rappelons les conditions d'appropriété contextuelle telles que Verschueren (1979) les définit (voir chapitre VI). Selon lui, les conditions d'appropriété sont des présuppositions logiques ainsi que pragmatiques; la félicité des actes de discours déterminé au moyen de conditions et de règles; les implications conversationnelles, par exemple, quand le locuteur affirme "Je n'ai plus d'essence" et, à son tour, l'auditeur répond "Il y a un garage de l'autre côté de l'église", l'implication conversationnelle que le garage est ouvert et qu'il y a de l'essence, est indispensable pour l'appropriété de la réponse; le principe que le locuteur respecte les maximes conversationnelles de Grice, etc.

Alors, d'après nous, ces conditions d'appropriété contextuelle contiennent aussi le respect du principe de pertinence de Sperber et Wilson, c'est-à-dire, nous considérons tel principe comme une condition d'appropriété contextuelle d'un énoncé (ou conformément à leur terminologie, d'un "acte de communication ostensive"). En fait, dans le principe de pertinence de Sperber et Wilson, il existe une "**présomption**"³² de pertinence optimale d'un acte de

³² C'est nous qui soulignons.

communication” et qui pourrait *ipso facto* être considérée comme un phénomène relevant des présuppositions pragmatiques, quoique chez Sperber et Wilson cette présupposition concerne la *pertinence* d’un acte de communication par rapport à un contexte, tandis que chez Keenan, la présupposition concerne l’*appropriété* d’une énonciation par rapport à un contexte.

En tant que présupposition pragmatique, le principe de pertinence de Sperber et Wilson devient *ipso facto* l’une des conditions d’appropriété contextuelle d’un énoncé. Nous observons que tel principe est, d’une certaine façon, apparenté à la définition de présupposition pragmatique de Keenan; or, puisqu’il existe un rapport entre la définition de la présupposition pragmatique de Keenan et la définition de la présupposition de Fillmore (et autres), lesquelles ont à leur tour, des relations avec les actes de discours etc., nous indiquons qu’il serait enrichissant de joindre toutes ces approches, à celle de Sperber et Wilson, de façon à traiter adéquatement les phénomènes pragmatiques, avec la préoccupation majeure de décrire l’interprétation complète d’une énonciation (ayant pour base l’interprétation partielle de l’énoncé à son niveau phonologique, syntaxique et sémantique, et en faisant appel au contexte et à des mécanismes inférentiels que les locuteurs sont capables de réaliser).

Et ce sont toutes ces approches réunies qui nous ont guidé dans l'élaboration de notre travail, c'est-à-dire, c'est en prenant en considération cette convergence trouvée dans les approches de H. P. Grice (maximes et implications conversationnelles), J. Searle et D. Vanderveken (théorie des actes de discours), C. J. Fillmore, E. L. Keenan, D. Cooper, P. Harder, C. Kock, D. Sperber et D. Wilson (présuppositions), que nous avons formulé notre définition de la notion de pertinence pour la communication orale ou écrite.

BIBLIOGRAPHIE

- ANDERSON, A. R. et N. D. BELNAP (1975), *Entailment: the Logic of Relevance and Necessity*, vol. I, Princeton University Press.
- ARAÚJO, A. L. de, E. H. ALVES et J. A. D. GUERZONI (1987), "Some Relations Between Modal and Paraconsistent Logic", dans *The Journal of Non-Classical Logic*, vol. 2, number 2; 33-44.
- ARISTOTE (1997), *Lógica*, dans *Obras*, 2e. ed., traducción del griego, estudio preliminar, preámbulos y notas por Francisco de P. Samaranch, Madrid, Aguilar, 217-561.
- AUSTIN, John L. (1962), *How to do Things with Words*, Oxford, Clarendon Press.
- _____ (1962a), "Performatif-Constatif", dans *Cahiers de Royaumont, Philosophie* n° IV. La Philosophie Analytique, Minuit.
- BACH, Kent (1994), "Conversational Implication", dans *Mind & Language*, vol. 9, No, 2, 124- 162.
- BACH, K. et R. M. HARNISH (1979), *Linguistic Communication and Speech Acts*, Cambridge (Mass.), The MIT Press.
- BEAUGRANDE, R.-A. de et W. U. DRESSLER (1981), *Introduction to Text Linguistics*, London/New York, Logmam.
- BENNETT, J. (1969), "Entailment", dans *Phil. Rev.*, vol. 78, 197-235.
- BERRENDONNER, Alain (1981), *Éléments de pragmatique linguistique*, Paris, Minuit.
- BLANCHÉ, Robert et J. DUBUCS (1996), *La logique et son histoire*, Paris, Armand Colin.
- BROWN, G. et G. YULE (1983), *Discourse Analysis*, Cambridge, Cambridge University Press.
- CARNAP, R. (1956), *Meaning and Necessity*, Chicago, University of Chicago Press.
- CARSTON, R. et S. UCHIDA (eds.). (1998), *Relevance Theory: Applications and Implications*, Amsterdam, John Benjamins.
- CHAMBREUIL, Michel et Jean-Claude PARIENTE (1990), *Langue naturelle et logique. La sémantique intensionnelle de Richard Montague*, Berne, Peter Lang.

- CHERNIAK, C. (1986), *Minimal Rationality*, Cambridge Press, MIT Press (Bradford Books).
- CHURCH, A. (1951), "A Formulation of the Logic of Sense and Denotation", dans P. HENLE et al., *Structure, Method and Meaning*, New York.
- COOPER, David E. (1974), *Presupposition (Janua linguarum, series minor 203)*, La Hague, Mouton.
- CORNULIER, Benoît de (1995), *Effets de sens*, Paris, Minuit.
- DASCAL, Marcelo (1977), "Conversational Relevance", dans *Journal of Pragmatics*, 1, 309-328.
- DAVIDSON, D. (1984), *Inquiries into Truth and Interpretation*, Oxford, Clarendon Press.
- DENNET, Daniel C. (1990), *La stratégie de l'interprète: le sens commun et l'univers quotidien*; traduit de l'anglais par Pascal Engel, Paris, Gallimard.
- DESCARTES, René (1966), *Discours de la méthode* (1637); introduction et notes par Et. Gilson, Paris, J. Vrin.
- DIAS, Matias F (1994), "On Elementary Arithmetic" dans *The Journal of Symbolic Logic*, USA, volume 59, no. 2, june, 716-717.
- _____ (1996), "A New Generalization of Goedel's First Incompleteness Theorem", dans *The Bulletin of Symbolic Logic*, USA, volume 2, no. 4, December, 458-459.
- DOMINICY, Marc (1991), Compte-rendu à Dan Sperber et Dreidre Wilson, *La pertinence. Communication et cognition* (1989), *Revue de linguistique française*, LIX (1), 85-91.
- DUCROT, O. (1972), *Dire et ne pas dire. Principes de sémantique linguistique*, Paris, Hermann
- _____ (1980a), "Analyse de textes et linguistique de l'énonciation", dans O. DUCROT et al., *Les mots du discours*, Paris, Minuit, 7-56.
- DUMMETT, Michael (1993), *The Seas of Language*, New York, Oxford University Press.
- _____ (1980b), "Analyses pragmatiques", dans *Communications* 32: 11-60.
- ENGEL, Pascal (1994), *Davidson et la philosophie du langage*, Paris, PUF.

- EPSTEIN, Richard (1995), *The Semantic Foundations - Propositional Logics*, 2nd Edition, Oxford, Oxford University Press.
- FILLMORE, Charles J. (1971), "Types of lexical information", dans D. D. STEINBERG, et L. A. JAKOBOVITS (eds.), *Semantics: An Interdisciplinary Reader in Philosophy, Linguistics and Psychology*, Cambridge University Press, 370-392.
- FREGE, G. (1969), *Les fondements de l'arithmétique*, traduit de l'allemand par Claude Imbert, Paris, Seuil.
- _____ (1971), *Écrits logiques et philosophiques*, traduit par Claude Imbert, Paris, Seuil.
- GAZDAR, Gerald (1979), *Pragmatics - Implication, Presupposition and Logical Form*, New York, Academic Press.
- GENTZEN, Gerhard (1955), *Recherches sur la déduction logique*, traduction et commentaire par R. Feys et J. Ladrière, Paris, PUF.
- GHIGLIONE, Rodolphe et Alain TROGNON (1993), *Où va la pragmatique? De la pragmatique à la psychologie sociale*, préface de L. Sfez, Grenoble, Presses Universitaires de Grenoble.
- GOCHET, Paul et Pascal GRIBOMONT (1998), *Logique*; vol I - *Méthodes pour l'informatique fondamentale*, Paris, Hermes.
- _____ (2000), *Logique*; vol. III - *Méthodes pour l'intelligence artificielle*, Paris, Hermes.
- GRICE, H. P. (1957), "Meaning", dans *Philosophical Review* 66 : 377-388.
- _____ (1975), "Logic and Conversation", dans P. COLE et J.L. MORGAN (eds). *Syntax and Semantics*, vol. 3, *Speech Acts*, New York, Academic Press, 41-58.
- _____ (1978), "Futher notes on Logic and Conversation", dans P. COLE (ed.), *Syntax and Semantics 9 : Pragmatics*, New York , Academic Press, 113-127.
- _____ (1979), "Logique et conversation", dans *Communications - La conversation*, Paris, 30: 57-72.
- HARDER, Peter et Christian KOCK (1976), *The Theory of Presupposition Failure*, Copenhagen, Akademisk Forlag.
- HOBBES, Thomas (1974), *Leviatã ou Matéria, Forma e Poder de um Estado Eclesiástico e Civil (Leviathan, or Matter, Form, and Power of a Coomnwealth Ecclesiastical and Civil, 1651)*, traduit par J. P. Monteiro et M. B. N. da Silva, São Paulo, Abril Cultural.

- HOLDEROIT, David (1987), "Conversational Relevance", dans J. VERSCHUEREN et M. BERTUCCELLI-PAPI (eds.), *The Pragmatics Perspective. Selected Papers from the 1985 International Pragmatics Conference*, Amsterdam/Philadelphia, John Benjamins.
- HUGHES, G. E et M. J. CRESSWELL (1968), *An Introduction to Modal Logic*, London, Methuen and Co. Ltd.
- JACQUES, Francis (1985), "Du dialogisme à la forme dialoguée: sur les fondements de l'approche pragmatique", dans M. DASCAL, *Dialogue*, Amsterdam/Philadelphia, John Benjamins, 27-56.
- JEFFREY, Richard C. (1965), *The Logic of Decision*, The University of Chicago.
- KAPLAN, David. (1975), "How to Russell a Frege-Church", dans *Journal of Philosophy*, (19), 716-729.
- KASHER, Asa (1976), "Conversational Maxims and Rationality", dans A. KASHER (ed.), *Language in Focus*, Dordrecht-Holland, D. Reidel, 197-216.
- KEENAN, Edward L. (1971), "Two Kinds of Presupposition in Natural Language", dans C.J. FILLMORE et D. T. LANGENDOEN (eds.), *Studies in Linguistic Semantics*, New York, Holt, Rinehart & Winston, 45-52.
- KAUFMANN, J. Nicolas (1999), "Écueils des théories de rationalité", dans *Dialogue*, vol. XXXVIII, 801-826.
- LARGEAULT, Jean (1993), *La logique*, Paris, PUF (Que sais-je?, 225).
- LECLERC, André (1985), "La place réservée à la pragmatique dans 'Le procès de la métaphore' de Guy Bouchard", dans *Dialogue*, XXIV, 655-669.
- LEECH, E. (1983), *Principles of Pragmatics*, London, Longman.
- LEVINSON, S. C. (1983), *Pragmatics*, Cambridge, Cambridge University Press.
- LEWIS, C. I. (1912), "Implication and the Algebra of Logic", dans *Mind*, vol. 21, 522-531.
- LEWIS, David (1969), *Convention*, Cambridge, Mass, Harvard University Press.
- _____ (1972), "General Semantics", dans D. DAVIDSON et G. HERMAN (eds.) *Semantics of Natural Language*, Reidel, 169-218.
- _____ (1979), "Scorekeeping in a Language Game", dans *Journal of Philosophical Logic*,

8, 339-359.

MAINGUENEAU, Dominique (1989), *Novas tendências em análise do discurso*; tradução de Freda Indursky, Campinas, Pontes/EdUNICAMP.

MARCUSCHI, Luis A. (1986), *Análise da conversação*, São Paulo, Ática.

MARTINICH, A.P. (ed.). (1990), *The Philosophy of Language*; 2nd Edition, Oxford; Oxford University Press.

MENDELSON, Elliott (1979), *Introduction to Mathematical Logic*, 2nd Edition, New York, D. Van Nostrand Company.

MILLER, Seumas (1991), "Co-ordination, Salience and Rationality", dans *The Southern Journal of Philosophy*, vol. XXIX, No. 3, 359-370.

MOESCHLER, Jacques (1996), *Théorie pragmatique et pragmatique conversationnelle*, Paris, Armand Colin.

MOESCHLER, Jacques et A. REBOUL (1994), *Dictionnaire Encyclopédique de Pragmatique*, Paris, Seuil.

MONTAGUE, R. (1968), "Pragmatics", dans KLIBANSKY, R. (ed.), *Contemporary Philosophy/La philosophie contemporaine*, Florence, La Nuova Italia Editrice, 102-122.

_____ (1970), "Universal Grammar", dans *Theoria*, n° 36.

_____ (1974), *Formal Philosophy*, New Haven, Yale University Press.

PANACCIO, Claude (1991), *Les mots, les concepts et les choses: la sémantique de Guillaume d'Occam et le nominalisme d'aujourd'hui*, Sait-Laurent, Québec/Paris, Bellarmin/Vrin.

PARRET, Herman (1988), *Emunicação e Pragmática*; traduit par E. P. Orlandi et ali., Campinas, UNICAMP.

PEIRCE, Ch. S. (1956), *Philosophical Writings of Peirce*, selected and edited with an introduction by Justus Buchler, New York, Dover Publications.

PENROSE, Roger (1997), *A Mente Virtual (The Emperor's New Mind): sobre Computadores, Mentes e as Leis da Física*; traduit par Augusto J. F. de Oliveira et al., Lisboa, Gradiva.

POLLOCK, John L. (1982), *Language and Thought*, Princeton - New Jersey, Princeton

University Press.

PRIETO, Luis J. (1975), *Pertinence et pragmatique: essai de sémiologie*, Paris, Minuit.

PRIEST, G. et R. SYLVAN (1992), "Simplified Semantics for Basic Relevant Logic", dans *Journal of Philosophical Logic*, vol 21, 217-232.

QUINE, W. V. (1960), *Word and Object*, Cambridge, Mass, MIT Press.

RÉCANATI, François (1981), *Les énoncés performatifs. Contribution à la pragmatique*, Paris, Minuit.

ROUTLEY, Richard, R. K. MEYER, VAL PLUMWOOD et R. T. BRADY, (1982), *Relevant Logics and Their Rivals: 1, The Basic Philosophical and Semantical Theory*, Independence, Ohio, Ridgeview Publishing Company.

RUSSELL, B. (1905), "On Denoting", dans *Mind*.

SCHIFFER, S. (1972), *Meaning*, Oxford, Clarendon Press

SEARLE, John R. (1969), *Speech Acts*, Cambridge University Press.

_____ (1979), *Expression and Meaning*, Cambridge University Press (*Sens et expression*; traduction et préface par Joëlle Proust, Paris, Minuit, 1982).

_____ (1983), *Intentionality*, Cambridge University Press (*L'intentionnalité. Essai de philosophie des états mentaux*; traduit de l'américain par Claude Pichevin, Paris 1985).

_____ (1984), *Mind, Brains and Science*, Harvard University Press.

_____ (1989), "How Performatives Work", dans *Linguistics and Philosophy*, 12 (5), 535-558.

SEARLE, John R., et al. (eds.). (1980), *Speech Act Theory and Pragmatics*, Dordrecht, Netherlands, Reidel.

SEARLE, John R. et D. VANDERVEKEN (1985), *Foundations Illocutionary Logic*, Cambridge University Press.

SEARLE, John R. et al. (eds.). (1990), *(On) Searle on Conversations*, compiled and introduced by H. Parret and J. Verschueren, Amsterdam/Philadelphia, John Benjamins.

SHOENFIELD, J. R. (1973), *Mathematical Logic*, Reading, Mass., Addison-Wesley

- SPERBER, Dan et Deirdre WILSON (1987), "Précis of Relevance: Communication and Cognition", dans *Behavioral and Brain Sciences*, 10, 697-754.
- _____ (1989), *La pertinence. Communication et cognition; (Relevance, Communication and Cognition*, 1986), traduit de l'anglais par Abel Gerschenfeld et Dan Sperber, Paris, Minuit.
- STRAWSON, P.F. (1964), "Intention and Convention in Speech Acts", dans *Philosophical Review* 73: 469-460.
- _____ (1971), *Logico-linguistic Papers*, London, Methuen.
- TARSKI, A. (1956), *Logic, Semantics and Meta-mathematics*, Oxford, Clarendon Press.
- VANDERVEKEN, Daniel (1988), *Les actes de discours: essai de philosophie du langage et de l'esprit sur la signification des énonciations*, Liège-Bruxelles, Pierre Mardaga Editeur.
- _____ (1990), *Meaning and Speech Acts – Vol. I, Principles of Language Use*. Cambridge, Cambridge University Press.
- _____ (1990a), "On the Unification of Speech Act Theory and Formal Semantics", dans P. COHEN et al., *Intentions and Plans in Communication and Discourse*, Cambridge, Mass., MIT Press, 195-220..
- _____ (1991), *Meaning and Speech Acts – Vol. II, Formal Semantics of Success and Satisfaction*, Cambridge, Cambridge University Press
- _____ (1991a), "What is a Proposition?", *Cahiers d'Épistémologie*, Université du Québec à Montréal, Cahier 9103.
- _____ (1991b), "Non Literal Speech Acts and Conversational Maxims", dans E. LEPORE et R. VAN GULICK (eds.), *John Searle and his Critics*, Oxford, Blackwell, 371-384.
- _____ (1992), "La théorie des actes de discours et l'analyse de la conversation", dans *Cahiers de linguistique française*, 13 - *Théorie des actes de langage et analyse des conversations*, Genève, Université de Genève, 9-61.
- _____ (1995), "On the Ramification of the Fundamental Notions of Meaning, Analyticity, Consistency, Entailment and Commitment to Speech Acts in Formal Semantics: a reply to Brassac and Trognon", dans *Journal of Pragmatics*, 23 (5), 563-576.
- _____ (1996), "Illocutionary force", dans M. DASCAL et al. (eds), *Sprachphilosophie*,

Philosophy of Language, La philosophie du langage, Tome 2, Walter de Gruyter, Berlin, New York, 1359-1372.

- _____ (1997a), “La logique illocutoire et l’analyse du discours”, dans D. LUZZATI et al. (eds.), *Le dialogique*: Colloque international sur les formes philosophiques, linguistiques, littérales, et cognitive du langage, Bern/Berlin/Frankfurt/New York/Paris/Wien, Peter Lang, 59-94.
- _____ (1997b), “Formal Pragmatics on Non Literal Meaning”, dans *Linguistische Berichte*, vol. 3, 324-341.
- _____ (1999a), “Speech Act Theory and Universal Grammar”, dans *Manuscrito*, Campinas, CLE/UNICAMP, XXII (2), p. 445-468.
- _____ (1999b), “La structure logique des dialogues intelligents”, dans B. MOULIN, et al. (ed.), *Analyse et simulation des conversations*, Limonest, L’Interdisciplinaire, p. 61-100.
- _____ (1999c), “The Basic Logic of Action”, dans *Cahier d’épistémologie*, Montréal, UQAM, n° 9907.
- _____ (1999d), “Success, Satisfaction and Truth in the Logic of Speech Acts and Formal Semantics”, dans *Cahier d’épistémologie*, Montréal, UQAM, n° 9909.
- _____ (1999e), “Illocutionary Logic and Discourse Typology”, dans *Cahier d’épistémologie*, Montréal, UQAM, n° 9912.
- VERSCHUEREN, Jef (1979), “À la recherche d’une pragmatique unifiée”, dans *Communications*, 29, 274-284.
- WILSON, Dreidre et Dan SPERBER (1998), “Pragmatics and Time”, dans R. CARSTON et S. UCHIDA (eds.), *Relevance Theory: Applications and Implications*, Amsterdam, John Benjamins, 1-22.
- WITTGENSTEIN, L. (1961), *Tractatus Logico-Philosophicus*, London, Routledge and Kegan Paul.
- _____ (1968), *Philosophical Investigations*, Oxford, Blackwell.
- ZIV, Yael (1988), “On the Rationality of ‘Relevance’ and the Relevance of ‘Rationality’”, dans *Journal of Pragmatics*, 12, 535-545.