

CENTRE DE RECHERCHES SEMIOLOGIQUES

TRAVAUX DE LOGIQUE

INTRODUCTION AUX LOGIQUES NON CLASSIQUES

D. Miéville, éditeur

CdRS



Université de Neuchâtel

Centre de Recherches Sémiologiques
Travaux de logique
N° 11 - Octobre 1997

INTRODUCTION AUX LOGIQUES NON CLASSIQUES

D. Miéville, éditeur



Université de Neuchâtel

ISSN 1420-8520

SOMMAIRE

Denis MIÉVILLE

Avant-propos v-vi

Frédéric NEF

**Les logiques non classiques sont-elles des
logiques? Dans quelle mesure sont-elles
non classiques?** 1-13

Pierre JORAY

Les logiques modales 15-55

Marek BLASZCZYK

Logique multivalente 57-89

Jean-Pierre DESCLÉS

**La logique combinatoire, types,
preuves et langage naturel** 91-160

Denis MIÉVILLE

La logique développementale 161-187

Helmut LINNEWEBER-LAMMERSKITTEN

Paradoxal simplification 189-221

Bernadette BOUCHON-MEUNIER

**Spécificité et potentialités de
la logique floue** 223-234

Gérard CHAZAL

Logiques non monotones 235-255

Henri VOLKEN

**Le statut logique de la contradiction:
l'approche paraconsistante** 257-271

Giovanni SOMMARUGA-ROSOLEMOS

**Quelques aspects épistémologiques de
la théorie constructive des types** 273-288

AVANT-PROPOS

Grâce à la Conférence Universitaire de la Suisse Occidentale, les Universités de Lausanne et Neuchâtel ont pu organiser pendant l'année académique 1996-1997, un 3^e cycle romand de logique. Le thème de cette formation postgrade était le suivant: *logiques non classiques: limites et enjeux*. Celle-ci s'adressait avant toute chose aux étudiants-chercheurs et aux doctorants intéressés par l'étude de la logique contemporaine.

Si la théorie de la logique des prédicats du premier ordre reste un élément de référence incontournable par rapport auquel on situe d'autres théories logiques, elle a toutefois perdu, aujourd'hui, l'exclusivité de son statut de logique «pure, dure et authentique». Malgré de fameux mandarins, dont Quine n'est pas le moindre, qui n'ont eu de cesse de très clairement expliciter leur répulsion à l'égard de toute autre forme de logique que celle du premier ordre, d'autres théories ont émergé, puis marqué et influencé le développement de la logique contemporaine. Celles-ci peuvent, globalement, être classées en deux familles: il y a d'une part les logiques dites «conservatives» qui, tout en offrant une expansion de la logique classique, en conserve l'esprit et les acquis théoriques; il y a d'autre part les logiques dites malheureusement déviantes et qui s'écartent en fonction de la nature de leurs fondements ou de leurs objectifs, de la logique dite classique.

Dans le cadre de la recherche et de la formation en logique, il était temps d'organiser, au sein des Hautes Écoles de la Suisse occidentale, un cycle de formation postgrade sur ce sujet. Pour ce faire, les professeurs Marie-Jeanne Borel (Section de philosophie) et Henri Volken (Institut de mathématiques appliquées) de l'Université de Lausanne ainsi que Denis Miéville (Séminaire de logique) de l'Université de Neuchâtel ont invité diverses personnalités scientifiques pour animer différents cours, ateliers de recherche et travaux pratiques sur le thème des logiques non classiques. Ces présentations ont rencontré un tel succès qu'il nous a semblé pertinents de les mettre à la disposition d'un plus

grand nombre de chercheurs et d'étudiants avancés, intéressés par la problématique traitée. Grâce à l'amitié et à la générosité des différents auteurs, il est possible de les offrir aujourd'hui dans ce onzième fascicule des Travaux de logique du Centre de Recherches Sémiologiques de l'Université de Neuchâtel. Que chaque auteur soit, ici, encore une fois, chaleureusement remercié.

Denis Miéville
Séminaire de logique
Université de Neuchâtel

LES LOGIQUES NON CLASSIQUES SONT-ELLES DES LOGIQUES? DANS QUELLE MESURE SONT-ELLES NON CLASSIQUES?

Frédéric NEF

Préambule

Dans les pages qui vont suivre, nous présenterons les propriétés qui caractérisent une logique non classique.

1. Qu'est-ce qu'une logique?

On distingue habituellement des critères de démarcation externes et des critères de démarcation internes. Les premiers séparent la logique des autres sciences, alors que les seconds distinguent, à l'intérieur des logiques existantes, celles qui ont droit au titre de logique par excellence. Par exemple, la logique des propositions peut être considérée comme la seule logique.

P. Engel identifie ce problème de la démarcation et le fait de savoir ce qu'est une constante logique. Trois choses sont donc équivalentes:

- donner un critère de démarcation,
- délimiter la classe des constantes logiques,
- fixer la signification des constantes logiques.

Le critère de démarcation externe met en avant des propriétés qui distinguent la logique, ou les vérités logiques, qui sont:

- formelles,
- analytiques,

- non descriptives,
- normatives,
- évidentes,

.....

Dans la définition d'une logique non classique (LNC), deux questions se poseront alors:

- 1) Les constantes logiques des LNC ont-elles une signification correctement fixée?
- 2) Les vérités logiques des LNC sont-elles formelles, analytiques etc.?

2. Qu'est-ce qu'une LNC?

Pour continuer d'être une logique, une LNC doit conserver les propriétés de démarcation externe, mais elle peut modifier, réviser, enrichir les constantes logiques de la logique du premier ordre (LPO), ou en introduire de nouvelles.

Une constante logique est soit un connecteur ($\neg$, $\rightarrow$, $\wedge$, $\vee$), soit un quantificateur ($\exists$, $\forall$). A chaque constante est attachée une sémantique qui détermine les règles d'inférence. Dans une LNC, la sémantique d'une constante logique d'une LPO peut être modifiée.

Exemples: $\neg$ dans la logique intuitionniste;
 $\exists$ dans la quantification meinongienne.

De nouvelles constantes logiques peuvent être introduites:

Exemple: L (pour «nécessaire») en logique modale.

A côté de ces modifications ou innovations touchant les constantes logiques, les LNC peuvent différer de LPO uniquement par la déviation vis-à-vis de principes logiques.

Exemples: $\neg\neg p$ n'est pas équivalent à p en logique intuitionniste;
 $p \vee \neg p$ n'est pas valide en logique para-consistante.

En utilisant un vocabulaire éprouvé, on pourrait appeler *schismatiques* les LNC du premier type (celles qui se contentent d'introduire ou de modifier des constantes logiques) et *hérétiques* les autres (celles qui violent des principes).

Cette division des LNC schismatiques et hérétiques recoupe la classification implicite de l'encyclopédie des LNC qui fait autorité, c'est-à-dire l'ouvrage de D. Gabbay et F. Guenther *Handbook of Philosophical Logic*. En effet, ceux-ci distinguent entre alternatives à LPO (*Handbook*, t. III) et extensions de LPO (*Handbook*, t. II).

On rangera dans les LNC schismatiques les constantes logiques concernées entre parenthèses:

logique modale (L,M), logique temporelle (F,P,G,H), logique intensionnelle ($\wedge$), logique conditionnelle ($\Rightarrow$), logique déontique (O, P), logique érotétique (?)

On rangera dans les LNC hérétiques le principe violé entre parenthèses:

logique multivalente (principe de bivalence), logique paraconsistante (principe de contradiction), logique libre (principe de dénotation), logique non monotone (principe de monotonie de la relation de conséquence logique), ...

Cette division, dont l'introduction met un peu d'ordre dans le maquis ou la jungle des LNC, ne s'applique pas parfaitement pour certaines LNC, comme la logique de la pertinence ou la logique dynamique. Dans ces deux derniers cas, on pourrait discuter de leur distance par rapport à l'orthodoxie des LPO.

Au lieu de distinguer deux niveaux de déviation, on pourrait en distinguer trois, en discernant des déviations à l'intérieur même des deux classes de constantes logiques:

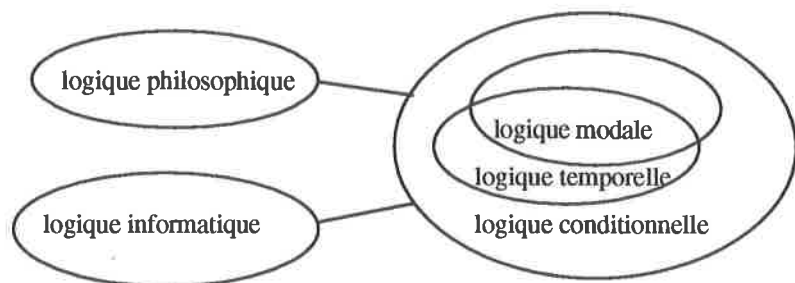
- modification des opérateurs logiques;
- modification de la quantification. Par exemple Quine considère dans la *Philosophie de la logique* les logiques à quantification substitutionnelle comme des LNC;
- violation des principes, comme ceux du tiers exclu ou de la contradiction.

On peut se demander où situer la logique intensionnelle générale dans cette carte des déviations. En effet, cette logique n'introduit pas de constante logique à proprement parler – ' $\wedge$ ' étant en fait plus un artifice notationnel pour marquer le sens d'une expression (exemple: ' $\wedge p$ ' désigne le sens de p) et elle ne viole pas de principe logique à proprement parler – en général, elle est bivalente, respecte le principe de contradiction, etc. Est-ce que le non-respect du principe d'extensionnalité la rejette du côté de l'hérésie logique? Rien n'est moins sûr. En tout cas, toutes les LNC ne sont pas intensionnelles (exemple: la logique multivalente, la logique non monotone, ...) et la logique intensionnelle n'est donc pas de manière certaine une LNC au sens le plus fort que nous venons de définir.

3. Architecture des LNC

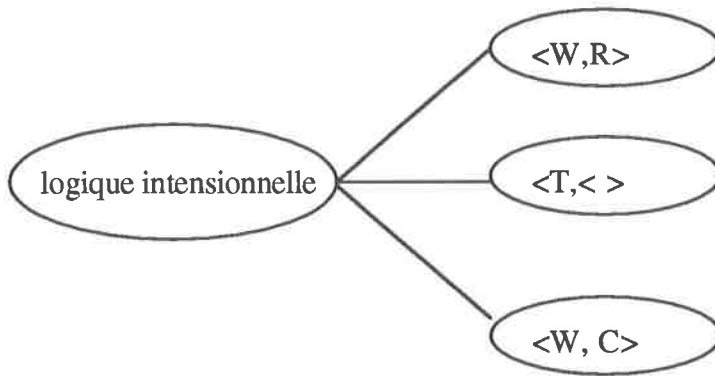
Pour un novice, la prolifération des LNC apparaît comme dénuée de plan. On peut cependant décrire l'architecture des LNC de deux façons au moins.

A. Situation de fait:



Logique philosophique et logique informatique sont séparées; dans les LNC, c'est la logique modale qui occupe un rôle central dans les applications, suivie par la logique temporelle et la logique conditionnelle (tout ceci étant très simplifié).

B. Par rapport à une relation interne de structuration du contenu théorique:



On a indiqué les cadres des trois logiques mentionnées ci-dessus:

- $\langle W, R \rangle$ pour la logique modale, W étant l'ensemble des mondes possibles et R la relation d'accessibilité entre mondes.
- $\langle T, < \rangle$ pour la logique temporelle, T étant l'ensemble des instants, et $<$ la relation d'antériorité sur les instants.
- $\langle W, C \rangle$ pour la logique conditionnelle, W étant comme ci-dessus et C une relation de condition.

Dans cette optique, on peut apercevoir une similarité de structure des cadres qu'ils soient modal, temporel ou conditionnel. C'est cette parenté de structure que va exploiter la théorie de la correspondance de Van Benthem.

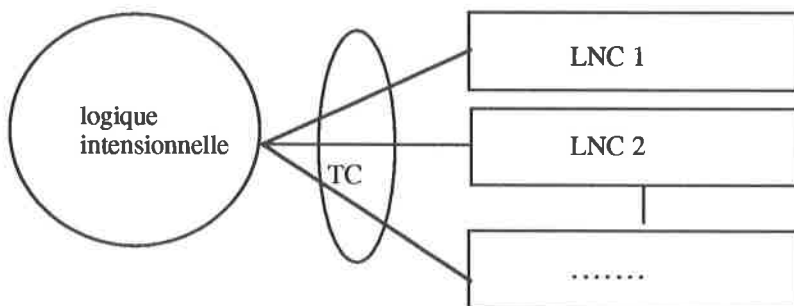
On a indiqué de plus des relations entre une logique intensionnelle générale et ces logiques intensionnelles spéciales.

4. Théorie de la correspondance de Van Benthem

La théorie de la correspondance (TC) a pour fonction d'explorer les correspondances formelles entre les différentes LNC schématiques. L'enjeu est grand par rapport à la question

posée au départ. Si les différentes LNC se laissent ordonner par la TC, alors nous devons poser la question à partir de cette mise en ordre. Epistémologiquement la situation est similaire à celle de l'axiomatisation des géométries. C'est à partir de celle-ci, qui dévoile les connexions entre les différentes géométries comme des structures abstraites, que l'on peut poser la question de la nature géométrique de tel ou tel formalisme – en dehors de cette démarche, la question n'a tout simplement pas de sens.

La TC a donc pour fonction de donner de la substance à la logique intensionnelle en opérant une mise en ordre sur les LNC:



TC agit à un niveau syntaxique en montrant des parallélismes dans l'usage des constantes logiques (exemple: le parallélisme de 'R' et de '<') et dans la quantification. On peut en effet concevoir par exemple la logique modale comme une quantification sur les mondes possibles. Elle agit également à un niveau sémantique en établissant des parallélismes entre les modèles d'interprétation.

La correspondance syntaxique établit donc une relation entre M et $\exists$, L et $\forall$, F, P et $\exists$, G et H et $\forall$, et donc des relations entre F,P et M et G,H et L:

M	L
$\exists$	$\forall$
F,P	G,H

N.B. On rappelle que M = possible que, L = nécessaire que, F = il sera le cas au moins une fois, P = il a été le cas au

moins une fois, $G =$ il sera toujours le cas que, $H =$ il a toujours été le cas que.

Par exemple, la correspondance sémantique de la logique modale (LM) et de la logique temporelle (LT) a l'allure suivante:

Soit le modèle pour LT: $\langle T, <, V \rangle$.

Soit le modèle pour LM: $\langle W, R, V \rangle$.

V est la fonction d'évaluation.

Correspondances: T et W , $<$ et R , et également entre cadres $\langle T, < \rangle$ et $\langle W, R \rangle$.

Cette correspondance sémantique fonde les ressemblances des propriétés de LT et LM, en permettant de comparer les propriétés de $<$ et de R , afin de distinguer des classes de logiques.

Exemple: la transitivité de $<$ caractérise une classe de logique temporelle, tout comme la transitivité de R caractérise un calcul modal (S4).

D'autres correspondances peuvent être mises en lumière. Citons-en deux. Tout d'abord, à un niveau axiomatique, on peut mettre en correspondance les systèmes minimaux de LM et de LT:

Exemples (on note # la correspondance):

$Fp \leftrightarrow \neg G\neg p \# Mp \leftrightarrow \neg L\neg p$
(interdéfinition des opérateurs).

$G(p \rightarrow q) \rightarrow (Gp \rightarrow Gq) \# (L(p \rightarrow q)) \rightarrow (Lp \rightarrow Lq)$
(distribution des opérateurs).

$p/Gp \# p/Lp$
(absolutisation).

Ensuite pour certaines formules, comme celles de Barcan:

$F \forall x Ax \rightarrow \forall x F Ax \# L \forall x Ax \rightarrow \forall x L Ax,$
 $\forall x F Ax \rightarrow F \forall x Ax \# \forall x L Ax \rightarrow L \forall x Ax.$

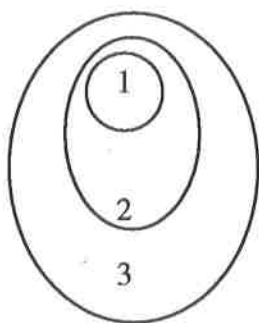
La TC peut être appliquée aussi pour comparer LT et LM à la logique conditionnelle. Cependant la relation C , définie ainsi:

$C(x, y, z) =_{\text{def}} y$ est plus semblable à x que z ne l'est,

est une relation ternaire, à la différence de $<$ et R ; cette propriété est irréflexive et transitive.

La TC permet donc de commencer à répondre à la question posée initialement. Si une LNC est une extension conservatrice de la logique classique, alors il existe une correspondance avec les autres LNC. La structure commune qui assure cette correspondance est un noyau quantificationnel. Les LNC qui sont des extensions conservatrices ont en commun la doctrine de la quantification, qui précisément assure la correspondance et sur la base de cette structure commune quantifie, mais selon les règles habituelles, sur des mondes, des instants, ce qui installe la quantification dans des contextes opaques, au sens de Quine, où le principe d'extensionnalité ne vaut plus. La réponse est donc: elles sont des logiques, parce qu'elles sont des extensions conservatrices des appareils quantificationnels, et elles sont non classiques, car elles sont intensionnelles. C'est dans ce sens que les logiques, dont nous avons parlé (LT, LM et la logique conditionnelle), sont conservatrices: elles préservent la bivalence, le principe de contradiction, les règles habituelles de la quantification et de la déduction dans des contextes qui n'autorisent pas la substitution des identiques.

En ce qui concerne la démarcation interne, la TC montre-t-elle que la logique modale est le noyau des LNC conservatrices? En effet, on peut hésiter entre deux représentations des relations entre les LNC conservatrice, hiérarchique et modulaire:



Hiérarchique

1 2 3

Modulaire

Il existe des arguments pour montrer que LT n'est qu'un décalque de LM (ce qui a motivé d'ailleurs l'abandon d'une LT fondée sur des instants, stricts analogues des mondes possibles, au profit d'une logique des événements, plus émancipée, mais que l'on intègre dans l'ensemble par des théorèmes de correspondance sur l'équivalence des structures d'instant et d'événements.

5. Quelques tendances en LNC conservatrice

On peut mentionner rapidement trois grandes tendances:

1) Une évolution de la théorie des modèles pour LM en termes de mondes possibles vers une sémantique des modèles partiels (situations: Barwise et Perry, représentations discursives: Kamp, etc.).

2) Un passage d'une ontologie pour LT en termes d'instant et d'intervalles à une logique des événements (Kamp). Il existe plusieurs types de motivations pour ce passage:

Motivation d'ordre sémantique: La sémantique du langage naturel réclame que l'on prenne en compte des intervalles, dans la mesure où l'idéalisation consistant à se limiter à des structures ponctuelles d'instant est trop forte. L'introduction des événements est justifiée par la description des aspects et des temps des langues naturelles.

Motivation d'ordre ontologique: Il est plus raisonnable de considérer les structures ponctuelles comme des limites des structures d'intervalles et, dans une ontologie qui admet des événements comme primitifs, de les considérer comme étendus sur des intervalles.

Motivation d'ordre psychologique: Beaucoup de faits militent pour une dualité du temps privé constitué d'intervalles (la durée bergsonienne) et d'un temps public constitué à partir de structures d'instant.

Les structures d'instant sont des cadres temporels composés de deux éléments:

< T, >

un ensemble d'instants T et une relation temporelle $<$.

Les structures d'intervalles sont des triplets composés d'ensembles d'intervalles I et de deux relations $<$ et $\subseteq$:

$$\langle I, <, \subseteq \rangle.$$

Les structures d'événements sont des triplets composés d'ensembles d'événements, E , et de deux relations, O (recouvrement) et $<$:

$$\langle E, <, O \rangle$$

Cette logique des événements provient de Russell.

3) Une modification de la théorie syntaxique de la logique conditionnelle vers une théorie de la révision (Gärdenfors).

On peut comprendre cette modification à partir de deux interprétations du test de Ramsey.

D'après Ramsey, $A \Rightarrow B$ (où $\Rightarrow$ est la relation conditionnelle) est vrai à un ensemble de croyances X si $A \Rightarrow B$ passe le test suivant:

Ajouter A à X . Si le résultat devient inconsistent, réaliser la révision minimale pour que A n'introduise pas d'inconsistance. Vérifier alors que B découle de X .

Ce test de Ramsey devient dans la version syntaxique de la logique conditionnelle:

Un ensemble X valide $A \Rightarrow B$, si pour chaque ensemble Γ maximal relativement à la propriété « $X \vdash \Gamma$ et $\Gamma + A$ est consistant», $\Gamma + A \vdash B$.

Dans la version de Gärdenfors (1985) ce test devient:

Accepter une phrase de la forme $A \Rightarrow B$ dans un état de croyances K , ssi le changement minimal de K nécessaire pour accepter A exige aussi d'accepter B .

A noter qu'on peut combiner le test de Ramsey et la théorie bayésienne de la croyance conditionnelle rationnelle. Par exemple, Stalnaker a montré que:

$$P(A \Rightarrow B) = P(A | B)$$

où $P(A \Rightarrow B) =_{\text{def}}$ degré de croyance de l'agent dans $A \Rightarrow B$ et
 $P(A | B) =_{\text{def}}$ degré de croyance conditionnelle en B étant donné A .

Ce qui est commun à (1), (2) et (3), c'est la mise en avant d'aspects dynamiques et partiels qui agissent sur la logique épistémique et la sémantique du discours et s'intègrent parfaitement dans une conception cognitive de la LNC. On entend par là une conception où la LNC sert de manière privilégiée à modéliser le raisonnement et les croyances rationnelles.

6. Conclusion

Comme on l'a vu, le domaine de la LNC peut être abordé de plusieurs manières: philosophique, informatique, sémantique et linguistique, cognitive. De chacun de ces points de vue, on privilégiera des aspects de LNC:

- Philosophiquement, la LNC est une partie de ce que l'on nomme «logique philosophique», avec quoi la «philosophie de la logique» n'a pas de frontière nettement tranchée. De ce point de vue, on privilégie les motivations ontologiques.

Exemple: les structures d'événements dont on préfère dériver les structures d'instant, plutôt que l'inverse.

- Informatiquement, surtout avec LT et LM, la LNC fournit des modèles de logique de la computation et on privilégie de ce point de vue les aspects calculatoires et syntaxiques.

- Sémantiquement, la LNC fournit des modèles pour une interprétation de phénomènes de la langue naturelle. De ce point de vue, et contrairement à ce qui précède, on peut privilégier la logique comme langage plutôt que la logique comme calcul. Associée à un λ -calcul et à une grammaire catégorielle, la logique intensionnelle est un bon outil pour la sémantique des langues naturelles.

- Cognitivement, la LNC peut être envisagée comme une modélisation de cartes cognitives sous-jacentes aux comportements intelligents (inférences, croyance, perception,...). Dans cette

perspective, on sera tenté de chercher une vue de la LNC où soient mis en avant les aspects dynamiques.

Université de Rennes 1
Institut de philosophie
 261, Av. du Général Leclerc
 F 35042 RENNES

7. Bibliographie sélective

N.B. Une bibliographie couvrant ce domaine et les questions évoquées ci-dessus excéderait de loin la taille d'un article. Les titres qui suivent ont été choisis pour leur capacité à fournir un approfondissement de la réflexion fondationnelle. On ne mentionne pas les grands classiques (Peirce, Frege, Russell, Prior,...).

BARWISE J. (1989). *The Situation in Logic*. Stanford: CSLI.

COOPER R., MUKAI K. & PERRY J. (eds) (1990-1993). *Situation Theory and its Applications*. Stanford: CSLI, 3 vols.

CRESSWELL M.J. (1990). *Entities and Indices*. Dordrecht: Kluwer.

CRESSWELL M.J. (1994). *Language and the World. A Philosophical Inquiry*. Cambridge: Cambridge U.P.

DUBUCS J. & LEPAGE F. (éds) (1995). *Méthodes logiques pour les sciences cognitives*. Paris: Hermès.

ENGEL P. (1989). *La norme du vrai: philosophie de la logique*. Paris: Gallimard.

GABBAY D. & GUENTHNER F. (eds) (1983-89). *Handbook of Philosophical Logic*. Dordrecht: Reidel, 4 vols.

GALIN (1975). *Intensional Logic and Higher-order Logic*. Los Angeles.

GÄRDENFÖRS P. (1985). *Knowledge in Flux: Modeling the Dynamics of Epistemic States*. Cambridge, London: The MIT Press.

HAAS S. (1996). *Deviant Logic, Fuzzy Logic: Beyond the Formalism*. Chicago: Chicago U. P.

- MARCUS R. (1995). *Modalities: Philosophical Essays*. Oxford: Oxford U.P.
- MONTAGUE R. (1973). *Formal Philosophy*. New Haven, London: Yale U.P.
- NEF F. (1988). *Logique et langage. Essais de sémantique intensionnelle*. Paris: Hermès.
- THAYSE A. et. al. (éds) (1988-1991). *Approche logique de l'intelligence artificielle*. Paris: Dunod, 4 vols.
- VAN BENTHEM J. (1982). *Time Logic*. Dordrecht: Reidel.
- VAN BENTHEM J. (1982). *Modal Logic and Classical Logic*. Napoli: Bibliopolis.
- VAN BENTHEM J. (1984). Correspondence Theory. In: Gabbay D. & Guenther F. (eds) 1984, vol. 2, 167-249.
- VAN BENTHEM J. (1986). *Essays in Logical Semantics*. Dordrecht: Reidel.
- VAN BENTHEM J. (1988). *A Manual of Intensional Logic*. Stanford: CSLI, 2e ed.
- VAN BENTHEM J. (1995). *Language in Action. Categories, Lambdas and Dynamic Logic*. Amsterdam: North Holland.
- ZALTA E. (1988). *Intensional Logic and the Metaphysics of Intentionality*. Cambridge, London: The MIT Press.

LES LOGIQUES MODALES

Pierre JORAY

Ainsi pourrait-on définir simplement le sens du possible comme la faculté de penser tout ce qui pourrait être "aussi bien", et de ne pas accorder plus d'importance à ce qui est qu'à ce qui n'est pas. On voit que les conséquences de cette disposition créatrice peuvent être remarquables; malheureusement, il n'est pas rare qu'elles fassent apparaître faux ce que les hommes admirent et licite ce qu'ils interdisent, ou indifférents l'un et l'autre...

R. Musil, *L'homme sans qualités*.

Introduction

Les raisons des modifications que l'on peut apporter à la logique classique sont de plusieurs ordres. Si les motivations liées aux développements de logiques rivales sont généralement associées au rejet de tel ou tel principe fondateur de la logique classique, le projet de construire une logique des modalités s'appuie sur un constat plus ordinaire: la logique classique reste inadaptée pour rendre compte de certaines inférences du raisonnement non formel, inférences qui portent cependant une validité qui nous apparaît intuitivement incontestable.

On admet généralement que, parmi la diversité de ces raisonnements, certains possèdent un caractère privilégié au regard de la logique conçue comme science de l'inférence. Autrement dit, certains s'avèrent s'articuler sur des notions dont une telle science ne peut envisager de se passer. Ainsi, au-delà de la volonté générale de couvrir le domaine le plus large des infé-

rences jugées valides, le projet d'intégrer les modalités dans la logique relève de visées plus spécifiques. Il s'agit principalement de tenter de donner à l'idée de *nécessité logique* un contour formel précis.

Sans vouloir entrer dans la problématique du sens qu'il faut attribuer à l'énoncé d'un locuteur affirmant qu'un objet possède une propriété de manière *nécessaire* ou seulement *contingente*, on constate simplement que certaines entités de la logique semblent intimement liées aux modalités dites *classiques*¹. Il s'agit principalement du cas des tautologies, dont on dit non seulement qu'elles sont vraies, mais qu'elles possèdent ce caractère de vérité d'une manière nécessaire. Parallèlement, les propositions inconsistantes relèvent de l'impossibilité.

En outre, plusieurs des concepts centraux de la logique comme science de la déduction ne peuvent se définir sans un appel explicite à des notions modales: je pense en particulier au concept de *validité* associé aux arguments. On dit qu'un argument est valide si et seulement si il est *logiquement impossible* que toutes ses prémisses soient vraies et sa conclusion fausse. Je relève encore la notion d'*inconsistance* d'un ensemble de propositions, qu'on définit comme l'*impossibilité logique* que la conjonction de toutes les propositions de cet ensemble soit vraie.

Nécessité et *impossibilité* apparaissent ici, dans le cadre classique, comme des notions clairement métalinguistiques; on est en présence de deux niveaux distincts car ce que disent ces modalités se rapporte à des éléments d'un langage dont elles ne font pas partie – en l'occurrence les propositions.

Cependant, le besoin peut se faire sentir d'amener ces notions hors du domaine exclusif de la pratique intuitive et d'en fixer l'usage en y associant des règles formelles précises. Et bien que parler de propositions modales semble relever d'une confusion de niveaux linguistiques, la construction d'une logique modale passe précisément par l'idée de ramener les modalités dans le langage objet de la logique, faisant de ces modificateurs de propositions de véritables opérateurs. L'idée est d'aug-

1 Ces modalités, que l'on nomme aussi *aléthiques* ou parfois même *ontiques*, sont principalement la nécessité, la possibilité, la contingence.

menter le vocabulaire de la logique classique en y adjoignant, à côté de la négation $\sim$, de nouveaux symboles d'opérateurs unaires: L et M, respectivement les symboles de la nécessité et de la possibilité logiques.

Cela dit, la construction d'un langage logique à partir d'un vocabulaire élargi relève d'une liberté presque totale. Mais si le but poursuivi par une telle entreprise est de circonscrire formellement la perception intuitive des modalités qui ont été évoquées, il paraît raisonnable de fixer certaines conditions élémentaires:

1. Une logique modale sera une *expansion conservative* de la logique classique C. Elle devra, autrement dit, contenir C mais ne pas ajouter à celle-ci des théorèmes dans lesquels ne figure-raient aucun des nouveaux symboles.

2. Elle devra disposer d'une correspondance entre les deux nouveaux opérateurs, telle que l'un puisse s'exprimer à l'aide de l'autre. Dans la pratique, on ajoutera au vocabulaire un seul des symboles – généralement L –, le second étant défini ensuite de la manière suivante:

$$MP = \text{df } \sim L \sim P$$

3. L'expression suivante y sera une thèse:

$$\vdash LP \supset P$$

4. Une logique modale devra satisfaire au *principe de nécessité* suivant: toute thèse du système – et en particulier, selon la condition 1, toute thèse du calcul classique – sera encore une thèse lorsqu'elle est précédée du symbole L de nécessité, offrant ainsi la garantie que toute tautologie sera nécessaire.

$$\text{Si } \vdash P \text{ alors } \vdash LP$$

Ces conditions générales laissent cependant encore la voie ouverte à de nombreuses réalisations axiomatiques ou des systèmes de règles d'inférences forts différents. Je me contenterai dans les pages qui suivent d'en présenter un petit nombre parmi les plus significatifs, en particulier les systèmes connus sous les étiquettes de T, S4 et S5.

D'un point de vue historique cependant, le cheminement parcouru jusqu'à l'élaboration de ceux-ci est passablement différent de celui que j'ai rapidement esquissé ici. Si, comme je l'évoquerai dans la suite, on rencontre dès les premières systématisations de la logique dans l'Antiquité des tentatives de traitement des modalités, c'est principalement à C. I. Lewis que l'on doit leur entrée dans le champ de la logique moderne. L'idée développée par Lewis fut de proposer comme alternative à l'implication matérielle des *Principia Mathematica* de Whitehead et Russell un concept d'implication dit *implication stricte*, qui inclut explicitement une notion modale et dont le but était d'éviter ce qu'on a nommé les «paradoxes de l'implication matérielle».

Certes les développements les plus récents, attachés principalement à une approche sémantique des modalités, ont pratiquement rendu à la logique modale son autonomie par rapport à cette problématique de l'implication². Il m'a cependant paru opportun de lui donner dans mon propos la place qui lui revient lorsqu'on privilégie, comme je le ferai ici, une approche à dominante syntaxique.

Par ailleurs, on ne saurait parler de logique modale sans mentionner la grande famille des logiques dites intensionnelles. Leur apparentement aux logiques modales se fonde sur une manière analogue de traiter d'autres opérateurs de type modal. On trouve ainsi des logiques déontiques, des logiques épistémiques, des logiques temporelles ou encore des logiques de la causalité. Toutes ces tentatives de formalisation constituent des outils importants tant dans le domaine de la logique philosophique que pour l'étude des arguments du discours ordinaire³. Je ne ferai cependant ici que les nommer, la présentation qui suit étant exclusivement consacrée aux modalités classiques.

Bien entendu, de par sa brièveté, le présent travail ne peut prétendre proposer une vue large sur une discipline qui reste aujourd'hui en plein essor. Il s'agit ici plus modestement de mettre l'accent sur certains aspects fondamentaux propres à éclairer les problèmes syntaxiques ainsi que la démarche

2 Les principaux développements issus spécifiquement de la problématique de Lewis sont connus sous le nom de logique de la pertinence (ou *relevant logic*). Cf. Read 1988.

3 Cf. Van Benthem 1988, Zalta 1988 et Gardies 1979.

sémantique, désormais dominante. Pour le reste, je me contenterai de renvoyer aux traités que l'on trouve dans une abondante littérature⁴.

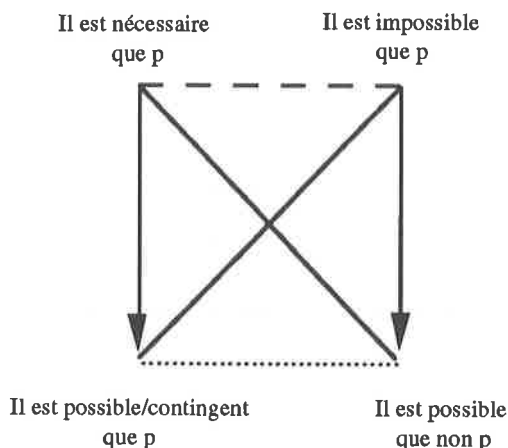
1. Les modalités dans la logique traditionnelle

Comme souvent en logique, le poids de la tradition aristotélicienne paraît tout à fait déterminant. Non seulement les travaux d'Aristote constituent la première attestation d'une logique des modalités⁵, mais on peut faire remonter au Stagiritte la prédominance accordée aux modalités dites classiques, ainsi que dans une large mesure l'usage dominant des termes *nécessaire*, *impossible*, *possible* et *contingent*; c'est ce dernier point que j'aborderai principalement ici.

L'hésitation d'Aristote entre plusieurs façons d'organiser les modalités est bien connue. Mais la concurrence d'acceptions différentes pour certaines modalités semble être une conséquence de sa philosophie essentialiste. La distinction entre deux sortes de *puissance* entraîne en effet deux interprétations des notions de *possibilité* et de *contingence*. D'une part, lorsqu'il s'agit de la *puissance* rapportée aux êtres éternels, loin de s'opposer au caractère constamment nécessaire des propriétés essentielles, la possibilité y est impliquée par la nécessité. Elle est ainsi comprise comme la négation contradictoire de l'impossibilité. Cette première acception, dite *unilatérale*, conduit Aristote à une organisation des modalités que l'on représente par un carré des oppositions tout à fait analogue à celui des propositions catégoriques:

4 On retiendra particulièrement Konyndyk 1986, Chellas 1980, ainsi que l'excellent Hughes & Cresswell 1996.

5 La théorie des propositions modales est développée dans le *De Interpretatione* et les *Premiers Analytiques*, ce dernier incluant une théorie des syllogismes modaux.



Jouant sur les deux manières d'appliquer la négation aux propositions modales – soit sur la modalité elle-même, soit sur son *dictum* –, Aristote montre que toutes les modalités peuvent y être définies sur la base d'une seule, par exemple la nécessité⁶. Bien que cette figure reprenne toutes les oppositions du traditionnel carré et présente une grande cohérence, on y remarque un certain manque d'équilibre dû à un usage constamment concurrent de la possibilité et de la contingence⁷. Et surtout, il manque encore à Aristote un terme simple pour exprimer la négation contradictoire de la nécessité.

Cependant Aristote accorde dans sa syllogistique une grande importance à une autre sorte de *puissance* que l'on trouve dans sa métaphysique. Il s'agit de la *puissance* que l'on rapporte aux êtres qui connaissent des propriétés accidentelles et pour lesquels la possibilité ou la contingence s'opposent d'un bloc au caractère nécessaire de la présence ou de l'absence de propriétés essentielles. L'opposition centrale entre des propositions *apodictiques* (nécessaires ou impossibles) et des propositions *problématiques* (possibles/contingentes) amène Aristote à proposer dans son étude des syllogismes modaux une organisation

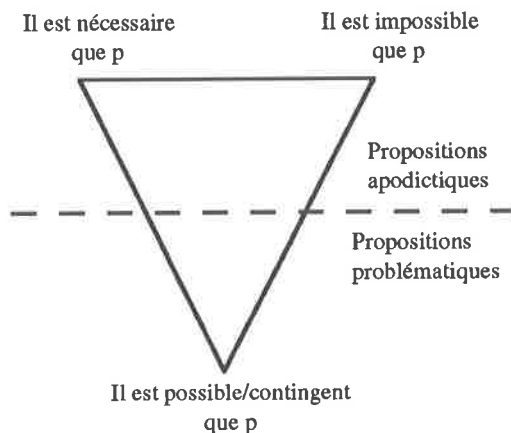
6 On y trouve en particulier les fameuses équivalences:

Il est possible que p $\leftrightarrow$ Il n'est pas nécessaire que non p

Il est impossible que p $\leftrightarrow$ Il est nécessaire que non p.

7 Quelle qu'en soit l'interprétation, Aristote semble pratiquement toujours utiliser possibilité et contingence comme des synonymes.

concurrente des modalités qui, cette fois, peut être présentée sous la figure d'un triangle dont les sommets sont des contraires pris deux à deux:

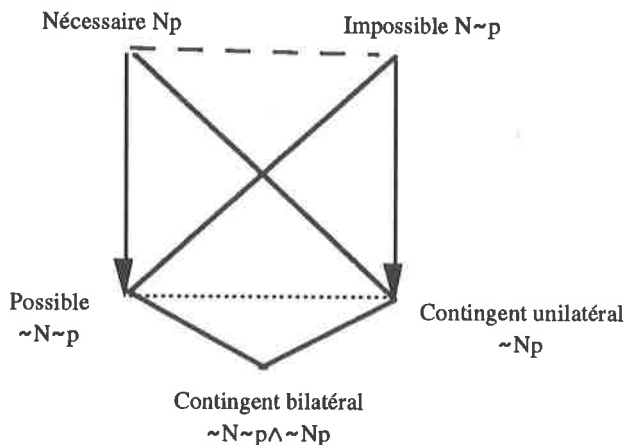


Mais, choisissant d'opposer le possible aux deux autres modalités, Aristote fait *de facto* reposer l'ensemble sur des bases qui manquent d'homogénéité; ce que le jeu de la négation va une fois de plus mettre en évidence. La modalité caractéristique des propositions problématiques – nommée possible ou contingent *bilatéral* – s'avère en effet être une modalité composée. Comme le remarquera Aristote lui-même, nier le *dictum* d'une proposition problématique c'est revenir à cette même proposition⁸, alors que la négation de la modalité conduit fâcheusement à une disjonction. La négation du possible/contingent c'est ce qui est nécessaire *ou* impossible. Le possible/contingent est bien une conjonction de modalités: ce qui à la fois n'est pas nécessaire *et* n'est pas impossible.

Bien qu'en un sens la figure triangulaire soit plus conforme à la distinction entre propriétés essentielles et accidentelles, ses conséquences logiques négatives, en particulier l'absence de la relation fondamentale de contradiction, ont permis à la conception bilatérale de prendre historiquement le dessus. Aujourd'hui encore, c'est cette version d'un possible impliqué par le neces-

8 En effet, il est contingent que p *ssi* il est contingent que non p.

saire qui prévaut. Seul persiste un flottement dans l'usage du terme contingent. Tantôt il désigne la simple négation du nécessaire, tantôt la modalité bilatérale d'Aristote. Si bien que l'organisation qui nous est restée de la tradition peut être représentée par un schéma reprenant toutes les modalités mises en évidence par Aristote:



Cette organisation traditionnelle des modalités, et en particulier la considération des propositions apodictiques et problématiques, reste cependant le produit direct du système aristotélicien. Or le projet même d'une syllogistique modale, comme le relèvent Kneale et Kneale, semble manquer de légitimité. Car il tombe dans un dilemme relatif aux deux manières d'attribuer la modalité: la manière externe ou *de dicto*, que l'on trouve dans la version du carré et la manière interne ou *de re*, plus conforme à la philosophie essentialiste, et qui prévaut dans l'organisation triangulaire des *Premiers analytiques*. Le dilemme se pose en ces termes:

If modal words modify predicates [interprétation *de re*], there is no need for a special theory of *modal* syllogisms. For these are only ordinary assertoric syllogisms of which the premisses have peculiar predicates. On the other hand, if modal words modify the whole statements to which they are attached [interprétation *de dicto*], there is no need for a special modal *syllogistic*, since the rules determining the logical

relations between modal statements are independant of the character of the propositions governed by the modal words (Kneale and Kneale 1962: 91).

Le reproche fait ainsi aux modalités *de re* relève de leur caractère éminemment intensionnel, peu conforme au projet d'une syllogistique formelle, dont un des buts est précisément de faire abstraction des contenus. Les modalités *de dicto* semblent, elles, entachées d'une confusion entre le niveau de l'assertion modale et celui du *dictum*, objet de cette assertion. Selon Kneale et Kneale, ce qui manque au projet d'Aristote pour échapper à ce problème, c'est précisément la base sur laquelle les mégaro-stoïciens vont construire leur logique modale: une logique des propositions inanalysées.

Comme on le sait, les mégariques et les stoïciens rejetaient très clairement l'essentialisme d'Aristote. Ils ne pouvaient ainsi s'accommoder de modalités fondées sur une distinction qu'ils refusaient entre des propriétés essentielles et des propriétés accidentelles. La plus grande originalité dans leur tentative d'intégrer les notions modales à leur logique est sans doute d'en avoir donné une interprétation temporelle⁹.

Diodore Cronos donnait ainsi les définitions suivantes, qui sont restées parmi les plus célèbres:

Le *possible* est ce qui est vrai ou sera vrai,
 l'*impossible* ce qui étant faux ne sera pas vrai,
 le *nécessaire* ce qui étant vrai ne sera pas faux, et
 le *non nécessaire* ce qui est faux ou sera faux.

Il ressort de ces définitions que les modalités ne portent pas sur les états de choses décrits par les énoncés, mais sur les énoncés eux-mêmes. On ne saurait cependant dire s'il s'agit de véritables propositions du fait de la présence d'indices temporels¹⁰. Il faut ici se souvenir que Diodore admettait dans sa logique que les énoncés puissent changer de valeur de vérité. Le fameux «Il fait jour» est vrai à tel instant, alors qu'il est faux quelques

9 On trouvera une description des sources, ainsi qu'un certain nombre de textes en traduction anglaise dans Mates 1961.

10 Sur les manières de représenter les indices temporels des propositions des stoïciens, Cf. Mignucci 1978: 336-337.

heures plus tard. Les modalités apparaissent ainsi comme des sortes de quantifications portant sur les indices temporels, le nécessaire et le possible jouant alors le rôle respectivement d'universelle et de particulière temporelles. Et, de même qu'ils peuvent changer de valeur de vérité, les énoncés peuvent aussi changer de modalités, ce qui montre que Diodore entend par nécessité ou possibilité non pas des termes absolus mais des termes relatifs à un instant. En outre, les changements de modalités ne se font pas sans règles: les énoncés possibles ou non nécessaires pourront passer à toute autre modalité, mais une fois devenu nécessaire ou impossible aucun énoncé ne pourra plus changer, ni sa modalité, ni même sa valeur de vérité.

Une telle réduction des modalités au temporel comporte inévitablement certains défauts. En bien des points c'est notre intuition des modalités qui est mise à mal. Par exemple, comment accepter que tout ce qui est possible est ou sera vrai au moins une fois. Ne doit-on pas admettre qu'il est possible que je sois un assassin, même dans le cas où jamais je ne tuerai personne, et qu'il nous faut pour exprimer cela une modalité spécifique? Mais il ne faudrait pas pour autant négliger les qualités de l'entreprise, qui ne vont pas, pour certaines, sans se rapprocher de préoccupations modernes. Et comme le relève J.-L. Gardies:

L'originalité de l'analyse mégaro-stoïcienne des modalités ontiques tient à ce qu'elle réussit à dépouiller les représentations ordinaires de ces modalités de ce qu'elles pouvaient apparemment comporter de rigoureusement intensionnel. Dans une telle transcription, la modalité s'analyse, au même titre que la simple universalité, en strictes valeurs de vérité et de fausseté (1979: 39).

En effet, pour montrer par exemple qu'une proposition comme «Il est possible que je sois à Corinthe» est vraie, les mégaro-stoïciens disposent d'une méthode «quasi extensionnelle»: exhiber un moment de la chaîne du temps où la proposition «Je suis à Corinthe» est vraie. D'une telle possibilité, qui ne va pas sans rappeler les évaluations des sémantiques modernes, la logique aristotélicienne n'en dispose pas; elle doit, dans un cas comme celui-ci, inévitablement passer par une analyse du

contenu de la proposition¹¹. Mais voyons maintenant par quel biais la logique moderne est revenue à la question des modalités.

2. C. I. Lewis et le problème de l'implication

Si la logique peut être conçue comme un langage capable de mettre en évidence l'ensemble des tautologies ou des formules valides d'un calcul vérifonctionnel, elle peut aussi l'être d'une manière plus large comme science de l'inférence. Pris en ce second sens, l'objet central de son étude est la notion de validité des arguments, et plus précisément elle a pour but la mise en évidence des conditions sous lesquelles il est possible d'inférer une conclusion à partir d'une prémisses ou d'un ensemble de prémisses. La relation qui existe alors entre la (les) prémisses(s) et la conclusion est la relation que l'on nomme *relation d'implication*. Cette relation, à partir de laquelle on définit une relation d'équivalence propositionnelle, joue un rôle tout à fait central dans tout système d'inférence, car elle permet d'organiser l'ensemble de ses expressions selon des classes d'équivalence ordonnées¹². Dans la logique classique, on définit comme suit la relation $\rightarrow$ dite d'*implication vérifonctionnelle*, à partir de l'opérateur $\supset$ de conditionnelle:

$$P \rightarrow Q =_{df} \vdash P \supset Q$$

Autrement dit, *P implique vérifonctionnellement Q* si et seulement si $P \supset Q$ est une *tautologie* ou, ce qui revient ici au même, un *théorème* du système.

Le problème soulevé par Lewis consiste à se demander si une telle relation est adéquate pour recouvrir l'usage ordinaire de l'implication dans la pratique déductive¹³. Apporter une réponse à cette question relève inévitablement de la gageure dans la

11 Sur les modalités chez Aristote, Kneale & Kneale, 1962: 81-96, Blanché 1970: 67-77, Gardies 1979: 11-28 et pour une étude plus précise, Patterson 1995. Concernant les stoïciens, cf. Mates 1961, Mignucci 1978, Kneale & Kneale 1962: 117-128.

12 Cf. Grize 1971: 13-18.

13 La présentation qui suit s'éloigne passablement de celle de Lewis, le but poursuivi étant de rendre plus explicite la distinction des différents niveaux linguistiques. On trouvera un exposé de Lewis lui-même dans Lewis & Langford 1959, en particulier, chap. VIII.

mesure où une définition générale de la notion d'implication fait défaut. Cependant, on peut reconnaître pour toute relation d'implication certaines conditions minimales.

D'un point de vue sémantique, une telle relation doit préserver la vérité dans le passage des prémisses à la conclusion. Elle doit ainsi empêcher tout passage du vrai au faux. Cette condition découle tout naturellement de la définition de la validité des arguments¹⁴. De part sa définition, $\rightarrow$ satisfait à cette condition. On peut aisément le vérifier à l'aide de la table de vérité de la conditionnelle; exiger que $P \supset Q$ soit une tautologie, c'est précisément exclure le cas où P est vraie et Q fausse:

P	Q	$P \supset Q$
1	1	1
1	0	0
0	1	1
0	0	1

D'un point de vue syntaxique, si P implique Q , il faut que Q soit *déductible* de P . Cette condition est elle aussi satisfaite par $\rightarrow$. Car (si on prend le cas de la déduction naturelle), montrer que $P \supset Q$ est un théorème consiste justement, comme l'exige la règle d'introduction de la conditionnelle, à exhiber une sous-déduction dans laquelle Q est déduite de P .

m		<u>P</u>	
.		.	
.		.	
.		.	
n		Q	
		$P \supset Q$	m-n, $\supset$ intro

Cela dit, ces conditions restent minimales. Notons que la relation d'implication vérifonctionnelle $\rightarrow$ s'appuie sur la notion de *tautologie*. Cela signifie que ne sont pas prises en compte les valeurs effectives des propositions en jeu, mais seulement l'ensemble des valeurs qu'elles sont susceptibles de prendre. Dans la pratique déductive cependant, cette généralité

14 Voir ici même, p. 16.

est rarement visée. La relation d'implication intervient le plus souvent entre des énoncés auxquels, à tort ou à raison, on attribue des valeurs de vérité déterminées. On rencontre en effet des exemples comme:

Jean est un homme *implique* Jean a un coeur.

Il fait jour *implique* Il fait clair.

Il pleut *implique* Un parapluie est utile.

Comme la relation $\rightarrow$ ne porte jamais entre des propositions atomiques, elle ne peut convenir à aucun de ces exemples. Désignons les propositions atomiques en jeu par les symboles p et q . Soutenir les implications qui précèdent revient en fait, non pas à admettre que $\vdash p \supset q$ (ce qui n'est jamais le cas avec p et q atomiques), mais simplement que $\text{val}(p \supset q) = \text{vrai}$.

On définit ainsi une relation plus faible qui ne fait pas appel à la notion de tautologie, et que l'on nommera, pour la distinguer de $\rightarrow$, la relation d'*implication de fait*:

$\text{PIQ} = \text{df } \text{val}(P \supset Q) = \text{vrai}$

Celle-ci semble en effet plus adéquate pour représenter les relations présentes dans les exemples ci-dessus. Cependant, comme le fait remarquer Lewis, il ne faut pas entendre par une telle relation plus que ce qu'elle signifie effectivement et que la table de vérité de la conditionnelle peut nous montrer: PIQ signifie uniquement que le cas où $\text{val}(P) = \text{vrai}$ et $\text{val}(Q) = \text{faux}$ est exclu. La relation d'implication de fait revêt ainsi un caractère peu intuitif qui est pourtant une conséquence immédiate des propriétés de la conditionnelle.

On trouve en effet dans la logique classique certains théorèmes liés à la conditionnelle qui peuvent, une fois interprétés, être ressentis comme peu en accord avec l'idée intuitive d'implication:

(1) $\vdash P \supset . Q \supset P$

(2) $\vdash \sim P \supset . P \supset Q$

Lewis illustre ces expressions, connues sous le nom de «paradoxes de l'implication matérielle», par un exemple devenu

fameux. Prenons une paire de propositions atomiques (indépendantes) p et q , telles que:

p =df «Le vinaigre est acide»,

q =df «Il y a des hommes barbus».

On sait, de par notre expérience des faits, qu'il s'agit de deux propositions vraies. Ainsi, bien que $p \supset q$ n'exprime pas une tautologie, val $(p \supset q) = \text{vrai}$. On dispose donc d'une implication de fait pIq . Autrement dit, «Le vinaigre est acide» implique (de fait) «Il y a des hommes barbus». Le caractère paradoxal d'un tel exemple tient clairement au fait que les deux propositions n'entretiennent aucun lien sémantique.

Comme l'expriment (1) et (2), il suffit pour qu'une conditionnelle soit vraie que son conséquent le soit, ou encore que son antécédent soit faux. La relation I peut porter entre des propositions qui n'entretiennent aucun lien de pertinence. En particulier, assumer une relation comme PIQ ne signifie aucunement être en mesure, d'une manière générale, de déduire Q de P .

Comme le relève Lewis, la différence entre les deux relations $\rightarrow$ et I tient au fait que, là où $P \rightarrow Q$ signifie que Q est déductible de P , PIQ signifie seulement que la vérité de Q peut être déduite de la vérité supposée de P . Ce n'est pas une vérité abstraite et nécessaire que l'on puisse déduire *Jean a un coeur* de *Jean est un homme*, mais il y a un univers qui nous est familier - le nôtre - dans lequel cette implication exprime une relation effective. On peut dire formellement la chose de la manière suivante:

(3) $[\text{val}(P) = \text{vrai et } PIQ] \rightarrow \text{val}(Q) = \text{vrai}$

Selon Lewis, l'obscurité qui entoure les notions de déductibilité et d'implication est due à une confusion entre les deux sortes de relations:

Inability to distinguish between (1) "If p is true and $p \supset q$ holds, then q is true" and (2) "When $p \supset q$ holds, q is deducible from p " is responsible for most of the confusion about material implication and deducibility (Lewis and Langford 1959: 246).

La distinction entre les deux types d'implication reste cependant difficile car la discussion porte sans cesse sur des niveaux linguistiques différents. Si la conditionnelle est bien un élément du langage objet, tout ce que nous pouvons dire à l'aide de $\vdash$, $\rightarrow$ et I demeure métalinguistique.

Le projet de Lewis consiste précisément à ramener tous les éléments du problème au niveau de la langue objet, en construisant un langage logique capable de représenter les différents types de liens que nous avons rencontrés.

S'interrogeant sur la relation $\rightarrow$, Lewis constate qu'elle échappe pratiquement au caractère paradoxal de I . La tautologie apporte effectivement la garantie d'un lien de pertinence minimal. Car, lorsqu'on a $\vdash P \supset Q$, l'une au moins des trois conditions suivantes est satisfaite:

- i. P est inconsistante
- ii. Q est tautologique
- iii. P et Q possèdent au moins une proposition atomique en commun¹⁵.

Si un certain caractère paradoxal demeure avec les conditions i. et ii., cela ne peut être que dans une moindre mesure. Quoi de surprenant dans le fait qu'une proposition inconsistante implique toute proposition puisqu'elle enferme une contradiction, de même que dans celui que la tautologie, qui ne peut pas être fausse, soit déduite de n'importe quel ensemble de prémisses? Quant à la condition iii., elle garantit justement le lien de pertinence qui manquait parfois à l'implication de fait.

15 Pour i et ii la chose étant triviale, il reste à se convaincre que si aucune de ces deux conditions n'est satisfaite P et Q possèdent au moins une proposition atomique en commun (par exemple: $\vdash p \supset p \vee q$). La démonstration consiste à voir que sous l'hypothèse qu'aucune des trois conditions n'est satisfaite, il est impossible que $\vdash P \supset Q$. Par hypothèse, P n'est pas inconsistante; parmi les 2^m familles d'interprétations possibles des m propositions atomiques de P , l'une au moins I est telle que $\text{val}_I(P)=1$. Comme P et Q ne possèdent aucune proposition atomique en commun, les 2^n familles d'interprétations possibles des n propositions atomiques de Q sont indépendantes des familles associées à P . Calculer les valeurs possibles de $P \supset Q$ consiste ainsi à examiner, pour *chacune* des valeurs possibles de P , *toutes* les valeurs possibles de Q . Or, parmi ces 2^n familles d'interprétations associées à Q , une au moins J est telle que $\text{val}_J(Q) = 0$, car Q , par hypothèse, n'est pas tautologique. Par conséquent, il existe au moins une famille d'interprétation K des $m+n$ propositions atomiques de $P \supset Q$ (K étant une sous-famille à la fois de I et de J) telle que $\text{val}_K(P) = 1$ et $\text{val}_K(Q) = 0$. Donc $P \supset Q$ n'est pas une tautologie et on a pas $\vdash P \supset Q$.

La logique des propositions S (*the Survey system*) que Lewis va proposer se construit naturellement à partir de variables et d'opérateurs propositionnels. Mais l'idée de Lewis sera d'ajouter aux opérateurs vérifonctionnels classiques un opérateur susceptible de rendre compte de l'idée de *tautologie* et qui permette d'échapper au caractère paradoxal lié à la simple vérifonctionnalité des opérateurs classiques.

La tautologie, comme nous l'avons vu, est intimement associée à l'idée de *nécessité logique*. La démarche consistera ainsi à introduire dans S un opérateur modal L, que l'on interprétera de la manière suivante:

LP: «P est nécessaire», «P est tautologique».

L'utilisation métalinguistique de $P \rightarrow Q$ peut alors être remplacée par celle, dans S, de l'expression $L(P \supset Q)$. Sur cette base, Lewis va poser la définition d'un nouvel opérateur de conditionnelle dont les propriétés diffèrent passablement de celles de $\supset$. Il s'agit de l'*implication stricte*¹⁶:

$$P \rightarrow Q =_{df} L(P \supset Q)$$

Comme tous les systèmes qui suivront¹⁷, ce premier système moderne de logique modale ne peut disposer de l'extensionnalité qui caractérise la logique classique. Car L et $\rightarrow$ ne sont pas des opérateurs vérifonctionnels; il est bien entendu impensable de les définir à l'aide des opérateurs classiques, qui couvrent toute la combinatoire des valeurs de vérité.

Avant de passer à une présentation plus systématique d'exemples de logiques modales, proposons encore pour terminer trois théorèmes caractéristiques du système S. Ils constituent en effet le pendant «stricte» de nos expressions (1), (2) et (3) et illustrent ainsi la manière dont les problèmes que nous avons évoqués peuvent désormais être formellement représentés.

$$(1') \vdash_S P \rightarrow . Q \supset P$$

16 Nous en conservons ici le nom attribué par Lewis bien que *conditionnelle stricte* serait plus adéquat puisqu'il s'agit d'un opérateur et non d'une relation.

17 Précisons pour la conformité historique que Lewis propose une série de systèmes. Le premier, S1, est construit à partir de l'ensemble d'opérateurs $\{\sim, \wedge, M\}$. On y définit les opérateurs que nous avons utilisés comme suit:

LP = $df \sim M \sim P$ et $P \rightarrow Q = df \sim M (\sim P \wedge Q)$. Cf. Lewis 1918.

$$(2') \vdash_S \sim P \rightarrow . P \supset Q$$

$$(3') \vdash_S P \wedge (P \supset Q) \rightarrow Q$$

Il apparaît ainsi que, dans le langage de S, «what is tautological is distinguishable from what is merely true, whereas this difference does not ordinarily appear in the symbols of [classical] truth-value system» (Lewis and Langford 1959: 247)¹⁸. Il est temps désormais d'aborder la construction effective de quelques systèmes.

3. Quelques systèmes modaux

Comme je l'ai déjà indiqué, proposer un système de logique modale relève d'une grande liberté. En effet, même en se restreignant à une famille de systèmes dits *standard*, on peut encore concevoir de nombreuses logiques différentes. Ces systèmes sont ceux qui satisfont aux conditions que nous avons déjà rencontrées en introduction et que je rappelle ici:

Soit S, un système standard de logique modale,

1. S est une expansion conservatrice de C,
2. MP se définit dans S comme $\sim L \sim P$,
3. S inclut le théorème $\vdash_S LP \supset P$,
4. Si $\vdash_S P$ alors $\vdash_S LP$.

Non seulement, il existe plusieurs manières d'accéder aux systèmes, par des axiomatiques différentes ou des ensembles différents de règles, mais on rencontre bien entendu des systèmes qui déterminent des ensembles différents de théorèmes, et donc qui diffèrent par les inférences modales qu'ils valident.

Notons tout d'abord que les premiers S1 à S3 de Lewis, auxquels nous venons de faire allusion, ne sont pas standard au sens ci-dessus. Ils ne satisfont pas, en particulier, au principe de nécessité de la condition 4. Le premier système que nous allons voir est dû à Feys (1937). Connue sous l'étiquette de T, ce système, qui est un peu plus fort que les systèmes S1 et S2 de

18 On trouvera une description axiomatique simple des premiers systèmes de Lewis dans Hughes & Cresswell 1996: 193-209.

Lewis, s'impose ici par son caractère intuitif et élémentaire comme une base minimale pour la construction des systèmes standard. On en trouvera dans ce qui suit une présentation par des règles de déduction naturelle. Pour des raisons de clarté, il sera aussi fait appel ponctuellement à des considérations qui relèvent de la présentation axiomatique¹⁹.

Mais rappelons tout d'abord la situation de départ. Conformément à la première des conditions, T est une expansion de C, la logique classique des propositions. On y admettra ainsi comme base le vocabulaire ainsi que toutes les règles de C, auxquels on ajoutera pour commencer, l'unique symbole L – compris comme «il est nécessaire que...». On introduira aussi M – «il est possible que...» – mais ceci sans en faire un symbole primitif, grâce à la définition conventionnelle suivante:

M-déf.: $MP =_{df} \sim L \sim P$

Cette manière de procéder, au demeurant conforme à la condition 2, nous dispense de donner une caractérisation de chacune des modalités, puisqu'il suffit alors de donner des règles relatives à l'unique opérateur primitif nouveau: L²⁰.

La condition 3, associée à l'élimination de la conditionnelle, indique très clairement la forme que prendra la règle d'élimination de L. On pose simplement:

Règle Le	n	LP	
		P	n, Le

La règle d'introduction, quant à elle, s'appuie sur le principe de nécessitation de la condition 4. Sa conception demande cependant un certain détour. Selon le principe de nécessitation, si l'on dispose d'une preuve de P, on est en droit d'introduire une expression comme LP. Par exemple, pour montrer que $L(P \supset P)$ est un théorème de T, il suffit de disposer d'une preuve de $P \supset P$, autrement dit de la déduction suivante:

19 La présentation qui suit n'est pas celle de Feys. La notation et les règles utilisées pour la logique C sont celles que l'on trouve dans Grize 1969.

20 Il n'est évidemment pas exclu en pratique de donner des règles pour M, mais celles-ci auront le statut de règles dérivées.

1.			P	hyp.
2.			P	1, rep
3.		P	$\supset$ P	1-2, $\supset$ i

Afin de condenser cette démarche sous la forme d'un schéma de règle de déduction naturelle, l'idée consiste à présenter la preuve attendue sous la forme d'une sous-déduction assortie de deux conditions:

- i. elle ne dépendra d'aucune hypothèse,
- ii. toute réitération y sera interdite.

Ces conditions visent à assurer un caractère de preuve à la sous-déduction. On parle alors de sous-déduction *stricte* et notre exemple devient:

1.				P	hyp.
2.				P	1, rep
3.			P	$\supset$ P	1-2, $\supset$ i
4.		L(P $\supset$ P)			1-3, Li

Si la démarche semble satisfaisante, un schéma générale de règle conçu sur ce modèle s'avère cependant trop faible. Un système qui serait ainsi construit ne donnerait en effet accès qu'à des instances substitutionnelles de théorèmes de C, avec pour seule différence qu'on pourrait les faire précéder de symboles de nécessité. Par exemple l'expression suivante, dont la validité paraît tout à fait intuitive, n'y serait pas démontrable:

$$L(P \supset Q) \supset (LP \supset LQ).$$

Pour accéder à de telles expressions, il faudrait encore se donner les moyens d'utiliser dans la sous-déduction stricte une partie au moins des informations qui peuvent la précéder dans la démarche déductive. Il s'agit en somme d'affaiblir l'interdiction de réitération. Bien que le caractère de preuve de la sous-déduction est alors compromis, ce ne peut être que dans une moindre mesure, car la réitération reste assortie d'une condition forte. On ne réitérera que des expressions dont le caractère de nécessité aura été, sinon démontré, du moins explicitement pris en charge.

Dans ce dernier cas, bien entendu, le résultat de la sous-déduction stricte restera conditionné par la supposition faite.

Cette manière de faire revient à admettre de manière implicite un principe très généralement associé à la nécessité logique: *ce qui s'ensuit logiquement de ce qui est nécessaire est également nécessaire*. Dans la pratique, on passera par une nouvelle règle de réitération, qui sera la seule autorisée dans une sous-déduction stricte²¹.

Règle T-réit.	n	LP	
		L	
		P	n, T-réit.

On peut désormais proposer le schéma de règle suivant pour l'introduction de L, où l'indication L placée au départ de la sous-déduction en rappelle les conditions particulières:

Règle Li	m	L		Sans hypothèse et réitération unique-
				ment avec T-réit.
	n	LP	P	m-n, Li

Voyons maintenant quelques exemples de preuves dans T:

1. $\vdash_T P \supset MP$

1		P	hyp.
2		L ~ P	hyp.
3		P	1, reit
4		~P	2, Le
5		~L ~ P	2, 3, 4, ~i
6		MP	5, M-déf
7		P $\supset$ MP	1-6, $\supset$ i

21 La règle habituelle reste bien entendu valable dans toute autre sous-déduction.

2. $\vdash_T L(P \supset Q) \supset (LP \supset LQ)$

1			$L(P \supset Q)$	
2			LP	
3			$\overline{L}$	$P \supset Q$
4				P
5				Q
6			LQ	
7			$LP \supset LQ$	
8			$L(P \supset Q) \supset (LP \supset LQ)$	

hyp.
hyp.
1, T-réit
2, T-réit.
3,4, $\supset e$
3-5, Li
2-6, $\supset i$
1-7, $\supset i$

3. $\vdash_T L(P \supset Q) \supset (MP \supset MQ)$

1			$L(P \supset Q)$	
2			MP	
3			$\sim L \sim P$	
4			$L \sim Q$	
5			$\overline{L}$	P
6				$P \supset Q$
7				Q
8				$\sim Q$
9				$\sim P$
10			$L \sim P$	
11			$\sim L \sim P$	
12			$\sim L \sim Q$	
13			MQ	
14			$MP \supset MQ$	
15			$L(P \supset Q) \supset (MP \supset MQ)$	

hyp.
hyp.
2, M-déf
hyp.
hyp.
1, T-réit.
5,6, $\supset e$
4, T-réit
5, 7, 8, $\sim i$
5-9, Li
3, réit.
4, 10, 11, $\sim i$
12, M-déf
2-13, $\supset i$
1-14, $\supset i$

4. $\vdash_T \text{MLL}(\text{LP} \supset \text{P})$

1	L	L	LP	hyp.
2			P	1, Le
3			LP $\supset$ P	1-2, $\supset$ i
4		L(LP $\supset$ P)		1-3, Li
5		LL(LP $\supset$ P)		1-4, Li
6		LL(LP $\supset$ P) $\supset$ MLL(LP $\supset$ P)		Th.1, P/LL(LP $\supset$ P)
7		MLL(LP $\supset$ P)		5, 6, $\supset$ e

Ce dernier exemple montre de quelle manière, lorsqu'on part d'une expression qui est un théorème du système – ici $\text{LP} \supset \text{P}$ – on peut toujours la faire précéder d'un L dominant (grâce à la règle Li), ou alors d'un M dominant (grâce au schéma de théorème 1). Remarquons encore que rien ne nous empêche de faire ces introductions successivement, puisqu'à chaque étape le résultat est encore un théorème. Il n'y a ainsi aucun doute qu'une expression comme

$\text{LMMLMLMLL}(\text{LP} \supset \text{P})$

est bien un théorème de T.

Ce phénomène, touchant les théorèmes de T, n'est en fait qu'une manifestation d'un problème plus général concernant l'interprétation que l'on peut donner des expressions du système. Si notre intuition des modalités peut nous guider tant qu'il s'agit de manipuler des formules contenant au plus des modalités simples – celles correspondant en fait à l'expression des modalités classiques: LP, MP, $\sim\text{LP}$, $\sim\text{MP}$ –, elle n'est que d'un faible secours dès que les expressions comprennent des chaînes de modalités²². Si l'on peut encore concevoir une nuance entre LP et MLP, que pouvons nous dire de la distinction entre LMP et LLMP? L'approche «naïve» qui préside à l'élaboration d'un système comme T conduit à une richesse expressive qui dépasse très largement ce qui était visé.

Il paraît ainsi souhaitable d'examiner si les expressions contenant des chaînes de modalités ne sont pas susceptibles

22 On parle de *chaîne de modalités* dans une expression lorsqu'au moins un de ses opérateurs modaux porte sur une partie d'expression dont l'opérateur dominant est soit une modalité, soit une modalité précédée d'un ou plusieurs symboles de négation. On reconnaît par exemple une telle chaîne dans $p \supset L\sim\text{Mq}$, alors que $L(p \supset \sim\text{Mq})$ n'en possède pas.

d'être réduites à des expressions équivalentes n'en comportant aucune. Si tel était le cas, elles ne constitueraient plus que des variantes notationnelles pour des significations modales identiques et le problème de leur lecture tomberait *ipso facto*.

Tout schéma d'équivalence permettant une telle réduction est connu sous le nom de *loi de réduction*, et l'on voit clairement qu'il serait utile de disposer des deux suivants:

$$R1: LP \leftrightarrow LLP$$

$$R2: MP \leftrightarrow MMP$$

Malheureusement, T ne possède aucune de ces deux lois, car si les expressions a et b suivantes y sont bien des théorèmes, ce n'est en revanche pas le cas de leurs converses a' et b':

$$\begin{array}{ll} a: \vdash_T LLP \supset LP & a': \vdash_T LP \supset LLP \\ b: \vdash_T MP \supset MMP & b': \vdash_T MMP \supset MP. \end{array}$$

Il est donc impossible dans T de réduire les chaînes de L successifs ainsi que celles de M. En fait, T ne comporte aucune loi de réduction et si, en un sens large, on nomme modalité toute suite formée des symboles L, M et ~, on dira que T comporte une infinité de modalités irréductibles.

Pour mener à bien la formalisation des modalités, il serait pourtant intéressant de disposer d'un système dont le pouvoir expressif soit limité à un nombre fini de modalités différentes. C'est la principale motivation de la construction du système S4. L'enjeu consiste à renforcer T de manière à obtenir un système où l'expression a' soit un théorème. On se convaincra dans la suite que dans un tel système b' est aussi dérivable, ce qui garantit l'accès aux deux lois R1 et R2.

L'expression a' est connue généralement comme la formule caractéristique de S4 (désormais FC[S4]). Pour construire S4, on modifie T en y ajoutant FC[S4] comme un axiome, ce qu'on note:

$$S4 = T + FC[S4]$$

S4 étant une expansion de T, tous les théorèmes de T sont *a fortiori* des théorèmes de S4. En revanche, le nouvel axiome

permet de déduire des expressions qui n'étaient pas théorèmes de T. On obtient un système parfaitement équivalent, mais par le biais de la déduction naturelle, en modifiant uniquement la règle spéciale de réitération T-réit pour en faire la règle suivante:

Règle S4-réit	$\frac{n \mid LP}{\frac{L \mid LP}{n, S4-réit}}$
----------------------	--

Par cette règle, la seule réitération autorisée dans une sous-déduction stricte est celle de propositions nécessaires, mais cette fois sans modification, autrement dit sans ôter le symbole de nécessité²³. Voyons maintenant quelques schémas de théorèmes spécifiques de S4 (certaines démonstrations étant laissées au lecteur):

5. $(FC[S4]) \vdash_{S4} LP \supset LLP$

6. $\vdash_{S4} MMP \supset MP$

1		MMP		hyp.
2		L~P		hyp.
3		MMP		1, reit
4		L		MP
5				~L~P
6				L~P
7				~MP
8		L~MP		4-7, Li
9		~L~MP		3, M-déf
10		~L~P		2, 8, 9, ~i
11		MP		10, M-déf
12		MMP $\supset$ MP		1-11, $\supset$ i

7. $\vdash_{S4} LP \supset LMLP$ (utiliser le th. 1)

23 Bien que restant valable, la règle T-réit est désormais superflue puisqu'elle est aisément dérivable de S4-réit et Le.

8. $\vdash_{S4} MLMP \supset MP$

1		MLMP		hyp.
2		—		hyp.
3			L~P	1, réit
4			MLMP	hyp.
5			L	4, Le
6				5, M-déf
7				2, S4-réit
8				4,6,7, ~i
9				4-8, Li
10				3, M-déf
11				2, 9, 10, ~i
12				11, M-déf
13				1-12, $\supset i$
		MLMP $\supset$ MP		

9. $\vdash_{S4} LMLMP \supset LMP$

Les deux lois R1 et R2 de S4 autorisent un grand nombre de réductions. Clairement, toute succession d'un même opérateur modal peut être réduite à une seule occurrence de cet opérateur. Mais plus encore, toute chaîne composée d'un mélange des deux opérateurs peut être réduite à une suite de trois, voire de deux, de ces opérateurs. Si bien que seules demeurent irréductibles les chaînes de trois et deux opérateurs sans répétition. Comme S4 ne contient aucune loi de réduction plus forte, on y dénombre les quatorze modalités irréductibles suivantes (l'absence d'opérateur modal étant comprise comme une modalité au sens large du terme)²⁴:

24 Grâce aux équivalences que le jeu de la négation permet entre L et M, on peut se convaincre que toutes les chaînes se réduisent à une modalité irréductible négative lorsqu'elles comprennent un nombre impair de négation et à une modalité affirmative sinon.

Modalités irréductibles de S4:

Affirmatives	Négatives
P	$\sim P$
LP	$\sim LP$
MP	$\sim MP$
LMP	$\sim LMP$
MLP	$\sim MLP$
LMLP	$\sim LMLP$
MLMP	$\sim MLMP$

Si l'interprétation des expressions modales en chaîne reste intuitivement difficile, le système S4 offre une organisation fort intéressante, qui trouve par ailleurs des applications dans des domaines non modaux²⁵. On peut représenter cette organisation par le diagramme relationnel suivant, où $\rightarrow$ représente la relation d'implication, les lignes passant par le centre c des relations de contradiction et les lignes horizontales, soit des relations de contrariété, lorsqu'elles sont situées en dessus de c, soit de sub-contrariété, lorsqu'elles sont en dessous:

25 Sur la possibilité d'interprétations temporelles de S4, cf. Prior 1967: 20-31, Hughes & Cresswell 1968: 127 ss., ainsi que 225-230 pour des versions probabilistes et algébriques

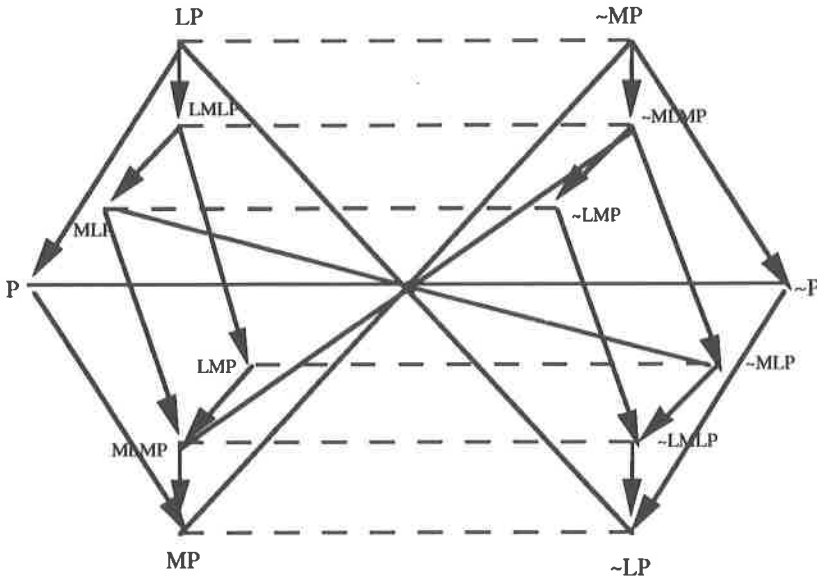


Diagramme A

Quatre chaînes de modalités ainsi que leurs négations demeurent cependant irréductibles. Le système S5, que nous allons aborder maintenant, est conçu pour permettre leur réduction à des modalités simples. Il devra, pour ce faire, disposer des deux lois de réduction suivantes:

$$R3: MP \leftrightarrow LMP$$

$$R4: LP \leftrightarrow MLP$$

Comme R1 et R2, ces deux lois ne sont pas indépendantes, et l'on se convaincra assez aisément qu'il suffit pour en disposer d'avoir comme théorème l'expression suivante, dite formule caractéristique de S5, formule qui reste indémontrable dans les systèmes précédents:

$$FC[S5]: MP \supset LMP$$

Un système muni des quatre lois de réduction R1-R4 permet de réduire toute chaîne de modalité à des modalités simples. D'une manière pratique, toute chaîne se réduira à son dernier

opérateur, précédé ou non d'un symbole de négation. Par exemple:

$$\sim \text{LMMLMP} \leftrightarrow \sim \text{MP}$$

Comme S4, S5 est conçu comme une expansion du système précédent. Il est construit par l'ajout à S4, comme nouvel axiome, de la formule FC[S5]²⁶:

$$\text{S5} = \text{S4} + \text{FC}[\text{S5}]$$

Avec S5, on obtient un système où seules six modalités restent irréductibles (en fait, les modalités de la tradition). On peut représenter leur organisation par un diagramme, plus simple que celui de S4, où l'on retrouve les oppositions du carré traditionnel:

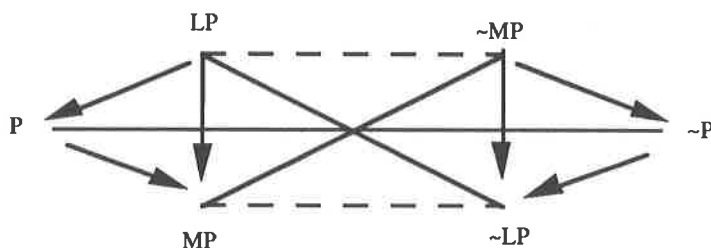


Diagramme B

Par le biais de la déduction naturelle, on obtient le même résultat en modifiant une fois encore les conditions de réitération dans les sous-déductions strictes. Dans S4, la formule caractéristique peut se lire: *si une proposition est nécessaire, alors sa nécessité est nécessaire*, et ceci s'est traduit en déduction naturelle par la règle S4-réit, autorisant la réitération de toute expression nécessaire, sans lui enlever son symbole de nécessité. Mais qu'en est-il de FC[S5]? On peut la lire de la manière suivante: *si une proposition est possible, alors sa possibilité est nécessaire*. En somme, possibilité et nécessité dans S5 sont toujours elles-mêmes nécessaires. La nouvelle règle de

26 On obtient aussi S5 en ajoutant directement FC[S5] à T, FC[S4] étant alors dérivable.

réitération autorisera ainsi également la réitération de propositions possibles. En voici le schéma:

Règle S5-réit	n		MP	
			L	
			MP	n, S5-réit

La règle S4-réit reste bien entendu valable, on est donc autorisé dans S5 à réitérer dans une sous-déduction stricte toute expression modalisée. Le lecteur pourra s'en convaincre en démontrant les deux règles dérivées suivantes, qui facilitent d'ailleurs grandement la déduction dans S5:

Règle S5-réit*	n		~LP	
			L	
			~LP	n, S5-réit*

Règle S5-réit**	n		~MP	
			L	
			~MP	n, S5-réit**

Mais voyons encore quelques exemples de schémas de théorèmes:

10. $(FC[S5]) \vdash_{S5} MP \supset LMP$

11. $\vdash_{S5} L(MP \supset Q) \supset L(P \supset LQ)$

12. $\vdash_{S5} MLP \supset LP$

1			MLP	hyp.
2			~LP	hyp.
3			MLP	1, réit
4			L ~LP	2, S5-réit*
5			L~LP	4-4, Li
6			~L~LP	3, M-déf
7			~~LP	2,5,6, ~i
8			LP	7, ~e
9			MLP $\supset$ LP	1-8, $\supset$ i

Le système S5 apparaît, du point de vue de l'expression des modalités, comme une limite. Toute addition d'un axiome supplémentaire qui donnerait accès à des lois de réduction plus fortes que R3 ou R4 rendrait le système inadéquat pour une formalisation des modalités. En effet, une expansion conçue de cette manière mènerait soit à un système inconsistent, soit à une logique n'exprimant rien d'autre que la logique des propositions classique. Pour se convaincre de ce fait, il suffit de partir du diagramme B des modalités irréductibles de S5. Toute nouvelle équivalence entre deux expressions du diagramme qui sont en relation d'opposition (contraires, subcontraires ou contradictoires) mènerait évidemment à une contradiction. Reste donc la possibilité de nouvelles équivalences entre des expressions dont l'une implique l'autre. Une tentative de ce genre reviendrait à réduire toute modalité à l'une des deux expressions P et $\sim P$.

On trouve un tel exemple, sous le nom de système Triv, dans Hughes et Cresswell (1968, 64-68). Ce système est obtenu par l'ajout à S5 d'un nouvel axiome:

$$\text{Triv} = \text{S5} + P \supset LP$$

Ce système est consistant et admet les deux lois de réduction suivantes:

$$\text{R5: } P \leftrightarrow LP$$

$$\text{R6: } P \leftrightarrow MP$$

Ces lois montrent que les deux opérateurs modaux, non seulement ne sont plus distincts, mais que leur nature modale est compromise. Ils peuvent désormais recevoir une interprétation vérifonctionnelle, la même en fait que celle de l'opérateur unaire connu sous le nom de Verum. Triv est donc une version redondante de la logique classique C; toute expression contenant des opérateurs modaux peut y être réduite à une expression équivalente construite par simple effacement de ces opérateurs.

Si ce système est assurément inapte à formaliser une quelconque intuition des modalités, il présente cependant un intérêt métathéorique. Triv constitue en effet une borne pour les systèmes standard, car tout autre système standard y est contenu.

C'est aussi un système maximal, en ce sens que, comme C, tout ajout d'un nouvel axiome indépendant le rendrait inconsistant²⁷. Ainsi, dans la visée proposée ici, tout système de logique modale devra être contenu dans Triv, tout en étant plus faible que ce dernier.

4. Où il est question de sémantique

Nous avons esquissé jusqu'ici une série de systèmes dont chacun se construit sur la base des précédents et dont l'ensemble s'organise en une chaîne allant du système le plus faible au système le plus fort. Bien qu'il s'agisse de constructions au caractère purement syntaxique, il faut remarquer que les différents choix qui ont prévalu à leur développement n'ont pu se faire que sur la base de motivations sémantiques fortes. De la mise en évidence des conditions élémentaires des systèmes standard à la problématique de la réduction des modalités, le seul guide disponible dans le travail syntaxique est celui du sens que l'on envisage d'attribuer aux expressions modales. Pour Lewis, le but des systèmes modaux consistait à approcher au mieux le sens qu'il attendait pour sa conditionnelle stricte $P \rightarrow Q$: «Q est déductible de P». Le nôtre a été, dans les pages qui précèdent, de suivre notre intuition de la nécessité logique, en s'appuyant au départ sur la notion de tautologie.

Cette approche, que l'on peut qualifier de *présémantique*, est certainement indispensable à l'élaboration d'une syntaxe. Elle ne va pas cependant sans soulever certaines questions. Dès le début, avec Lewis, ce n'est pas un seul système qui est proposé mais une série de systèmes qui diffèrent par les théorèmes mis en évidence et donc par le sens des opérateurs qu'ils tentent de représenter. La question se pose ainsi de savoir comment préciser le sens visé et surtout comment déterminer lequel des systèmes peut être adéquat pour le capturer. Si une certaine lecture des opérateurs peut servir à soutenir l'intuition, elle se révèle

27 On parle généralement de *complétude syntaxique*, et on voit aisément qu'aucun des systèmes plus faibles que Triv que nous avons vu ne peut posséder cette propriété. On peut en effet donner de chacun d'eux une expansion consistante par l'ajout d'un axiome indépendant.

souvent inadéquate par la suite. La lecture de L comme un opérateur de tautologie par exemple s'avère inadéquate dans T dès qu'on fait usage du principe de nécessitation. On peut bien comprendre l'expression $LP \supset P$ comme «si P est une tautologie, alors P est vraie». Cette interprétation mène cependant à l'incohérence avec une expression comme $L(LP \supset P)$, car le *dictum* du premier L ne peut être compris comme une tautologie.

Ces exemples montrent pourquoi il s'est agi de dépasser le stade de l'intuition pour aller à la recherche d'une sémantique précise pour les systèmes modaux. Les premiers essais dans cette direction ont consisté à proposer pour chaque système une lecture différente du symbole de nécessité. C'est ainsi qu'a été élaboré par exemple, sur le modèle des tables de vérité, un calcul sémantique des propositions modales à quatre valeurs²⁸. Malheureusement, les expressions validées par de telles sémantiques ne correspondent pas systématiquement aux théorèmes de la syntaxe. Si les tables à quatre valeurs constituent une bonne méthode pour montrer qu'une expression est indémontrable, elles ne sont malheureusement pas fiables pour la mise en évidence des théorèmes. En termes plus précis, on dit que les systèmes sont *fondés* sous une telle interprétation, car tout théorème de la syntaxe est valide dans la sémantique, mais qu'il leur manque encore la propriété de *complétude sémantique*. On n'est en effet pas assuré que toute expression valide de la sémantique soit bien un théorème de la syntaxe.

Pour disposer d'une sémantique parfaitement adéquate, rendant les systèmes fondés et complets, il faudra attendre la fin des années cinquante, avec principalement Kripke (1959 et 1963), et plusieurs preuves de complétude dont le point commun est de s'appuyer sur une sémantique dite *des mondes possibles*.

Il n'est pas question, dans une présentation comme celle-ci, d'exposer formellement une telle sémantique. Je me contenterai dans ce qui suit d'en décrire l'esprit. On reprend dans ce type de sémantique une idée qui remonte à Leibniz et qui consiste à définir les notions modales en termes de mondes possibles. Les

28 On trouve dans Puttill 1971: 243-251 un tel calcul dit des "*quasi truth tables*" pour les systèmes S3, S4 et S5 de Lewis.

mondes possibles sont conçus comme des alternatives non réalisées du monde réel. On peut les décrire d'une manière relativement précise comme des états de faits concevables et complets. On dit d'un état de fait E qu'il est complet lorsque pour tout autre état de faits, soit il est inclus dans E , soit il est exclu par E . Les mondes possibles sont des entités parfaitement abstraites qui, d'une manière effective, apparaissent dans la sémantique comme des ensembles consistants de propositions décrivant des états de faits complets. Ces ensembles sont ainsi infinis; toute proposition envisageable soit y figurera, soit y sera explicitement niée.

Cependant, lorsqu'il s'agira d'évaluer des propositions modales, on ne considérera des états de faits que les parties concernant effectivement ces propositions. Ceci revêt une certaine importance pratique, car en travaillant ainsi il suffit de considérer des familles d'états de faits. Par ce biais, on se donne la possibilité d'envisager l'ensemble des états de faits à travers une combinatoire finie qui suffit à montrer la validité ou la non-validité d'une expression donnée.

Dans une telle sémantique, l'interprétation d'une expression E consiste en la donnée d'un certain nombre de mondes possibles dans lesquels on attribue des valeurs de vérité différentes aux propositions atomiques de E . Chacun des mondes de l'interprétation peut être relié ou non à d'autres mondes par une relation R que l'on nomme *relation d'accessibilité*²⁹. On dira que l'expression E est valide dans une interprétation i si et seulement si E est vraie dans tous les mondes de i .

Pour déterminer, dans chaque monde, la valeur des propositions composées – celles des propositions atomiques étant données –, on agira comme dans la logique classique pour les opérateurs non modaux, et on appliquera les clauses suivantes pour L et M :

- (a) LP est vraie dans le monde m , ssi P est vraie dans tous les mondes accessibles à m .
- (b) MP est vraie dans le monde m , ssi P est vraie dans au moins un des mondes accessibles à m .

29 Cette relation binaire est au moins réflexive, tout monde est accessible à lui-même.

Le diagramme 1 représente une interprétation qui falsifie l'expression $FC[S4] \text{ } Lp \supset LLp$. Chacun des trois mondes considérés y est représenté par un rectangle où figurent, à gauche, les propositions vraies, à droite, les fausses. Les flèches représentent les relations d'accessibilité.

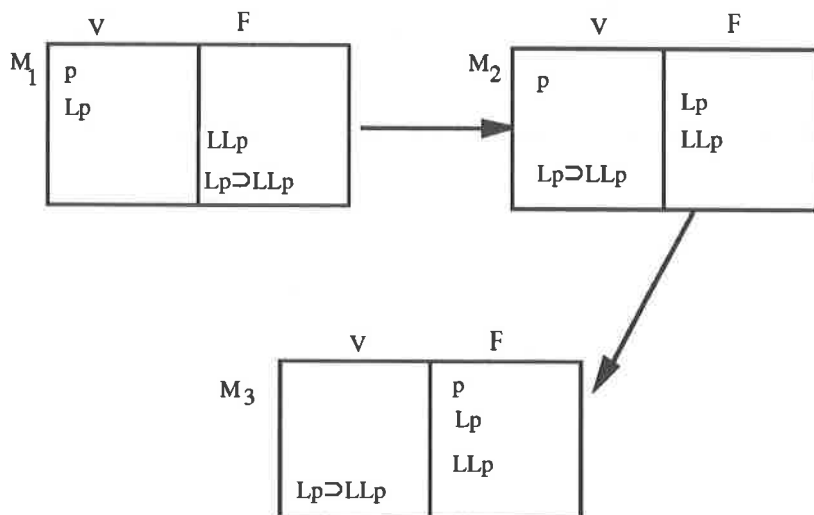


Diagramme 1

A la première ligne, figure dans chacun des mondes la valeur attribuée à p , la seule proposition atomique. Les lignes suivantes représentent respectivement les valeurs calculées – en particulier avec la clause (a) – des expressions Lp , LLp et $Lp \supset LLp$. Le fait de voir figurer dans un des mondes l'expression $Lp \supset LLp$ dans la colonne des propositions fausses suffit à montrer que cette expression n'est pas valide dans cette interprétation.

Bien entendu, le résultat de l'évaluation dans les différents mondes dépend de la manière de distribuer la relation R entre ces mondes. On verra plus loin qu'en modifiant l'interprétation par l'ajout d'une troisième flèche allant de M_1 à M_3 , l'expression considérée y est alors valide.

Remarquons simplement que, la plus grande généralité étant recherchée, une seule interprétation falsifiant une expression

suffit à montrer qu'elle n'est pas *logiquement valide*. On pose en effet la définition suivante:

L'expression E est *logiquement valide* ssi elle est valide dans toute interprétation.

Par l'exemple ci-dessus, nous avons donc montré que l'expression $Lp \supset LLp$ n'est pas logiquement valide, et ceci n'est pas surprenant puisque la sémantique que nous avons décrite est une sémantique adéquate pour le système T. Pour obtenir des sémantiques de S4 et S5, il suffit en fait d'ajouter des conditions sur la relation R d'accessibilité. La seule condition pesant sur R dans T est la réflexivité; on obtient une sémantique pour S4 en ajoutant la transitivité et une sémantique pour S5 en ajoutant encore la symétrie. Mais examinons quelques exemples:

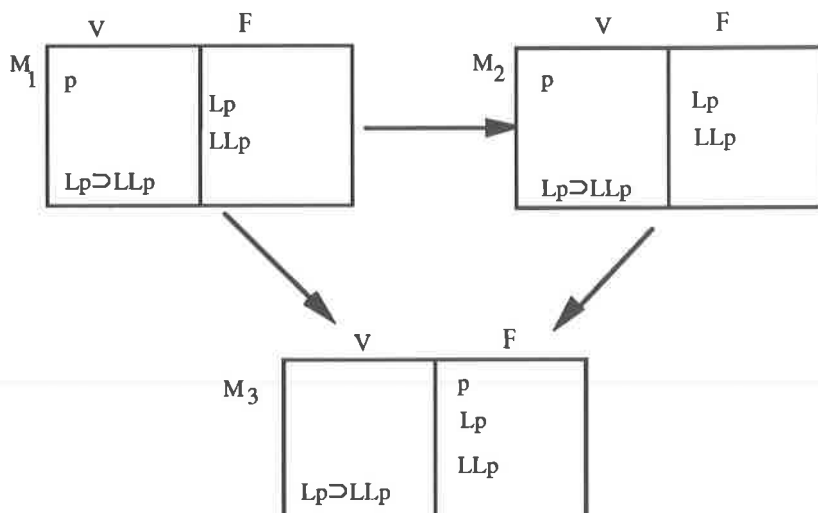


Diagramme 2

En modifiant le diagramme 1 par l'ajout de la condition de transitivité de la relation R , on obtient le diagramme 2. Cette interprétation est désormais adéquate pour S4 et la formule $FC[S4]$ y est valide. L'exemple suivant représente une interpré-

tation qui falsifie $FC[S5]$, c'est-à-dire l'expression $Mp \supset LMp$:

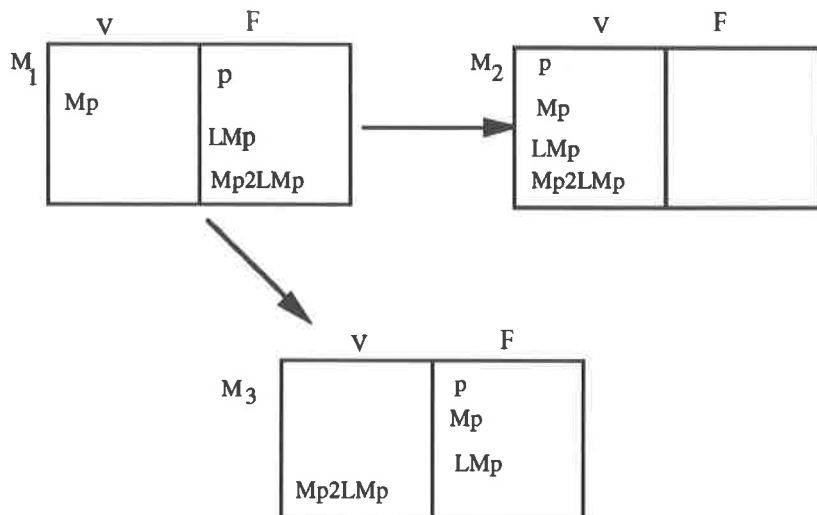


Diagramme 3

Lorsqu'on complète cette interprétation en vue d'une relation R , non seulement réflexive, mais aussi transitive et symétrique, on obtient une interprétation adéquate pour $S5$, qui valide $FC[S5]$:

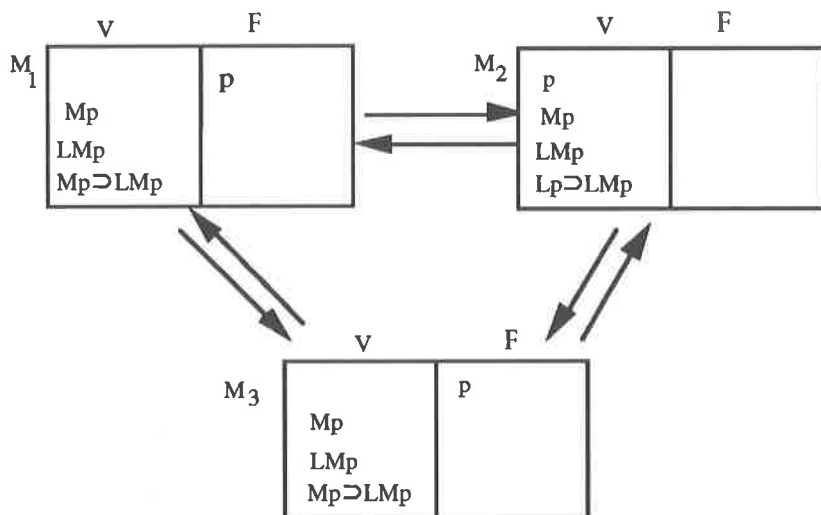


Diagramme 4

Les trois sémantiques auxquelles nous avons fait allusion constituent respectivement des sémantiques pour T, S4 et S5³⁰. Celles-ci se distinguent par le type de liens que l'on peut établir entre les mondes d'une interprétation. D'un point de vue formel, il ne faut pas comprendre par ces liens autre chose qu'une relation binaire R qui détermine un certain nombre de couples à l'intérieur de l'ensemble des mondes d'une interprétation.

Il est cependant utile pour la compréhension de s'aider d'une image intuitive de R. On dira par exemple qu'à partir d'un état de faits, il est possible de «concevoir» ou «d'envisager» un certain nombre d'états de faits différents, de même que dans le monde réel, un sujet peut raisonner sur des situations décrivant le monde tel qu'il aurait pu être. Avec cette analogie, on dispose d'un moyen de comprendre intuitivement ce qui distingue les systèmes. Les différentes conditions qui pèsent sur R peuvent ainsi être comprises comme relevant de capacités plus ou moins fortes de «concevoir», dans une certaine situation, des états de faits différents.

30 On trouvera des démonstrations du caractère fondé et sémantiquement complet des systèmes examinés, dans Hughes & Cresswell 1996: *passim*.

Avec S5, il n'y a aucune limite, c'est l'ensemble des autres situations qui est toujours accessible à partir de n'importe laquelle d'entre elles. Toute situation «concevable», dans ce système, l'est indépendamment d'un quelconque monde de référence. S5 représente au mieux la nécessité logique en son sens le plus large. Avec T et S4, on a affaire à des sens plus faibles de la nécessité car la validité y dépend des limites imposées par l'accessibilité entre mondes. Dans T, l'unique situation qui est toujours accessible est le monde de référence lui-même. Les autres peuvent l'être ou non. Dans S4 par contre, l'accessibilité est plus forte, dans la mesure où «concevoir» une situation c'est non seulement accéder à un autre monde, mais aussi à toutes les situations «concevables» à partir de cet autre monde. La logique modale apparaît comme un cas exemplaire d'investigation logique, dans la mesure où elle propose au philosophe un outil précis pour mesurer les écarts qui séparent les différents usages souvent entremêlés d'une même notion.

Pour conclure

Au terme d'une brève présentation comme celle-ci, on est bien loin de pouvoir conclure avec l'assurance d'avoir dit l'essentiel. Le domaine des logiques modales – certainement le plus classique parmi ceux des logiques non classiques – a pris aujourd'hui une telle ampleur qu'il serait impensable, même dans un vaste traité, de vouloir en dresser un tableau exhaustif. J'aimerais cependant pour terminer dire quelques mots sur un sujet que je n'ai pas abordé et qui mériterait pourtant un long développement: il s'agit des logiques modales des prédicats.

Le traitement logique des modalités ne s'arrête pas au niveau d'une logique des propositions. Sur la base de ce qui précède, on peut en effet construire, avec une relative facilité, un système modal des prédicats. Il suffit pour ce faire de procéder aux expansions modales, non pas à partir de la logique propositionnelle C, mais sur la base de la logique classique des prédicats C*. Ces expansions se font d'une manière tout à fait similaire à ce que nous avons vu pour T, S4 et S5. On obtient par exemple

le système S5* des prédicats, en ajoutant l'ensemble suivant de règles à la logique C*:

$$S5^* = C^* + \{Le, Li, S4\text{-réit}, S5\text{-réit}\}.$$

Si la construction en est aisée, le système obtenu ne va pas sans poser un certain nombre de problèmes difficiles, dès qu'il s'agit d'en donner une interprétation. Je ne relèverai ici qu'un seul exemple parmi ces problèmes, peut-être le plus connu. Examinons deux expressions bien formées de S5*:

1. $L(\forall x)(ax \supset bx)$
2. $(\forall x)L(ax \supset bx)$.

Si dans la première le symbole de nécessité porte sur une proposition, on constate que dans la seconde il porte sur une fonction propositionnelle. Avec l'unique symbole L, on est donc en mesure dans S5* d'exprimer deux sortes de modalités différentes. Dans un cas, la nécessité qui est en cause est celle d'une proposition ou d'un *dictum*. Ce qui est qualifié de nécessaire, comme dans les systèmes propositionnels, est la vérité de la proposition. On parle alors, en suivant en cela la tradition, de modalité *de dicto*. Dans l'autre cas, la modalité est forcément d'une nature différente puisqu'elle porte sur une entité qui n'a pas de valeur de vérité. Ce qui est qualifié de nécessaire dans l'expression 2, c'est le lien qu'entretiennent les prédicats a et b. On dira ici que la modalité est *de re*.

Ces deux modalités, comme l'ont abondamment montré les logiciens médiévaux, sont fort différentes. Et c'est une question difficile et débattue que de savoir si l'une peut être ramenée à l'autre.

Sans vouloir entrer dans plus de détails, relevons simplement que, par la considération de plusieurs systèmes, on est capable de représenter, d'une manière précise, les différentes solutions qui peuvent être apportées à un tel problème, ainsi que les conséquences qui en découlent. La formalisation logique, une fois de plus, constitue dans le domaine des modalités un instrument privilégié. Et l'on peut même dire qu'avec les systèmes quantifiés, elle a permis d'élargir considérablement un des champs

d'investigation les plus fructueux de la philosophie contemporaine³¹.

Séminaire de logique
Faculté des lettres
Université de Neuchâtel
Espace Louis-Agassiz 1
CH 2000 NEUCHÂTEL

Bibliographie

- ARISTOTE (1946). *De l'interprétation*. éd. Tricot. Paris: Vrin.
 ARISTOTE (1947). *Les Premiers Analytiques*. éd. Tricot. Paris: Vrin.
 BLANCHÉ R. (1970). *La logique et son histoire, d'Aristote à Russell*. Paris: Armand Colin (éd. augmentée par J. Dubucs, 1996).
 CHELLAS B. F. (1980). *Modal Logic: An Introduction*. Cambridge: Cambridge University Press.
 FEYS R. (1937). Les logiques nouvelles des modalités. *Revue Néoscholastique de Philosophie* 40, 517-553 et 41, 217-252.
 FEYS R. (1965). *Modal Logics*. Louvain, Paris: Nauwelaerts, Gauthier-Villars.
 GARDIES J.-L. (1979). *Essai sur la logique des modalités*. Paris: PUF.
 GRIZE J.-B. (1969, 1971). *Logique moderne I et II*. La Haye, Paris: Mouton, Gauthier-Villars.
 HUGHES G. E., CRESSWELL M. J. (1996). *A New Introduction to Modal Logic*. London: Routledge.
 KNEALE W. & KNEALE M. (1962). *The Development of Logic*. Oxford: Clarendon Press.
 KONYNDYK K. (1986). *Introductory Modal Logic*. Louvain: University of Notre-Dame Press.

31 Sur les logiques des prédicats, cf. en particulier Konyndyk 1986: 77 ss. et Hughes and Cresswell 1996: 235 ss. On trouve aussi une anthologie d'articles célèbres sur les aspects philosophiques des logiques modales des prédicats dans Linsky 1971.

- KRIPKE S. A. (1959). A completeness theorem in modal logic. *Journal of Symbolic Logic* 24, 1-14.
- KRIPKE S. A. (1963). Semantical considerations on modal logic. *Acta Philosophica Fennica* 16, 83-94.)Rééd. dans Linsky (1971).
- LEWIS C. I. (1918). *A Survey of Symbolic Logic*. Berkeley: University of California Press.
- LEWIS C. I. & LANGFORD C. H. (1932). *Symbolic Logic*. New York: Dover.
- LINSKY L. (1971). *Reference and Modality*. Oxford: Oxford University Press.
- MATES B. (1961). *Stoic Logic*. Berkeley: University of California Press.
- MIGNUCCI M. (1978). Sur la logique modale des Stoïciens. In: Ouvrage collectif. *Les Stoïciens et leur logique*. Paris: Vrin, 317-346.
- PATTERSON R. (1995). *Aristotle's Modal Logic. Essence and Entailment in the Organon*. Cambridge: Cambridge University Press.
- PRIOR A. (1967). *Past, Present and Futur*. Oxford: Clarendon Press.
- PURTILL R. L. (1971). *Logic for Philosophers*. New York: Harper & Row.
- READ S. (1988). *Relevant Logic. A Philosophical Examination of Inference*. Oxford: Blackwell.
- RUSSELL B., WHITEHEAD A. N. (1910-13). *Principia Mathematica*. Cambridge: Cambridge University Press.
- VAN BENTHEM J. (1988). *A Manual of Intensional Logic*. Stanford: CSLI. (2e ed.).
- ZALTA E. N. (1988). *Intensional Logic and the Metaphysics of Intensionality*. Cambridge, London: The MIT Press.

LOGIQUE MULTIVALENTE

Marek BLASZCZYK

Jan Lukasiewicz

Jan Lukasiewicz¹ naît le 21 décembre 1878 à Lwów² et y demeure jusqu'en 1915. Passées les années de gymnase, il entre en 1897 à l'Université de Jan Kazimierz où il fait des études d'histoire de la philosophie sous la direction de Kazimierz Twardowski (disciple de Brentano). Devenu docteur (1902), il passe trois ans dans l'enseignement privé tout en occupant les fonctions de bibliothécaire adjoint à l'Université. En 1905, ayant reçu une bourse du gouvernement autonome de Galicie, il complète sa formation philosophique à Berlin, puis à Louvain, où il assiste au cours du futur cardinal Mercier. Nommé privatdozent à Lwów (automne 1906) il assiste au séminaire de Meinong, à Graz. Il reçoit en 1911 le titre de professeur extraordinaire de philosophie à l'Université de Lwów, où il enseigne jusqu'en 1915.

En 1915, il est invité à donner à Varsovie une série de conférences à l'occasion de la réouverture de l'Université. En 1917, il est élu pro-recteur de l'université, après s'être occupé du département de l'enseignement supérieur au ministère polonais de l'éducation publique, comme ministre dans le cabinet Paderewski (après la proclamation de l'indépendance polonaise), il est nommé deux fois recteur en 1922-1923 et en 1931-1932. En 1923, il reçoit le titre de grand commandeur de l'ordre polonais *Polonia Restituta*, et celui de grand commandeur de l'ordre

1 La bibliographie d'après l'article de B. Soboci «In memoriam J. Lukasiewicz», *Philosophical Studies*, Maymooth, XII 1956, T. VI.

2 Lwów (Leopolis, Lemberg) la capitale de la province Galicie, alors sous la domination autrichienne, demeurait profondément polonaise, surtout dans son université. La langue maternelle de J. Lukasiewicz est le polonais

Travaux de logique, 11, 1997.

hongrois du mérite. En 1926, il abandonne la philosophie pour la logique mathématique. Il se consacre désormais à celle-ci d'une façon exclusive, et fonde avec Stanisław Lesniewski un centre de recherches logiques qui devient l'«École de Varsovie». En 1932, on lui donne le titre inhabituel de «professeur ordinaire et honoraire». En 1935, la ville de Varsovie lui offre un prix pour son oeuvre scientifique. En 1937, il entre à l'Académie polonaise des sciences et des arts de Cracovie; il fait également partie des Sociétés des arts et des sciences de Lwów et de Varsovie. En 1938, il reçoit le titre de docteur *honoris causa* de l'Université de Munster («distinction d'ordre purement scientifique et sans arrière-plan politique»). La guerre interrompt ses travaux jusqu'en 1946. Le 17 juillet 1944, il quitte Varsovie avec l'espoir de gagner la Suisse, mais il est bloqué en Allemagne, où il vit dans la clandestinité jusqu'en 1945, date à laquelle il gagne Bruxelles, il accepte, en 1946, l'offre du gouvernement irlandais qui accueille un certain nombre de savants polonais. Dans la dernière période de sa vie (1946-1956), il vit hors de Pologne. En 1953, il reçoit un prix d'une fondation philosophique anglaise pour l'avancement qu'il a apporté à la logique philosophique, puis, en 1955, Trinity College (Dublin) lui confère le titre de docteur *honoris causa*. Łukasiewicz meurt le 13 février 1956 à Dublin.

L'oeuvre de Jan Łukasiewicz

L'oeuvre de Łukasiewicz, comme sa vie, peut être «partagée» en trois parties³:

1. Philosophique (Lwów-Warszawa) jusqu'à 1920.
2. Logique en Pologne (de publication de la logique trivalente jusqu'à la guerre).
3. Logique à Dublin (1946-1956).

Dans la première période, jusqu'en 1920, les publications de Łukasiewicz sont consacrées aux problèmes de la philosophie et à l'histoire de la philosophie. A la suite de Twardowski, il pense

3 Cette classification est subjective.

à la philosophie «d'un point de vue logique», cela demeura pendant toute sa vie. Autrement dit, il s'agissait d'appliquer à ces problèmes la clarté, la rigueur et la précision de l'analyse logique. Au cours de cette période, Lukasiewicz met au clair sa conception de la méthode philosophique adéquate, qu'il appelle «scientifique». Il l'applique à trois sujets:

- la notion de cause (Analyse et construction du concept de cause 1906),
- le principe de contradiction (Sur le principe de contradiction dans l'oeuvre d'Aristote 1910),
- l'idée de probabilité (Sur les fondements logiques du calcul des probabilités 1913).

Mais ces trois sujets n'en font qu'un dans la pensée de Lukasiewicz et sont tous organisés autour d'une idée unique. Cette idée n'appartient ni à la logique, ni à la méthodologie, mais bien à la métaphysique et, plus précisément, à une métaphysique que l'on pourrait appeler «existentielle» puisqu'elle concerne la liberté de l'homme dans le monde. Cet article écrit en 1927 et publié pour la première fois en 1961 sous le titre *Sur le déterminisme* dans ses Oeuvres choisies (direction de Jerzy Slupecki *Z zagadnień Logiki i Filozofii* PWN, Warszawa 1961; traduction anglaise dans [10]). Durant cette période, Lukasiewicz publie aussi un certain nombre d'articles consacrés à l'histoire de la logique et surtout à la logique d'Aristote ainsi qu'à la logique stoïcienne.

La deuxième période est marquée par l'invention, ou plutôt la découverte, de la logique multivalente. Cette logique est une formalisation de ses travaux précédents. Les chefs-d'oeuvre de cette période sont: *Sur la logique trivalente* (1920), *Untersuchungen Über den Aussagenkalkül* (avec Tarski en 1930), *Philosophische Bemerkungen zu mehrwertigen Systemen des Aussagenkalküls* (1930). Cette période apporte aussi des travaux sur l'histoire de la logique; deux publications sur la logique mono-axiomatique (*Remarques sur l'axiome de Nicod*, 1931 et *Der Äquivalenzkalkül*, 1939) sont particulièrement intéressantes. C'est à ce moment-là que Lukasiewicz axiomatise la syllogistique aristotélicienne, ses réflexions aboutissent à la rédaction

d'un livre qui paraît en 1951 sous le titre *Aristotle's syllogistic from the standpoint of modern formal logic* (traduit en français en 1972 sous le titre: *La syllogistique d'Aristote dans la perspective de la logique formelle moderne*).

La dernière période, à Dublin, apporte deux importants systèmes de logique modale (quadrivalente): *A system of modal logic I-II* (1953) et de logique intuitionniste: *On the intuitionistic theory of deduction* (1952).

Les travaux, que je présente dans cette section, ont été choisis de manière totalement subjective. J'ai choisi de mettre en lumière les questions les plus importantes, celles qui mettent en évidence l'oeuvre du maître dans tous les aspects de sa recherche. On trouvera en annexe une bibliographie de Lukasiewicz.

Notations de Lukasiewicz

Jan Lukasiewicz crée une notation simplifiant l'écriture des propositions symboliques. Cette dernière n'utilise pas de parenthèses et change la place «traditionnelle» du foncteur.

– Les énoncés: $p, q, r, \dots$

– Les connecteurs:

négation (non)	N
conjonction (et)	K
disjonction (ou inclusif)	A
implication matérielle (si ... alors)	C
biconditionnelle (si et seulement si)	E

Règle générale: tous les connecteurs (foncteurs) sont placés devant les opérandes ou l'opérande.

Quelques exemples:

$\neg p$	non p	Np
$p \wedge q$	p et q	Kpq
$p \vee q$	p ou q	Apq
$p \rightarrow q$	si p alors q	Cpq

$p \equiv q$	p si et seulement si	E_{pq}
$[(p \rightarrow q) \wedge p] \rightarrow q$	$CKCpqqq$	
$[(p \rightarrow q) \wedge \neg q] \rightarrow \neg p$	$CKCpqNqNp$	
$\neg(p \vee q) \equiv (\neg p \wedge \neg q)$	$ENApqKNpNq$	

Exemple de preuve:

La preuve du théorème d'identité «si p alors p » (Bochenski «Formale Logik» [1956])

Les axiomes:

1. $CCpqCCqrCpr$
2. $CCNppp$
3. $CpCNpq$

La preuve: Voyons maintenant comment il est possible de dériver à partir des axiomes, et à l'aide des règles d'inférence, de nouvelles thèses. Nous allons déduire des axiomes 1–3 la loi d'identité Cpp . Cette déduction requiert deux applications de la règle de substitution et deux autres de la règle de détachement. Elle se fait de la façon suivante:

1. $q/CNpq \times C3-4$
4. $CCCNpqrCpr$
4. $q/p, r/p \times C2-5$
5. Cpp .

La première ligne s'appelle «ligne de dérivation». Elle est composée de deux parties que sépare le signe $\times$. La première partie, 1. $q/CNpq$ signifie qu'il faut substituer $CNpq$ à q dans (1.). Pour abréger la démonstration, nous ne mentionnons pas la thèse qui résulte de cette substitution. Elle aurait la forme suivante:

1a. $CCpCNpqCCCNpqrCpr$

Quant à la seconde partie, $C3-4$, elle indique comment est construite cette thèse que nous avons omise, et montre clairement qu'il est permis de lui appliquer la règle de détachement: la thèse (1a.) commence par C ; viennent ensuite les axiomes (3.) comme antécédent, et (4.) comme conséquent. Nous pouvons donc détacher (4.) à titre de thèse nouvelle. La ligne de dériva-

tion qui se trouve avant (5.) a une explication semblable à celle de (1.). La barre oblique (/) est le signe de la substitution et le tiret (–) le signe indiquant le détachement.

Logiques multivalentes

Les systèmes de logique propositionnelle peuvent être définis de plusieurs façons. Pour rendre les choses plus «claires», je définis ci-dessous les systèmes proposés par Lukasiewicz en 1929 dans ses *Eléments de logique mathématique* [4].

Système $\mathfrak{B}_2$ ou C–N

Le calcul propositionnel est un système déductif. Dans chaque système déductif il y a un certain nombre de termes indéfinis dits *termes primitifs* de la théorie. La signification de ces termes est issue de leurs insertions dans les axiomes. Les propositions, qu'on appelle les *axiomes* de la théorie, sont reconnues sans preuve, mais elles doivent être évidentes par elles-mêmes. Toutes les propositions acceptées comme axiomes dans un système peuvent être des thèses et devenir prouvables dans d'autres systèmes. Le système présenté ici ne contient pas que des termes primitifs mais aussi des termes définis. Dans les définitions, le symbole « $\equiv$ » signifie «est» et la définition a la forme suivante *definiendum* = *definiens*. Chaque côté du signe « $\equiv$ », c'est-à-dire le *definiendum* et le *definiens*, a exactement la même signification et peut être remplacé par l'autre sans aucune limitation. Les définitions doivent obéir à toutes les règles de construction de la définition explicite correcte. D'après Lukasiewicz, les définitions ne peuvent pas jouer de rôle créatif dans le système. Elles servent à abrégé certaines expressions appartenant au système ou, par l'introduction d'un terme nouveau, elles peuvent contribuer à une nouvelle intuition dans la théorie.

Dans le système présenté il y a trois règles: une *règle de remplacement* sous entendue, signifiant simplement que chaque forme de *definiendum* peut être remplacée par son *definiens* et

vice-versa, deux règles d'inférence: *règle de substitution* et *règle de détachement*. La première nous permet de déduire des thèses nouvelles à partir d'une assertion du système, en substituant, à une variable, une expression α douée de sens toujours la même pour une même variable. Les expressions douées de sens sont définies inductivement, de la façon suivante:

1. Toute variable propositionnelle est une expression douée de sens;
2. $N\alpha$ est une expression douée de sens si α en est une;
3. $C\alpha\beta$ est une expression douée de sens si α et β sont des expressions douées de sens.

La règle de détachement est le *modus ponens* des stoïciens. Si l'on fait l'assertion d'une proposition de type $C\alpha\beta$ ainsi que de son antécédent α , on peut également faire celle de son conséquent β et le détacher de l'implication à titre de nouvelle thèse.

Ces deux règles nous permettent de déduire, de notre ensemble d'axiomes, toutes les thèses (propositions vraies) du système C-N. Si l'on veut avoir, outre C et N, d'autres foncteurs dans notre système, il faut les introduire par l'intermédiaire de définitions. Pour la conjonction K par exemple, il y a deux façons de le faire. 1) La conjonction «p et q» et «il-n'est-pas-vrai-que (si p alors non-q)» ont le même sens, et l'on pourrait exprimer cette connexion entre Kpq et $NCpNq$ par la formule suivante:

$$Kpq = NCpNq$$

nous autorisant à remplacer le *definiens* par le *definiendum* et vice versa. 2) Nous pourrions également exprimer la connexion existant entre Kpq et $NCpNq$ par deux implications converses:

$$CKpqNCpNq \text{ et } CNCpNqKpq.$$

Il est nécessaire d'être familiarisé avec la pratique des démonstrations si l'on désire déduire à partir des axiomes, par exemple la loi de simplification $CpCqp$ ou la loi de commutation $CCpCqrCqCpr$. Une autre méthode plus simple existe, inventée par le logicien américain Charles S. Peirce vers 1885. Elle se fonde sur le «principe de bivalence», qui établit que toute proposition est vraie ou fausse, c'est-à-dire que des deux valeurs

de vérité possibles – le vrai ou le faux –, elle en choisit une et uniquement une. Il s'agit de distinguer ce principe de la loi du tiers-exclu (*tertium non datur*) selon laquelle de deux propositions contradictoires l'une doit nécessairement être vraie. Cette loi fut considérée comme la base de la logique par les stoïciens, en particulier par Chrisippe.

Toutes les fonctions de la théorie de la déduction sont des fonctions de vérité, c'est-à-dire que leur vérité et leur fausseté ne dépendent que de la vérité et de la fausseté de leurs arguments. Créons maintenant une nouvelle constante appelée «proposition fausse» et désignons-la par 0; désignons de façon analogue par 1 toute proposition vraie. Nous pouvons alors définir la négation de la manière suivante:

$$N1 = 0 \text{ et } N0 = 1.$$

Ces symboles signifient que la négation d'une proposition fausse a le même sens qu'une proposition vraie (ou plus brièvement: est vraie), et la négation d'une proposition vraie est fausse. Pour l'implication, nous avons les quatre définitions que voici:

$$C00 = 1, C01 = 1, C10 = 0, C11 = 1.$$

Autrement dit, une implication est fausse dans le seul cas où son antécédent est vrai et son conséquent faux; dans tous les autres cas, elle est vraie. Cette définition, la plus ancienne que l'on ait de l'implication, est attribuée à Philon de Mégare et a été adoptée par les stoïciens.

Or, si nous voulons vérifier, dans la théorie de la déduction, une expression douée de sens et contenant au moins les deux foncteurs N, C, il faut substituer, aux variables figurant dans cette expression, selon toutes les combinaisons possibles, les symboles 1 et 0, et réduire les formules ainsi obtenues. Si, après réduction, toutes les formules ont 1 comme résultat final, l'expression est vraie: c'est une thèse; mais si une quelconque de ces formules a 0 pour résultat final, l'expression est fausse. Exemple: la loi de la transposition $CCpqCNqNp$ donne:

$$\text{Pour } p/0, q/0: CC00CN0N0 = C1C11 = C11 = 1,$$

$$\text{Pour } p/0, q/1: CC01CN1N0 = C1C01 = C11 = 1,$$

Pour $p/1, q/0$: $CC10CN0N1 = C0C10 = C00 = 1$,

Pour $p/1, q/1$: $CC11CN1N1 = C1C00 = C11 = 1$.

On obtient, au terme de toutes les substitutions, le résultat: 1, la loi de transposition est donc bien une thèse de notre système. Autre exemple: $CNCpqq$. Il suffit ici d'un seul essai de substitution:

$p/1, q/0$: $CNC100 = CN00 = C10 = 0$.

Le résultat final étant 0, l'expression $CNCpqq$ est fausse.

Les quantificateurs sont désignés par des caractères grecs majuscules, le quantificateur universel par Π et le quantificateur particulier ou existentiel par Σ .

Définition de la logique bivalente

Termes indéfinis:

p, q, r, s
 $N, C,$

variables propositionnelles
négation, implication.

C	0	1	N
0	1	1	1
1	0	1	0

Termes définis:

$Kpq = NCpNq$

$Apq = CNpq$

$Epq = NCCpqNCqp$

conjonction
alternative (ou inclusive)
équivalence.

Axiomes:

L1. $CCpqCCqrCpr$

L2. $CCNppp$

L3. $CpCNpq$

système C-N de Lukasiewicz (1929)

loi du syllogisme hypothétique

loi de Clavius⁴

loi de Duns Scot

4 L'expression «Si (si non-p alors p), alors p» a été utilisée par Euclide pour la démonstration du théorème mathématique suivant: «si le produit de deux entiers, a et b, est divisible par un nombre premier n, alors si a n'est pas divisible par n, b doit être divisible par n».

Règles d'inférence:

Règle de substitution

Règle de détachement

modus ponens.

Une logique bivalente mono-axiome

En 1931, Lukasiewicz a publié dans [7] (traduction anglaise dans [10]) les résultats de ses recherches sur l'axiome de Nicod ainsi que son système ne contenant qu'un axiome. Le système de Nicod est un point de départ pour des réflexions plus profondes sur la déduction généralisée. Nous présentons ce système pour son caractère inhabituel et parce qu'à l'époque il était vivement discuté par tous les membres de l'«école polonaise de logique».

Termes indéfinis:

p, q, r, s ...

variables propositionnelles

D

foncteur⁵ de Sheffer |

D	0	1
0	1	1
1	1	0

Axiome: $(\mathfrak{N}) \text{ DDpDqrDDtDttDDsqDDpsDps.}$ *Règles d'inférence:*

Règle de substitution

Règle de détachement

 $D\alpha D\beta\gamma$

L'axiome de Nicod n'est pas une proposition évidente. Lukasiewicz en propose sa version. Il démontre qu'un axiome à quatre variables a la même signification qu'un d'axiome à cinq variables. Voici l'axiome:

 $(\mathfrak{Z}) \text{ DDpDqrDDsDssDDsqDDpsDps.}$

5 Le créateur du terme «foncteur» est T. Kotarbinski.

Il pose ensuite l'équivalence de ces deux propositions en montrant que l'axiome $\mathfrak{N}$ est plus général que $\mathfrak{B}$, et que l'on peut déduire $\mathfrak{N}$ à partir de $\mathfrak{B}$. Notre attention est attirée par la règle de détachement appliquée par Nicod. Si la thèse $D\alpha D\beta\gamma$ fait partie du système et que la thèse α appartient aussi au système, alors on peut ajouter l'expression γ comme une thèse du système. Lukasiewicz ajoute également β et son raisonnement est le suivant: on considère que la règle est vraie, c'est-à-dire:

$$D\alpha D\beta\gamma = 1$$

on admet que la proposition α est aussi vraie en obtenant

$$D1D\beta\gamma = 1.$$

Pour simplifier le raisonnement et surtout le rendre plus intuitif, on substitue au foncteur D sa propre définition conçue sur la conjonction $Dpq = NKpq$ (D/NK), on obtient ainsi:

$$NK1NK\beta\gamma = 1.$$

Pour que la proposition $NK1NK\beta\gamma$ soit vraie, il faut que la proposition $K1NK\beta\gamma$ soit fausse en raison de la négation au début de la thèse. On voit que ce n'est possible que dans le seul cas où la proposition $K\beta\gamma$ est vraie, c'est-à-dire dans le cas où β et γ sont vraies. On peut alors non seulement ajouter la proposition γ au système comme Nicod le souhaite, mais également β .

Puis Lukasiewicz présente une version avec le quantificateur universel pour prouver que son axiome est plus court que celui de Nicod:

$$(No) \Pi p \Pi q \Pi r \Pi s \Pi t DDpDqrDDtDttDDsqDDpsDps$$

$$(Lo) \Pi p \Pi q \Pi r \Pi s DDpDqrDDsDssDDsqDDpsDps.$$

Cette recherche, confiée à M. Wajsberg comme travail de diplôme, a donné un résultat inattendu. En 1927, ce dernier définit la notion «organique» et l'attache à l'axiome. L'axiome de Nicod et celui de Lukasiewicz ne sont pas «organiques» parce qu'ils contiennent les parties $DtDtt$ et $DsDss$ qui sont elles-mêmes des thèses du système. Pour éviter cela, Wajsberg propose son axiome et montre qu'on peut remplacer le $\mathfrak{B}$ et le $\mathfrak{N}$ dans le système de Nicod:

($\mathcal{W}$) $DDpDqrDDDsrrDDpsDpsDpDpq$.

L'intuition de Wajsberg est prophétique car, quelques années plus tard, S. Lesniewski démontre que, dans la déduction de la thèse $DtDtt$, Nicod a commis une erreur, erreur répétée par Lukasiewicz qui écrivit un erratum où il montre qu'il avait trouvé un axiome ne contenant pas cette thèse, axiome différent de celui de Wajsberg.

($\mathfrak{L}1$) $DDpDqrDDpDrpDDpDrpDDsqDDpsDps$.

Les travaux sur le système de Nicod ont provoqué la naissance de la déduction généralisée et toute une série de systèmes mono-axiomatiques basés sur l'implication (sans négation) et aussi, à mon avis, le système équivalentiel (ou biconditionnel) [8] de Lukasiewicz.

Logiques multivalentes de Lukasiewicz

Le premier système de la logique trivalente est publié par Lukasiewicz [3] en 1920. Puis, avec tous les membres de l'«Ecole de Lwów et Varsovie», il publie plusieurs systèmes de la même classe de logique à ensemble fini de valeurs $\mathfrak{L}_3$, $\mathfrak{L}_4$, ..., $\mathfrak{L}_n$, ainsi qu'un système de logique à ensemble infini (mais dénombrable) de valeurs $\mathfrak{L}_{\aleph_0}$.

Logique trivalente $\mathfrak{L}_3$

Il s'agit du premier système multivalent, créé par Lukasiewicz en 1920 [3].

Termes indéfinis:

$p, q, r, s \dots$
 $N, C,$

variables propositionnelles
négation, implication

C	0	1/2	1	N
0	1	1	1	1
1/2	1/2	1	1	1/2
1	0	1/2	1	0

Termes définis:

$Apq = CCqpp$

alternative (ou inclusive)

$Kpq = NANpNq$

conjonction

$Epq = KCpqCqp$

équivalence.

Axiomes: système $\mathfrak{L}_3$ de Lukasiewicz (1920)

L1. $CqCpq$

loi d'affirmation de conséquent

L2. $CCpqCCqrCpr$

loi de syllogisme hypothétique

L3. $CCCpNppp$

L4. $CCNqCNpCpq.$

Règles d'inférence:

Règle de substitution

Règle de détachement

modus ponens.

Au début, le système $\mathfrak{L}_3$ n'était pas fonctionnellement complet, au sens où on ne parvient pas à définir tous les foncteurs possibles⁶. En 1936, J. Slupecki [15] propose un foncteur T et deux axiomes:

p	Tp
0	1/2
1/2	1/2
1	1/2

L5. $CTpNTp$

L6. $CNTpTp$

⁶ Système à trois valeurs a $3^{61} = 27$ foncteurs possibles à une variable et $3^{33} = 19683$ à deux variables.

Suite à ces modifications le système est devenu complet dans les deux sens de cette notion.

Au début de ses réflexions, la troisième valeur était pour Łukasiewicz synonyme de «possibilité» ou d'«indéterminé». Mais, il a abandonné cette première signification après une étude systématique de la notion «possible» alors déjà bien définie.

Remarques. On peut remarquer que certaines lois qui sont valides en $\mathfrak{B}_2$ ne le sont pas dans $\mathfrak{B}_3$, par exemple le *tertium non datur* $ApNp$ ou principe de contradiction $NKpNp$. Ceci est conforme à l'idée de l'auteur. De plus, il y a aussi un certain nombre de contre tautologies (propositions toujours fausses) en logique bivalente et vraies en trivalente par exemple $EpNp$. Łukasiewicz a défini sa logique en conformité avec sa propre métaphysique héritée de Brentano, Twardowski et Meinong.

Une autre caractéristique qu'on peut y voir, c'est que les propositions dites «indéterminées» dans la conjonction et la disjonction donnent des propositions «indéterminées», ce qui est conforme à l'intuition.

Trivalence et modalité: la logique modale a toujours intéressé Łukasiewicz. Il a essayé par deux fois de construire des systèmes modaux basés sur la logique multivalente. L'idée était simple: tenter de définir les foncteurs modaux comme tous les autres foncteurs du système. Le premier système créé en 1920 était le suivant:

Il a défini le foncteur «il-est-possible-que» par Mp (voir le tableau) et le foncteur «il-est-nécessaire-que» par $Lp = NMNp$. Łukasiewicz voulait rendre compte de toutes les phrases «modales» attestées comme vraies par la tradition.

Il a démontré la non-contradiction des propositions suivantes dans son système:

(M1) $CNMpNp$,

(M2) $CpNMNp$,

(M3) $KMpMNp$ pour certains p , c'est-à-dire $\Sigma p KMpMNp$,

et proposé les interprétations suivantes:

NM – «il n'est pas possible que»;

NMN – «il est nécessaire que»;

MN – «il est par hasard que».

En 1921, Tarski propose la définition de M dans le système $\mathfrak{I}_3$ comme $Mp = CNpp$, le troisième foncteur n'ayant pas de signification modale $Ip = KMpNLp$.

p	Mp	Lp	Ip
0	0	0	0
1/2	1	0	1
1	1	1	0

Lukasiewicz abandonne les travaux sur les systèmes modaux en 1953 quand il publie son système modal quadrivalent.

Généralisation de $\mathfrak{I}_3$ en $\mathfrak{I}_n$

Termes indéfinis:

p, q, r, s ...

variables propositionnelles

N, C,

négation, implication.

Ces logiques utilisent les valeurs 0 et 1 et les nombres rationnels. En général, pour la logique à $n = (3, 4, \dots)$ valeurs, l'ensemble des valeurs est le suivant:

$$A = \{0, \frac{1}{n-1}, \frac{2}{n-2}, \dots, \frac{n-2}{n-1}\}, B = \{1\}$$

$$Cpq = \min(1, 1 - p + q)$$

$$Np = 1 - p.$$

Termes définis:

$$Kpq = NCCNpNqNq = \min(p, q)$$

conjonction

$$Apq = CCpq = \max(p, q)$$

alternative (ou inclusive)

$$Epq = NCCNCpqNCqpNCqp =$$

$$\min(\min(1, 1 - p + q), \min(1, 1 - q + p))$$

équivalence.

Axiomes: système Ln $1 \leq n \leq \aleph_0$ de Lukasiewicz (Wajsberg)⁷

- W1. $CpCqp$
- W2. $CCpqCCqrCpr$
- W3. $CCNpNqCqp$
- W4. $CCCpNppp$.

Règles d'inférence:

Règle de substitution

Règle de détachement

modus ponens.

Lukasiewicz a proposé pour la logique $n = \aleph_0$, $\mathfrak{L}_{\aleph_0}$ les axiomes suivants étudiés par ailleurs par Wajsberg:

- L1. $CCpCqp$ pour la logique $\mathfrak{L}_{\aleph_0}$
- L2. $CCpqCCqrCpr$
- L3. $CCCpqqCCqpp$
- L4. $CCCpqCqCqpCqp$
- L5. $CCNpNqCqp$.

Logiques de Post

En 1921, indépendamment de Lukasiewicz, Emil L. Post définit [16] une classe de logiques à n valeurs (n est fini). Contrairement à Lukasiewicz, Post n'est pas intéressé par les significations et les interprétations des valeurs «ajoutées», mais il est fasciné par les relations et liaisons purement formelles entre ces valeurs. La définition de Post est basée sur les foncteurs des *Principia Mathematica* [Whitehead, Russell 1910] et sur la méthode matricielle de Peirce.

Définition de la logique P_n^k de Post

Termes indéfinis:

$p, q, r, s \dots$

N, A

variables propositionnelles
négation, disjonction.

⁷ Les axiomes avec la preuve proposée par Mordchaj Wajsberg 1931 [24].

Les valeurs de ces logiques sont les nombres naturels. En général, pour la logique à $n = (3, 4, \dots)$ valeurs l'ensemble des valeurs est:

$$A = \{0, 1, 2, \dots, n\} \quad B = \{0, 1, 2, \dots, k\}$$

où $k \leq n$. Ce sont les logiques n -valentes (finies) avec $1 \leq k \leq n$ valeurs distinctes.

$$Apq = \max(p, q)$$

$$Np = p+1 \text{ pour } p < n \text{ et } 1 \text{ pour } p = n \quad \text{négation cyclique.}$$

Termes définis:

$$Kpq = NANpNq$$

$$Cpq = ANpq$$

$$Epq = KCpqCqp$$

Axiomes:

n'existent pas!

La différence principale, par rapport à la logique de $\mathfrak{B}_n$, (exception des termes indéfinis, de l'interprétation sémantique, de l'axiomatisation, etc.), est que Post traite la négation d'une autre manière que Lukasiewicz. Il n'a attaché aucune signification à la troisième valeur. Sa négation est «cyclique» c'est-à-dire qu'il a rejeté la loi de la double négation qui est satisfaite dans la logique P_2^1 . Grâce à cette négation, le système de Post est complet sans avoir à lui ajouter aucun foncteur nouveau. Quant à la disjonction, elle est restée intuitivement claire et tout à fait «classique». Notre attention se porte aussi sur le choix des valeurs distinctes. Dans la logique classique ainsi que dans le système de Lukasiewicz, l'ensemble des valeurs distinctes est élémentaire, et la définition de la tautologie est exactement la même en logique classique que chez Lukasiewicz. Par contre chez Post, on a k -valeurs distinctes; ce qui change de la définition traditionnelle de la notion de tautologie et, par conséquent, empêche l'axiomatisation des logiques de Post ($n > 2$). Ces logiques sont aisément définies à l'aide du mécanisme des matrices et de l'algèbre. J'ajouterai que grâce au calcul matriciel, on a pu démontrer la complétude des logiques de Post.

On donne l'exemple des tables de la vérité en présentant le système P_3^2 de Post:

A	1	2	3	N	Np	K	1	2	3	C	1	2	3	E	1	2	3
1	1	1	1		2		3	3	2		1	2	2		3	3	3
2	1	2	2		3		3	1	2		1	2	3		3	1	2
3	1	2	3		1		2	2	2		1	1	1		3	2	3

La définition de la négation appliquée par Post, appelée aussi «rotation» nous permet de redéfinir un certain nombre de lois et principes dans leurs formes purement multivalentes.

Par exemple la loi *tertium non datur* $ApNp$, qui est une tautologie dans la logique bivalente, mais pas dans les systèmes de Lukasiewicz et de Post, peut être définie dans la logique P_3^1 ainsi:

$$ApANpNNp$$

et devenir la tautologie. De même, la loi de la double négation peut devenir loi de la n-uple négation dans le système P_3^1 . Elle prend la forme:

$$EpNNNp.$$

Ces quelques points différencient les systèmes de Post des autres systèmes multivalents et rendent cette logique intéressante également pour les sciences humaines – en particulier pour l'ontologie et l'épistémologie –, car elle leur fournit un outil pour explorer les autres «mondes» possibles.

Généralisation des systèmes finis de Post au système infini $P_{\aleph_0}$

Emil Post n'a pas donné de définition des systèmes infinis, mais on peut l'imaginer [17] comme l'extension du système P_k^1 au système $P_{\aleph_0}$ en changeant quelques éléments.

Les valeurs de la vérité sont définies ainsi:

$$1, \frac{1}{2}, \frac{1}{4}, \dots, \left(\frac{1}{2}\right)^k, \dots, 0.$$

Pour la négation on prend la formule suivante:

$$Np = 1 \text{ pour } p = 0 \text{ et } \frac{1}{2} \times p \text{ pour } p \neq 0.$$

Cette logique $P_{\frac{1}{2}}$ est très intéressante parce qu'elle ne contient aucune tautologie mais seulement des formules bien formées.

On peut conclure que des «manipulations» purement formelles, seule motivation de Post, conduisent aux logiques multivalentes avec un certain nombre d'idées totalement différentes de celles de Lukasiewicz. En effet, elles ne sont pas fondées sur une interprétation sémantique.

Logique de Sobocinski

En 1936, Boleslaw Sobocinski, membre de l'École polonaise de Lukasiewicz, présente une autre logique intéressante. Cette logique n -valentes est différente du courant dominant alors à l'École; elle est basée entre autres sur la négation cyclique.

Définition des logiques de Sobocinski

Termes indéfinis:

p, q, r, s

variables propositionnelles

$H, C,$

négation, implication.

Les valeurs de ces logiques sont les nombres naturels. En général, l'ensemble des valeurs de vérité, pour la logique à n -valeurs, est le suivant:

$$A = \{1, 2, \dots, n\} \quad B = \{2, \dots, n\}.$$

Il s'agit des logiques n -valentes possédant $n-1$ valeurs distinctes.

L'implication et la négation s'écrivent:

$$\begin{aligned} Cpq &= q, \text{ pour } p \neq q \text{ et } Cpq = n \text{ pour } p = q \\ Hp &= p + 1 \text{ pour } p < n \text{ et } Hp = 1 \text{ pour } p = n. \end{aligned}$$

Il s'agit de formules plus proches de l'idée de Post que de celle de Lukasiewicz. Pour exemple, on présente les tables de vérité de la logique S_3 ainsi:

Cpq	1	2	3	Hp
1	3	2	3	2
2	1	3	3	3
3	1	2	3	1

Les logiques de Sobocinski sont des systèmes complets et axiomatisés.

Axiomes de S_3 :

1. CCpCpq
2. CrsCtCCspCqp
3. CCHqCHHq
4. HpCHHpCpq
5. CCHHpHHqHHCpq
6. CHHCpqCHHpHHq.

Les différences par rapport à Post sont grandes, il suffit d'examiner la définition de l'implication pour s'en rendre compte.

Cette «voix polonaise», dans la ligne de Post, est un exemple des polémiques et du dynamisme qui régnaient alors. Ces vives polémiques portaient davantage sur des problèmes d'axiomatisation, de complétude que sur la notion d'interprétation. Jerzy Slupecki proposera en 1939 une synthèse de ces réflexions (voir ci-après).

Logiques de Kleene et Bochvar

En 1938, D.A. Bochvar et S.C. Kleene proposent d'autres logiques trivalentes. Ces logiques sont à la fois «proches» et «éloignées» des systèmes classiques et servent, selon les intentions des auteurs, à la résolution de problèmes théoriques en mathématiques. Elles font partie des logiques «non lukasiewiczziennes».

Logique de Kleene

S.C. Kleene, inspiré par la recherche sur les fondements de la mathématique, construit en 1938 [2] une logique qui permet d'analyser les prédicats partiellement définis. À titre d'exemple, citons le prédicat donné par l'équivalence suivante:

$$P(x) \text{ si et seulement si que } 1 \leq \frac{1}{x} \leq \neq 2,$$

on peut remarquer que $P(a)$ est vraie si $\frac{1}{2} \leq a \leq 1$, indéfinie ou non déterminée si $a = 0$ et fausse dans les autres cas.

Kleene est intéressé par ce type de propositions dont la valeur de vérité n'est pas déterminée par des algorithmes ou n'est pas importante dans l'application examinée.

Termes indéfinis:

$p, q, r, s \dots$

variables propositionnelles

N, K, A, C, E

négation, conjonction, disjonction,
implication et équivalence

T - «vrai»; I - «indéfini», F - «faux».

Foncteurs forts⁸

Kpq	T	I	F	Np	Apq	T	I	F	Cpq	T	I	F	Epq	T	I	F
T	T	I	F	F		T	T	T		T	I	F		T	I	F
I	I	I	F	I		T	I	I		T	I	I		I	I	I
F	F	F	F	T		T	I	F		T	T	T		F	I	T

On voit que le comportement des foncteurs des valeurs «vrai» (T) et «faux» (F) est le même que chez Lukasiewicz pour la négation, la conjonction et la disjonction. La différence réside dans la méthode d'interprétation de la troisième valeur et dans la définition de l'implication.

Un trait caractéristique de cette logique est qu'elle est non tautologique, malgré la valeur distincte (T). Cette situation est curieuse parce qu'on a l'extension de la logique classique avec un certain nombre des équivalences classiques, par exemple $ECpqANp$ ou $EApqCNpq$ mais les expressions telles Cpp et Epp

⁸ La notion de foncteur fort et faible est donnée par Kleene en 1952 [3].

n'apparaissent pas comme des tautologies. En 1952, Kleene dans son *Introduction aux Métamathématiques* [3] introduit les notions de foncteurs forts cités ci-dessus et faibles (ci-après).

Kpq	T	I	F	Np	Apq	T	I	F	Cpq	T	I	F	Epq	T	I	F
T	T	I	F	F		T	I	T		T	I	F		T	I	F
I	I	I	I	I		I	I	I		I	I	I		I	I	I
F	F	I	F	T		T	I	F		T	I	T		F	I	T

Ses motivations sont issues de l'arithmétique. Les nouveaux foncteurs servent à décrire des fonctions propositionnelles récursives où à chaque étape ils doivent être définis, l'indéfinition rend ainsi le calcul indéfini. C'est pour cela que tout foncteur appliqué à «I» donne «I». Il s'agit d'une réflexion de Bochvar.

Logique de Bochvar

En 1939, Bochvar propose [1] d'interpréter la troisième valeur «I» comme «indécidé». Par conséquent, dans la conjonction K, tous énoncés «indécidés», composés avec n'importe quelle valeur, donnent aussi comme résultat «indécidé». Cette idée est totalement différente de celle de Lukasiewicz. Cette troisième valeur ne peut pas être interprétée comme quelque chose d'«intermédiaire» entre la vérité et la fausseté, mais bien comme quelque chose de «paradoxale» ou «sans signification». Bochvar agit ainsi pour tenter de résoudre le paradoxe sémantique.

Termes indéfinis:

p, q, r, s, ...
N K

variables propositionnelle
négation, conjonction.

Il y a trois valeurs dans le système de Bochvar: T – «vrai», I – «indécidé», F – «faux».

Kpq	T	I	F	Np
T	T	I	F	F
I	I	I	I	I
F	F	I	F	T

Termes définis:

$Apq = NKNpNq$

$Cpq = NKpNq$

$E pq = KCpqCqp$.

Cette logique se caractérise par ses deux «niveaux». Le premier est tout à fait équiforme à la logique faible de Kleene et le deuxième métalogique est basé sur l'assertion externe définie comme $Ap =$ «il est vrai que p » selon la table suivante:

p	Assertion interne p	Assertion externe Ap
T	T	T
I	I	F
F	F	F

Suite à cette définition, Bochvar a défini les foncteurs externes de la manière suivante:

Forme interne	Forme externe
Np	$N_x p = NAp$
Kpq	$K_x pq = KApAq$
Apq	$A_x pq = AApAq$
Cpq	$C_x pq = CApAq$
E pq	$E_x pq = E ApAq$

La négation externe $N_x p$ – « p est faux»; la conjonction $K_x pq$ – « p est vrai et q est vrai»; l'alternative $A_x pq$ – « p est vrai ou q est vrai» et l'implication $C_x pq$ – «si p est vrai, alors q est vrai». Les valeurs de vérité sont données ci-dessous.

K_xpq	T	I	F	N_xp	A_xpq	T	I	F	C_xpq	T	I	F	E_xpq	T	I	F
T	T	F	F	F		T	T	T		T	F	F		T	F	F
I	F	F	F	F		T	F	F		T	T	T		F	T	T
F	F	F	F	T		T	F	F		T	T	T		F	T	T

Par conséquent, la logique interne ne contient pas de tautologie tout comme celle de Kleene et est une logique classique. Il existe une version de la logique faible, L_3^w , de Lukasiewicz, définie à la manière de Bochvar avec l'assertion W à la place de A .

En 1956, le logicien chinois Moh Swah-kwei propose une interprétation de la troisième valeur du système de Bochvar comme «paradoxal» et l'attribue aux propositions du type «cet énoncé est faux». Il semble que l'interprétation «sans signification» serait plus adéquate. Le système de Bochvar a trouvé sa place dans les travaux concernant les paradoxes sémantiques des théories.

Logique de Slupecki

La logique multivalente la plus générale est proposée par Jerzy Slupecki en 1939 [2], qui avait auparavant défini avec succès le foncteur «T» rendant ainsi complets les systèmes de Lukasiewicz.

Définition des logiques de Slupecki

L'idée de Slupecki est d'introduire des foncteurs généraux tels C , R , et S . Agissant ainsi, il définit la classe la plus générale des logiques multivalentes finies. Pour désigner ces logiques, il a utilisé les symboles L_n^k où $n = 2, 3, 4, \dots$ et $k < n$, parce qu'ils sont caractérisés par deux paramètres: le nombre de valeurs (n) et nombre des valeurs distinctes (k).

Termes indéfinis:

$p, q, r, s \dots$
 $C, R, S,$

variables propositionnelles

$Cpq = q$ pour $1 \leq p \leq k$, et 1 pour $k + 1 \leq p \leq n$

$Rp = p+1$ si $p = n$, alors $Rp = 1$

$Sp = p$ pour $3 \leq p \leq n$ $Sp = 1$ si $p = 2$, $Sp = 2$ si $p = 1$.

Pour k valeurs distinctes, le foncteur Cpq prend, comme résultat d'application, la valeur logique de la proposition q ; pour la suite, où $p > k$ jusqu'à n , prend la valeur distincte. Ce foncteur est une sorte d'implication généralisée. Pour le «R», Slupecki définit une notation postienne et pour le «S», il y a changement de la position des deux premières valeurs, la suite restant égale à «p».

Cpq	1 2 3 ... n	Rp	Sp
1*	1 2 3 ... n	2	2
2*	1 2 3 ... n	3	1
3*	1 2 3 ... n	4	3
...	...	...	...
k^*	1 2 3 ... n	$k + 1$	k
$k + 1$	1 1 1 ... 1	$k + 2$	$k + 1$
...	...	...	...
$n - 1$	1 1 1 ... 1	n	$n - 1$
n	1 1 1 ... 1	1	n

L'astérisque signifie que la valeur est distincte.

Ces démarches permettent à Slupecki d'inclure dans son système toutes les logiques à n -valeurs possibles, à la condition que n soit un nombre fini et de démontrer que le système général est complet et axiomatisé pour toutes les valeurs n .

Les résultats de Slupecki closent la discussion sur le problème de la construction des logiques multivalentes finies, mais ne donnent malheureusement pas de réponse sur les interprétations et les applications possibles de cette sorte de logique.

Conclusion

L'objectif de cet exposé était de présenter les logiques multivalentes dans leur dimension canonique de Lukasiewicz et Post à Slupecki qui clôt la question de l'axiomatisation et de la com-

plétude des logiques à un nombre fini de valeurs. Dans ce parcours, quelques événements historiques liés à la multivalence en logique ont été esquissés de manière sommaire. Cet exposé est destiné avant tout à préparer les étudiants à explorer les travaux consacrés aux logiques multivalentes. C'est la raison pour laquelle nous n'avons pas abordé les logiques conçues pour des applications spécifiques: tels les systèmes de Reichenbach servant à analyser des événements de la mécanique quantique, les systèmes probabilistes et possibilistes, utilisés dans le domaine de l'intelligence artificielle. Pour la même raison, nous n'avons pas traité des logiques modales de Lewis ou de la logique floue de Zadeh. Les problèmes de la logique multivalente des prédicats, logique à la fois difficile et intéressante, n'ont également pas été abordés. Des informations utiles sur ces différentes logiques peuvent être trouvées dans la monographie de Rescher [14].

Nous souhaitons que ce qui a été proposé rendra plus facile et agréable la découverte et la lecture des textes fondateurs.

*Département de philosophie
Université de Lausanne
1015 Lausanne*

Bibliographie

- [1] KOTARBINSKI T. (1964). *Wykłady z dziejow logiki* [Leçons sur l'histoire de la logique]. Paris: PUF.
- [2] LUKASIEWICZ J. (1920). Compte-rendu de sa propre conférence: *O pojeciu mozliwosci* [Du concept du possible], prononcée le 5.6.1920 à la 206^e séance de la Société Polonaise de Philosophie à Lwów; publié dans *Ruch filozoficzny* V^e année, n° 9.
- [3] LUKASIEWICZ J. (1920). Résumé de la conférence *O logice trojwartosciowej* [De la logique trivalente], prononcée le 5.6.1920 à la 207^e séance de la Société Polonaise de Philo-

- sophie à Lwów; publié dans *Ruch filozoficzny* V^e année, n° 9.
- [4] LUKASIEWICZ J. (1929). *Elementy logiki matematycznej* [Eléments de logique mathématique] Skrypt. (2 wyd.-Warszawa 1958, PWN).
- [5] LUKASIEWICZ J. (1930). *Uwagi filozoficzne o wielowrtosciowych systemach rachunku zdań* [Remarques philosophiques sur les systèmes multivalents du calcul propositionnel], en allemand, dans *Sprawozdania z posiedzen Towarzystwa Naukowego Warszawskiego* [Comptes-rendus des séances de la Société scientifique de Varsovie], Section III, XXIII^e année.
- [6] LUKASIEWICZ J., TARSKI A. (1930). *Untersuchungen über den Aussagenkalkul* *Sprawozdania z posiedzen Towarzystwa Naukowego Warszawskiego* [Comptes-rendus des séances de la Société scientifique de Varsovie], Section III, XXIII^e année.
- [7] LUKASIEWICZ J. (1931). *Uwagi o aksjomacie Nicoda i [dedukcji uogólniajacej]*. [Remarques sur l'axiome de Nicod] Księga pamiatkowa Polskiego Towarzystwa Filozoficznego. Lwów.
- [8] LUKASIEWICZ J. (1939),. *Der Aqiuvalenzkalkül. Collectanea logica* 1, 145-169.
- [9] LUKASIEWICZ J. (1961). *Z zagadnien Logiki i Filozofii*. Ed. J. Slupecki. Warszawa: PWN.
- [10] LUKASIEWICZ J. (1970). *Selected Works*. Amsterdam: North-Holland.
- [11] MALINOWSKI G. (1990). *Logiki wielowartosciowe*. Warszawa: PWN.
- [12] MC CALL S. (ed.) (1967). *Polish Logic 1920-1939*. Oxford: Clarendon.
- [13] POST E.L. (1921). Introduction to a general theory of elementary propositions, *American Journal of Mathematics* VLII.
- [14] RESCHER N. (1969). *Many-Valued Logic*. New York: McGraw-Hill.
- [15] SLUPECKI J. (1936). *Der volle dreiwertige Aussagenkalkul, Towarzystwa Naukowego Warszawskiego* [Comptes-rendus

des séances de la Société scientifique de Varsovie], Section III, XXIX^e année.

- [16] SOBOCINSKI B. (1936). *Aksjomatyzacja pewnych wielowartosciowych systemow teorii dedukcji* [Axiomatization of certain many-valued systems of the theory of deduction], *Roczniki prac naukowych zrzeszenia asystentow Uniwersytetu J. Pilsudskiego w Warszawie*, vol.1, 399-419 [Reviewed by Tarski in *The Journal of Symbolic Logic* 2 (1937), 93.
- [17] WOLF R.G. (1977). *A Survey of Many-Valued Logic (1966-1974)*. Dunn: Estein.

ANNEXE

Bibliographie complète de Lukasiewicz

(Source: *Studia Logica*, V, 1951 avec modifications du prof. A. Korcik).

Abréviations:

P.F. Przegląd Filozoficzny; † dans Polish Logic [12]

R.F. Ruch Filozoficzny; ‡ dans Collected Works [10].

1. Streszczenie: Vierteljahrschrift für wissenschaftliche Philosophie 1899, n° 3-4. P. F. V (1902), 232-236.
2. O indukcji jako inwersji dedukcji. P. F. VI (1903), 9-24; 138-152.
3. Recenzja: T. Mianowski: O tzw. pojeciach wrodzonych u Locke'a i Leibniza. P. F. VII (1904), 94-95.
4. O stosunkach logicznych. [On logical relations] P. F. VII (1904), 245.
5. Teza Husserla o stosunku logiki do psychologii. P. F. VII (1904), 476-477.
6. Z psychologii porownywania. P. F. VIII (1905), 290-291.
7. O dwóch rodzajach wniosków indukcyjnych. P. F. IX (1906), 83-84.

8. Analiza i konstrukcja pojecia przyczyny. [Analyse et construction du concept de cause] P. F. IX (1906), 105-179.
9. Tezy Höflera w sprawie przedstawien i sądów geometrycznych. P. F. IX (1906), 451-452.
10. Co poczac z pojeciem nieskonczonosci? P. F. X (1907), 135-137.
11. Recenzja: H. Struve. Die polnische Philosophie der letzten zehn. Jahre (1894-1904). To samo w przekladzie polskim K. Krola. P. F. X (1901), 336-346.
12. O wnioskowaniu indukcyjnym. P. F. X (1907), 474-475.
13. Logika a psychologia. P. F. X (1907), 489-491.
14. Pragmatyzm, nowa nazwa pewnych starych kierunków myslenia. P. F. XI (1908), 341-342.
15. Sprawozdanie z dwóch prac Stumpfa. P. F. XI (1908), 342-343.
16. Zagajenie pogadanki na temat rozprawy M. Borowskiego: Krytyka pojecia zwiazku przyczynowego. P. F. XI (1908), pp.343.
17. Zadania i znaczenie ogólnej teorii stosunków. P. F. XI (1908), 344-347.
18. O prawdopodobienstwie wniosków indukcyjnych. P. F. XII (1909), 209-210.
19. O poglądach filozoficznych Meinonga. P. F. XII (1909), 559.
20. O zasadzie wylaczonego srodka. P. F. XIII (1910), 372-373.
21. Über den Satz von Widerspruch bei Aristoteles. Bulletin international de l'Académie des Sciences de Cracovie, Classe de Philosophie (1910), 15-38.
22. O zasadzie sprzeczności u Arystotelesa. Studium krytyczne. [Sur le principe de contradiction dans l'oeuvre d'Aristote] Krakow (1910).
23. Recenzja: Wl. Tatarkiewicz. Die Disposition der aristotelischen Prinzipien. R. F. I (1911), 20-21.
24. O wartosciach logicznych. R. F. I (1911), 52.
25. O rodzajach rozumowania. Wstep do teorii stosunków. R. F. I (1911), 78.

26. Recenzja: P. Natorp. *Die logischen Grundlagen der exakten Wissenschaften*. R. F. I (1911), 101-112.
27. Recenzja: H. Struve. *Historia logiki jako teorii poznania w Polsce*. Wyd. drugie. R. F. I (1911), 115-117.
28. O potrzebie zalozenia instytutu metodologicznego. R. F. II (1912), 17-19.
29. Recenzja: Wl. Bieganski. *Czym jest logika?* R. F. II (1912), 145.
30. ↑ O tworczości w nauce. *Księga pamiątkowa ku uczczeniu 250 rocznicy założenia Uniwersytetu Lwowskiego*. Lwów 1912, 1-15.
31. *Nowa teoria prawdopodobieństwa*. R. F. III (1913), p. 22.
32. Recenzja: J. Kleiner. *Zygmunt Krasinski. Dzieje myśli*. R. F. III (1913), 109-111.
33. *Logiczne podstawy rachunku prawdopodobieństwa*. [Sur les fondements logiques du calcul des probabilités] Sprawozdania PAU, 1913, 5-7.
34. ↑ *Die logischen Grundlagen der Wahrscheinlichkeitsrechnung*. Kraków (1913), 75.
35. W sprawie odwracalności stosunku racji i następstwa. P. F. XVI (1913), 298-314.
36. *Rozumowanie a rzeczywistość*. R. F. IV (1914), 54.
37. *O nauce. Poradnik dla samouków*. Wyd. nowe. T. I. 1915, XV-XXXIX. Przedruk: Lwów 1934, 1936, 40.
38. *O nauce i filozofii*. P. F. XVIII (1915), 190-196.
39. ↑ O pojęciu wielkości [On the concept of magnitude] P. F. XIX (1916), 1-70.
40. ↑ Treść wykładu pozegegnalnego wygłoszonego w auli Uniwersytetu Warszawskiego 7 marca 1918. Warszawa 1918.
41. † O pojęciu możliwości. R. F. V (1919/20), 169-170.
42. †↑ O logice trójwartościowej. [Sur la logique trivalente] R. F. V (1920), 170-171.
43. ↑ *Logika dwuwartościowa*. P. F. XXIII (1921), 189-205.
44. *O przedmiocie logiki*. R. F. VI (1921), 26.
45. *Zagadnienia prawdy*. *Księga pamiątkowa XI Zjazdu lekarzy i przyrodników polskich*. 1922, 84-85, 87.

46. ↑ Interpretacja liczbowa teorii zdan. R. F. VII (1922/23), 92-93.
47. Recenzja: Jan Sleszynski. O logice tradycyjnej. R. F. VIII (1923), 107-108.
48. Kant i filozofia nowożytna. Wiadomosci Literackie I (1924), 19.
49. Dlaczego nie zadowala nas logika filozoficzna? R. F. IX (1925), 25.
50. O pewnym sposobie pojmowania teorii dedukcji. P. F. XXVIII (1925), 134-136.
51. Démonstration de la compatibilité des axiomes de la théorie de la déduction. Annales de la Société Polonaise de Mathématique 3 (1925), 149.
52. Sprawozdanie z dzialalnosci Uniwersytetu Warszawskiego za r. ak.1922/23. Warszawa 1925.
53. Z najnowszej niemieckiej literatury logicznej. R. F. X (1926/27), 197-198.
54. O logice stoikow. P. F. XXX (1927), 278-279.
55. O metodzie w filozofii. P. F. XXXI (1928), 3-5.
56. O pracy Fr. Weidauera: Zur Syllogistik. R. F. XI (1928), 178.
57. Rola definicji w systemach dedukcyjnych. R. F. XI (1928/29), 164.
58. Odefinicjach w teorii dedukcji. R. F. XI (1928/29), 177-178.
59. Wrazenia z VI Miedzynarodowego Zjazdu Filozoficznego. H. F. XI (1928/29), 1-5.
60. O znaczeniu i potrzebach logiki matematycznej. Nauka Polska 10 (1929), 604-620.
60. Elementy logiki matematycznej. Skrypt. 1929 (2 wyd.-Warszawa 1958, PWN).
61. (z A. Tarskim) Untersuchungen Über den Aussagenkalkül. Comptes-rendus de la Société des Sciences et des Lettres de Varsovie, Cl. III, 23 (1930), 1-21.
62. †↑ Philosophische Bemerkungen zu mehrwertigen Systemen des Aussagenkalküls. Comptes-rendus de la Société des Sciences et des Lettres de Varsovie, Cl. III, 23 (1930), 51-77.

63. ↓ Uwagi o aksjomacie Nicoda i [dedukcji uogólniającej]. [Remarques sur l'axiome de Nicod] Księga pamiątkowa Polskiego Towarzystwa Filozoficznego. Lwów 1931.
64. Ein Vollständigkeitsbeweis des zweiwertigen Aussagenkalküls. Comptes-rendus de la Société des Sciences et des Lettres de Varsovie, C1. III, 24 (1931), 153-183.
65. Z dziejów logiki starożytnej. R. F. XIII (1932-1936), 46.
66. †‡ Z historii logiki zdań. P. F. XXXVII (1934), 417-437.
67. Znaczenie analizy logicznej dla poznania. P. F. XXXVII (1934), 369-377.
68. Zur Geschichte der Aussagenlogik. Erkenntnis 5 (1935-1935), 111-131.
69. Zur vollen Aussagenlogik. Erkenntnis 5 (1935-1936), 176.
70. ↓ Logistyka a filozofia. P. F. XXXIX (1936), 115-131.
71. Bedeutung der logischen Analyse für die Erkenntnis. Actes du VIII Congrès International de Philosophie. Prague (1936), 75-84.
72. Co dała filozofii współczesna logika matematyczna? P. F. XXXIX (1936), 325-326.
73. ↓ W obronie logistyki. Myśl katolicka wobec logiki współczesnej. Studia Gnesnensia, nr 15 (1937), 22.
74. En défense de la logique. La pensée catholique et la logique moderne. Compte-rendu de la session spéciale tenue le 26.9.1936 pendant le IIe Congrès Polonais de Philosophie. Wydawnictwa Wydziału Teologicznego U.I., seria I, nr 2 (1937), 7-13.
75. Kartezjusz. Kwartalnik Filozoficzny XV (1938), 123-128.
76. O sylogistyce Arystotelesa. Sprawozdania PAU, 44, (1939), 220-227.
77. Der Äquivalenzkalkül. Collectanea logica 1 (1939), 145-169.
78. ↓ Die Logik und das Grundlagenproblem. Les entretiens de Zürich sur les fondements et la méthode des sciences mathématiques 6-9, XII (1938), Zurich 1 41, 82-100.
79. ↓ The shortest axiom of the implicational calculus of propositions. Proceeding of the Royal Irish Academy, Sect. A, 52 (1948), 25-33.

80. (↑) W sprawie aksjomatyki implikacyjnego rachunku zdań. *Annales de la Société Polonaise de Mathématique* 22 (1956), 87-92.
81. O zasadzie najmniejszej liczby. (Streszczenie referatu nadesłanego na Zjazd). Sprawozdanie z V Zjazdu matematyków polskich w Krakowie w dniach 29-31 maja 1947 r. oraz z Akademii poświęconej uczczeniu pamięci prof. Stanisława Zaremby. Dodatek do *Rocznika Polskiego Towarzystwa Matematycznego*. T. XXI, 1948-1949, Kraków 1950, 28-29.
82. ↑ On variable functors of propositional arguments. *Proceedings of the Royal Irish Academy, Sect. t A*, 54, (1951), 25-35.
83. Aristotle's syllogistic from the standpoint of modern formal logic. Oxford 1951.
84. ↑ On the intuitionistic theory of deduction. *Indagationes Mathematicae. Koninklijke Nederlandse Akademie van Wetenschappen, Proceedings, Series A* (1952), n° 3, 202-212.
85. Comment on K. J. Cohen's remark. *Indagationes Mathematicae. Koninklijke Nederlandse Akademie van Wetenschappen, Proceedings, Series A* (1953), n° 2, 113.
86. ↑ Sur la formalisation des théories mathématiques. *Colloques internationaux du Centre National de la Recherche Scientifique, XXXVI: Les méthodes formelles en axiomatique*. Paris 1953, 11-19.
87. ↑ A system of modal logic. *The Journal of Computing Systems*, vol. 1, n° 3, (1953), 111-149.
88. ↑ A system of modal logic. *Actes du XV^e Congrès International de Philosophie*, XIV, 72-78.
89. ↑ Arithmetic and modal logic. *The Journal of Computing System*, vol. 1, n° 4, (1954), 213- 219.
90. On a controversial problem of Aristotle's modal syllogistic. *Dominican Studies* 1954 (7), 114-128.
91. †↑ O determinizmie [On determinism] [9].

LOGIQUE COMBINATOIRE, TYPES, PREUVES ET LANGAGE NATUREL

Jean-Pierre DESCLÉS

Notre compréhension intime de la logique combinatoire de Curry nous la fait penser comme une *logique des processus opératoires abstraits*, ces processus opératoires étant, par ailleurs, éventuellement implantables sur les supports physiques d'une machine informatique et exécutables par les organes de calcul de cette machine. La logique combinatoire thématise les *modes intrinsèques de composition et de transformation* des opérateurs, analysés comme des opérateurs abstraits, appelés *combineurs*. Ces modes de composition et de transformation sont définis *indépendamment de toute interprétation* et donc indépendamment des domaines d'interprétation comme par exemple l'arithmétique, la logique des propositions et des prédicats, la construction de programmes, la linguistique, la psychologie ou ... la philosophie. Par son grand pouvoir d'abstraction et d'expressivité, la logique combinatoire tend à acquérir un rôle unificateur puisqu'elle nous permet de jeter des ponts assez solides entre des domaines qui sont *a priori* aussi éloignés l'un de l'autre que les fonctions récursives, l'analyse des paradoxes, les types intuitionnistes de Martin-Löf, la programmation fonctionnelle, la théorie des preuves, la logique linéaire de Girard, les domaines réflexifs de Scott, la théorie des catégories cartésiennes fermées et des $T[\Sigma]$ -algèbres, la théorie des topoï géométriques ... ou les théories des opérateurs linguistiques de Z.S. Harris ou de S.K. Shaumyan. A propos des sciences du langage, nous pouvons donc légitimement nous poser les questions suivantes: les langues naturelles, en tant que systèmes de représentations symboliques entretiennent-elles quelque rapport sémiotique avec d'autres systèmes sémiotiques que sont certains

langages logiques et certains langages de programmation? Les analyses linguistiques, en particulier les analyses syntaxiques et les constructions des interprétations sémantiques des énoncés, ont-elles quelque chose à voir avec les preuves pensées comme des objets d'étude? Les structurations internes des langues naturelles ont-elles quelque rapport avec des structurations logiques et géométriques (théorie des topoï par exemple)?

Historiquement, la logique combinatoire a pris naissance avec les travaux de M. Schönfinkel¹ (1924) qui cherchait à construire une logique de la quantification (ou logique du premier ordre) fondée sur la seule notion primitive de fonction. La logique combinatoire a été développée essentiellement sous forme d'une «prélogique» par H. B. Curry (1930, 1958, 1972) qui cherchait d'une part, à formaliser l'effectivité opératoire sous forme de règles analogues à «modus ponens» et d'autre part, à étudier, par des procédés logiques, l'émergence des paradoxes du genre de celui de Russell. Rappelons que la règle de «modus ponens» correspond, dans la «déduction naturelle» de Gentzen, à la *règle d'élimination* du symbole d'implication « $\rightarrow$ »:

$$\begin{array}{c} \text{[e-}\rightarrow\text{]} \quad \frac{A \rightarrow B, A}{B} \end{array}$$

Cette règle signifie que des deux prémisses « $A \rightarrow B$ » et « A », on peut en inférer la conclusion « B ». Ce genre de règle d'inférence sert de modèle aux règles formelles. Un autre règle relative à la flèche « $\rightarrow$ » correspond à la *règle d'introduction* qui se présente comme suit:

$$\begin{array}{c} \text{[i-}\rightarrow\text{]} \quad \frac{\begin{array}{c} \text{[A]} \\ \cdot \\ B \end{array}}{A \rightarrow B} \end{array}$$

1 Voir la traduction en français avec commentaires dans *Mathématiques, Informatique et Sciences Humaines* 112, 1990, 5-26.

Cette règle signifie que si d'une hypothèse A on peut déduire B, alors on peut se «libérer de l'hypothèse A» en introduisant, mais à un niveau hiérarchique supérieur, l'expression «A->B». Cette règle permet de relier la relation métalinguistique de déduction entre A et B à l'expression linguistique «A->B» interne au langage; c'est une version du théorème «classique» de la déduction.

M. Schönfinkel et H.B. Curry ont introduit des opérateurs abstraits, appelés «combinateurs», qui construisent dynamiquement des opérateurs complexes à partir d'opérateurs plus élémentaires. Indépendamment, A. Church (1944) est l'inventeur d'un calcul, le λ -calcul, formé d'expressions symboliques qui représentent les mécanismes opératoires de procédures associées à des opérateurs ou à des fonctions appliquées à des opérands. La logique combinatoire tout comme le λ -calcul sont des systèmes applicatifs car fondés sur l'opération d'application, considérée comme une opération primitive qui construit un résultat à partir d'un opérateur et d'un opérande. D'autres langages applicatifs existent. Mentionnons, par exemple, les différents langages de programmation fonctionnelle comme ML. Dans l'application d'un opérateur à un opérande, l'opérateur peut être dans certains cas interprété comme une fonction ensembliste, l'opérande étant l'argument de la fonction et le résultat l'image de cette fonction. Cependant, il y a une différence de point de vue. Dans la conceptualisation ensembliste, opérateur – ou fonction –, opérande – ou argument –, résultat – ou image – sont *co-existentiels*, ils sont situés hors du temps, ils sont donc *atemporels*; de plus, étant donné une fonction ensembliste, on peut affirmer l'existence d'une image à partir d'un argument sans que nécessairement on puisse toujours savoir construire effectivement cette image. Dans la conceptualisation applicative, l'opérateur et l'opérande doivent être d'abord construits avant de pouvoir appliquer l'un à l'autre pour en déduire la construction effective du résultat; de plus, l'application est conçue comme une *opération dynamique* qui implique une *asymétrie temporelle* entre opérateur, opérande d'une part et résultat, d'autre part. Un système applicatif, fondé sur l'application, est ainsi un système à la fois *effectif*, puisque les objets manipulés sont tous constructibles, et *dynamique*, puisque l'application est une opé-

ration qui s'exécute dans une certaine temporalité. Aucune restriction n'est posée *a priori* sur l'opération d'application. En particulier, on peut appliquer un opérateur à lui-même sans sortir pour autant des systèmes applicatifs. Cependant, on peut restreindre l'expressivité du formalisme applicatif en ne considérant que différents types d'entités manipulées; dans ce cas, les systèmes deviennent typés et alors les opérateurs typés ne peuvent s'appliquer qu'à certains types d'opérandes. Dans un langage typé, un opérateur quelconque ne peut généralement pas s'appliquer à lui-même sans violer certaines contraintes sur les types. Cependant, on peut étudier les mécanismes d'auto-application et faire apparaître les raisons profondes qui conduisent à une situation paradoxale.

Énumérons quelques caractéristiques de la logique combinatoire

1) *La logique combinatoire est une logique «sans variables» (liées).* Or, comme on sait, la notion de «variable» est complexe, délicate à manier et ambiguë. Historiquement, Schönfinkel et Curry ont du reste cherché à construire une logique «sans variables» pour rendre plus effectifs les systèmes formels. Bien que le λ -calcul de Church et la logique combinatoire de Curry soient extensionnellement équivalents, les expressions de la logique combinatoire, en dépit des difficultés apparentes et superficielles, sont plus simples à manipuler que les expressions du λ -calcul qui doit faire l'exacte gestion des variables liées pour éviter, entre autres problèmes, les «télescopages des variables» (comme dans le calcul des prédicats) et les «effets de bord» (en programmation).

2) *La logique combinatoire permet de ramener toutes les manipulations à des règles effectives et explicites* aussi simples, dans leurs formes, que les règles de «modus ponens» (élimination de l'implication dans le calcul propositionnel) et d'abstraction (introduction de l'implication dans le calcul propositionnel). Toutes les «constantes logiques» (d'un côté, les combinateurs et d'un autre côté, les opérateurs illatifs de

connexion propositionnelle et de quantification des prédicats par exemple) sont introduites ou éliminées par des règles dont la forme est très exactement analogue aux règles d'introduction et d'élimination de la flèche « $\rightarrow$ ». On obtient ainsi des présentations «dans le style de Gentzen» de la logique combinatoire (voir, entre autres, le livre de Fitch). Ce genre de présentation permet une discussion sur la nature des règles de la logique, aboutissant aux généralisations obtenues dans la logique linéaire de J.Y Girard.

3) *La logique combinatoire est puissante* puisqu'elle permet d'exprimer, entre autres, les opérations de l'arithmétique élémentaire et la définition des fonctions récursives; elle permet donc une conceptualisation de l'effectif et du calculable, selon la thèse de Church. Les théorèmes de Gödel trouvent également une expression dans ce cadre applicatif.

4) *La logique combinatoire a un pouvoir fondationnel*, puisqu'elle permet une analyse logique des paradoxes, aussi bien du paradoxe de Russell, que celui du menteur ou de la construction impossible de la surjection de Cantor entre un ensemble et l'ensemble de ses parties. Du reste, Curry a construit son système pour «analyser logiquement» les conditions qui faisaient émerger les paradoxes; il a ainsi montré que la plupart des paradoxes apparaissaient en construisant un opérateur particulier (le «combinateur paradoxal») qui met en jeu l'itération d'une «diagonalisation». Par ailleurs, la logique combinatoire est un langage interprétatif pour les programmes informatiques (dans la sémantique dénotationnelle de Strachey-Scott). Par son modèle topologique du λ -calcul, D. Scott a également montré quelle était l'importance du λ -calcul et des opérateurs constructeurs de point fixe dans les études sémantiques des langages de programmation de haut niveau, où certaines fonctions sont définies récursivement par des équations.

5) *La logique combinatoire a un pouvoir expressif d'analyse sémiotique des écritures symboliques*, en particulier, elle permet de passer de la notion extrinsèque d'opération à celle intrinsèque d'opérateur (Desclés 1980). Plusieurs exemples d'analyses sémantiques de concepts peuvent être donnés; citons par exemple le prédicat complexe «de qui rien de plus grand ne peut être

pensé» qu'utilise Anselme de Canterbury (X^e siècle) dans son fameux «unum argumentum» du *Proslogion* (Desclés 1991).

6) *La logique combinatoire se trouve à un carrefour* de disciplines formelles comme les mathématiques (théorie du point fixe, catégories cartésiennes fermées et topologies de Grothendieck), la logique (théorie de la calculabilité effective, théorie des démonstrations; théorie des fondements); philosophie des systèmes formels; approche intuitionniste des mathématiques; informatique fondamentale (programmation fonctionnelle; sémantique des langages de programmation).

7) *La logique combinatoire et le λ -calcul ont trouvé et trouvent des terrains d'application dans des disciplines qui s'appuient sur des données empiriques.* Ces applications sont certes encore timides mais cependant bien réelles dans certains secteurs des sciences humaines: psychologie cognitive (par exemple: analyse par L. Frey d'opérations cognitives du groupement de Piaget); philosophie (analyse de concepts philosophiques), musicologie (analyse de partitions, comme celle du *Clavier bien tempéré* de Bach); économie (travaux de Dupuy). Depuis les premiers travaux sur les Grammaires Catégorielles, à la suite des réflexions logico-philosophiques de Husserl et de S. Lesniewski, et surtout depuis les recherches menées par S.K. Shaumyan sur la Grammaire Applicative (1965), puis de Geatch, de Montague, de Dowty, de Moortgat, de Steedman, ..., les formalismes applicatifs de la logique combinatoire et du λ -calcul (avec types) ont été explicitement mis en oeuvre, de façon significative, dans plusieurs théories linguistiques: Grammaires Catégorielles étendues (avec des systèmes inférentiels de types fonctionnels de Lambek et une interprétation sémantique dans le λ -calcul par Moortgat, Steedman), syntaxe catégorielle et sémantique dénotative de la Grammaire Universelle de Montague, Grammaire des opérateurs et opérands de Z. Harris (1983), Grammaire Applicative Universelle de S.K. Shaumyan (1977, 1987), Grammaire Applicative et Cognitive (Desclés 1990), ... D'autres linguistes, par exemple: A. Culioli ont également entrevu l'importance de la logique combinatoire pour la linguistique:

On ramènera toutes les opérations unaires de prédication (à l'exclusion des transformations de composition de lexis) à une application (...) on ira jusqu'au bout de l'analyse, en y adjoignant [en en dérivant] une théorie des prédicats. (...) le linguiste aura parfois des concepts-clés à portée de main, parfois l'élaboration sera lente (ainsi en est-il de l'utilisation de la topologie en linguistique ou encore de la logique combinatoire) ...

Les questions qui nous préoccupent depuis quelques années peuvent se formuler de la façon suivante:

- (i) Les langues naturelles étant des systèmes sémiotiques particuliers, peuvent-elles être pensées et organisées comme des langages applicatifs?
- (ii) Est-il licite *a priori* de faire appel, pour une analyse linguistique, aux formalismes applicatifs de préférence à d'autres formalismes (certains étant mieux connus et plus utilisés comme les systèmes fondés sur la concaténation et sur les mécanismes d'unification de traits ou ceux qui, depuis quelques années, relèvent des réseaux connexionnistes)?
- (iii) Les principales opérations linguistiques dégagées par les linguistes (comme les opérations de prédication, de détermination, de thématisation, d'énonciation, ...) peuvent-elles prendre une forme opératoire effective dans un langage applicatif?
- (iv) Quelles sont les limites des formalismes applicatifs pour la description adéquate et fine des mécanismes langagiers?
- (v) Quelles sont les relations entre la construction de la signification encodée par un énoncé et la théorie de la démonstration puisque certains travaux actuels explorent, depuis quelques années avec profit semble-t-il, cette voie?
- (vi) Quelles conséquences peut-on déduire de l'analyse applicative des langues pour un traitement informatique des langues (par la programmation fonctionnelle, de nature applicative), à la fois sur le plan théorique et pratique?
- (vii) Les langages applicatifs sont-ils utilisables pour représenter plus adéquatement que les systèmes actuels – par

exemple, les réseaux sémantiques ou les graphes conceptuels de Sowa ou encore les modèles de Kamp – les connaissances et les représentations cognitives?

- (viii) Les langages applicatifs sont-ils utilisables pour une formalisation, au moins partielle, des conceptualisations élaborées par les Grammaires Cognitives?
- (ix) Quelles propriétés des formalismes applicatifs font progresser notre connaissance sur le fonctionnement du langage?
- (x) Comment les langages applicatifs peuvent-ils nous aider à appréhender, représenter et manipuler des représentations cognitives hypothétiques qui sont non directement observables et par conséquent (re-)constructibles seulement par des processus abductifs explicites?

Quelques *points de vue sur les langues* peuvent justifier *a priori* l'utilisation du λ -calcul et de la logique combinatoire:

1. L'opération de base qui est constitutive des expressions linguistiques est l'application et non la concaténation.
2. Les unités des langues sont des unités opératives (ou fonctionnelles).
3. Les unités linguistiques sont des opérateurs de différents types.
4. La «currification» des opérateurs a une interprétation linguistique ainsi que, peut-être l'analogie de Curry-Howard.
5. La construction de l'interprétation fonctionnelle des énoncés peut être effectuée à partir, d'une part, une spécification sous forme d'une formule de types syntaxiques et, d'autre part, de preuves que cette formule est «cohérente», c'est-à-dire que cette formule est un théorème d'un système inférentiel sur les types.
6. L'introduction des opérateurs illatifs devrait permettre une meilleure analyse de la quantification dans les langues naturelles.

7. Les «prédicats complexes» sont construits à l'aide des combinateurs, ils sont essentiels pour l'analyse et la représentation sémantique d'un grand nombre d'énoncés.

Après avoir introduit quelques notations utiles, nous examinerons donc successivement chacun de ces différents points qui touchent au statut sémiotique des langues naturelles. Dès lors qu'elles sont comprises et acceptées après discussion, ces affirmations justifieraient, au moins à nos yeux, qu'un programme de recherche sur le langage puisse utiliser à bon escient la logique combinatoire en tant que système métalinguistique d'analyse des langues.

0. Notations

0.1. Application. L'application d'une expression Y, considérée comme un opérateur par rapport à une autre expression X, considérée comme un opérande de Y, construit un résultat Z. Nous désignons par «App(X, Y)» l'opération d'application. La relation entre cette opération «à exécuter» et le résultat Z obtenu «après exécution» est une réduction, notée: $\text{App}(Y, X) \beta \rightarrow Z$.

0.2. Présentation sémiotique de l'application. Le résultat de l'application de Y à X pourra être désigné par la simple juxtaposition «YX» de l'opérateur et de l'opérande. L'opération d'application est alors représentable par un arbre, dit «arbre applicatif»:

$$\begin{array}{ccc} & Y & X \\ \text{[App]} & \frac{}{} & \\ & YX & \end{array}$$

0.3. Expressions applicatives. L'ensemble des expressions applicatives est construit récursivement à l'aide de la règle:

si Y et X sont des expressions applicatives, alors (YX) est une expression applicative.

Lorsqu'il n'y aura pas ambiguïté, les parenthèses externes seront omises. Il est clair que les deux expressions applicatives « $X(YZ)$ » et « $(XY)Z$ » (qui est égale à « XYZ ») n'ont pas la même signification.

0.4. λ -expressions. Les λ -expressions sont des expressions applicatives particulières obtenues par un processus d'abstraction. L'ensemble des λ -expressions est défini comme suit:

- (i) si X et Y sont des λ -expressions, alors (YX) est une λ -expression;
- (ii) si X et Y sont des λ -expressions, et si x est une variable libre qui présente des occurrences dans Y , alors $\lambda X.Y$ est une λ -expression.

La λ -expression « $\lambda x.Y$ » est un opérateur qui a été obtenue par abstraction de « x » à partir de « Y ». L'application de l'opérateur $\lambda x.Y$ à un opérande Z construit le résultat, noté « $Y[x := Z]$ », obtenu en substituant Z à toutes les occurrences libres de x dans Y , ce que nous désignons par la β -réduction:

$$((\lambda x.Y) (Z)) \beta \rightarrow Y[x := Z].$$

Une présentation rigoureuse des λ -expressions et du λ -calcul associé consisterait à prendre certaines précautions sur la «gestion des abstractions» pour éviter les télescopages.

0.5. Expressions typées. Dans un système applicatif typé, les opérateurs, les opérandes et les résultats sont typés. Une entité X à laquelle est assigné un type α (soit par assignation initiale, soit par inférence) sera notée, en suivant la littérature actuelle: $[X : \alpha]$ ². Désignons par la notation infixée « $\alpha \rightarrow \beta$ » le type fonctionnel «de α vers β »³. L'opération d'application, dans un système applicatif typé, est représentée par l'arbre applicatif typé:

2 Une notation plus appropriée serait $[\alpha : X]$ ou: ' αX ' notation de Curry qui considère le type comme un opérateur externe ou internalisable dans le système.

3 La notion préfixée ' $F\alpha\beta$ ' de Curry est beaucoup plus commode mais la littérature actuelle ne l'a pas suivi sur ce point.

$$\begin{array}{c}
 [Y : \alpha \rightarrow \beta] \qquad [X : \alpha] \\
 \hline
 [\text{App}_t] \qquad [YX : \beta]
 \end{array}$$

Un opérande absolu qui ne peut jamais être opérateur sera considéré comme un opérateur zéro-aire. Toutes les entités (typées ou non) d'un système applicatif sont donc des opérateurs. Chaque opérateur agit successivement sur éventuellement plusieurs opérandes pour construire un résultat (par exemple en linguistique un énoncé). Lorsqu'un, deux, ... n opérandes sont nécessaires pour qu'un opérateur construise une unité du système, on dit que l'opérateur est unaire, binaire, ... n aire.

Une λ -expression typée est construite comme suit: soit x une variable libre de type α et Y une expression de type β , la λ -expression $\lambda x. Y$ que l'on peut construire à partir de x et de Y est de type $\alpha \rightarrow \beta$, d'où l'abstraction:

$$\begin{array}{c}
 [Y : \beta], [x : \alpha] \\
 \hline
 [\text{Abst}] \qquad [\lambda x. Y : \alpha \rightarrow \beta]
 \end{array}
 \quad \begin{array}{l}
 \text{(avec } x \text{ une variable} \\
 \text{qui a des occurrences dans } Y)
 \end{array}$$

Soient $\alpha_1, \dots, \alpha_n$ des types. Le type *produit cartésien* sera noté: $\alpha_1 \times \dots \times \alpha_n$. Si l'on assigne les types α_i à chaque expression Y_i ($i=1, \dots, n$), alors nous avons la règle d'assignation du type «produit cartésien» au n-uple $\langle Y_1, \dots, Y_n \rangle$:

$$\begin{array}{c}
 [Y_1 : \alpha_1], \dots, [Y_n : \alpha_n] \\
 \hline
 [\text{Cart}_t] \qquad [\langle Y_1, \dots, Y_n \rangle : \alpha_1 \times \dots \times \alpha_n]
 \end{array}$$

Remarque: Les notations choisies font apparaître explicitement les analogies entre d'une part l'application et l'élimination de ' $\rightarrow$ ' et d'autre part, entre l'abstraction et l'introduction de ' $\rightarrow$ '. L'abstraction peut être exprimée, en logique combinatoire par l'introduction d'un combinateur. Cette analogie formelle est hautement significative, elle est à la base de l'isomorphisme de Curry-Howard, dont nous allons reparler, entre types et propositions, λ -expressions et preuves, spécifications et programmes.

1. Opération d'application en linguistique

Dans une analyse formelle des langues, deux points de vue se dégagent: (a) le point de vue strictement concaténationnel et (b) le point de vue applicatif. Selon le premier point de vue, les unités linguistiques sont des morphèmes et des mots qui sont juxtaposés à gauche ou à droite à d'autres unités linguistiques. Le linguiste doit alors définir les règles (c'est-à-dire la grammaire) qui engendrent, analysent et représentent tous, et rien que, les énoncés empiriques d'une langue. Dans ce cas, le linguiste travaille dans un monoïde libre engendré par l'opération de concaténation à partir des unités linguistiques de base. Selon cette conception, la grammaire devient un ensemble de règles de réécriture qui opèrent sur des suites du monoïde afin de caractériser exactement un langage formel comme une certaine partie d'un monoïde libre, ce langage correspond exactement aux suites finies qui sont en correspondance avec les phrases empiriques de la langue décrite. Les grammaires formelles syntagmatiques de N. Chomsky se sont constituées en adoptant explicitement ce point de vue. Les arbres syntagmatiques⁴ sont directement associés à des classes d'unités linguistiques équivalentes (syntagmes nominaux, syntagmes verbaux, parties du discours, ...) organisées entre elles par des relations hiérarchiques d'emboîtement: telle classe (par exemple une phrase) est décomposée en une concaténation de classes plus fines (par exemple un syntagme nominal *suivi d'un* syntagme verbal). Les grammaires transformationnelles de Chomsky ne rompent pas avec la concaténation puisqu'elles continuent à opérer avec des suites concaténées de catégories et d'unités linguistiques minimales. Les modèles post-chomskyens (GPSG, HPSG, TAG) ne semblent pas avoir rompu non plus avec le point de vue concaténationnel. Ces derniers modèles ajoutent à l'opération fondamentale de concaténation, l'opération informatique d'unification.

Selon le second point de vue, les unités linguistiques sont des opérateurs ou des opérands qui sont organisés dans un énoncé,

4 La structure d'arbre, en linguistique ou en informatique, est une structure que nous avons appelée dendron; deux ordres sont sous-jacents: un ordre hiérarchique entre une catégorie et une sous-catégorie, un ordre concaténationnel entre sous-catégories et unités linguistiques minimales.

plus généralement dans un discours, par l'opération d'application. Le linguiste doit alors définir comment les unités s'organisent par l'application en précisant les contraintes sur l'application et éventuellement sur les types des unités linguistiques. L'application étant l'opération de base, elle est constitutive de toutes les expressions linguistiques représentées par des expressions applicatives qui sont alors encodées et présentées sous forme d'expressions linéaires en juxtaposant, selon des règles précises d'encodage, opérateurs et opérandes. La structure linéaire des énoncés est une structure non seulement superficielle⁵ mais aussi apparente, elle ne correspond donc pas à une organisation significative des unités linguistiques. Selon cette approche non concaténationnelle, l'énoncé, plus généralement le discours, est une configuration organisée essentiellement par l'application (d'opérateurs à des opérandes), d'autres opérations pouvant venir s'ajouter aux opérations applicatives de base⁶. La juxtaposition des expressions linguistiques qui aboutit à une séquence apparente de symboles («la chaîne parlée») est simplement une *présentation* d'une organisation applicative sous-jacente que la grammaire, elle, doit révéler et analyser. De fait, cette présentation linéaire est imposée par le canal phonique qu'utilise le code oral. Quant aux langues des signes, qui ne font pas appel au même canal, ont, elles aussi, une structure applicative sous-jacente, elles diffèrent simplement des langues parlées par les *modes de présentation* imposées par les canaux utilisés mais les opérations constitutives des énoncés restent de même nature.

2. Unités opératives dans les langues

En adoptant le point de vue que les unités des langues sont des unités opératives (ou fonctionnelles), les unités linguistiques (selon le détail de l'analyse: des morphèmes, des mots, des

5 ... qui s'opposerait à une structure profonde; dans les premiers modèles chomskyens, la structure profonde a, elle aussi, une structure concaténationnelle.

6 Les opérations prosodiques, dans les langues parlées, viennent se superposer aux autres organisations analysables par l'opération d'application. Cependant, il n'est pas exclu que les opérations prosodiques elles-mêmes puissent être analysées dans des termes applicatifs.

lexies – au sens de B. Pottier –, des syntagmes, des propositions, des phrases, des énoncés, des discours, ...) sont des opérateurs ou des opérandes: un opérateur s'applique à un opérande pour former une expression qui fonctionnera ou bien comme un opérande par rapport à un autre opérateur ou bien comme un opérateur par rapport à un autre opérande. Cette prise de position par rapport au langage n'est pas neutre puisqu'elle modélise toutes les unités linguistiques, non comme des symboles amorphes qui sont juxtaposés avec d'autres symboles amorphes, la signification ne surgissant que par l'assemblage final, mais comme des symboles opératoires qui agissent les uns sur les autres pour construire une signification qui est, elle-même, représentable en termes d'opérateurs et d'opérandes. Les mots ne sont donc pas de simples séquences juxtaposées de morphèmes grammaticaux et lexicaux mais des constructions décomposables en opérateurs et opérandes; les énoncés ne sont pas des séquences de mots juxtaposés mais des configurations dynamiques analysables en opérateurs et opérandes et représentables par des organisations géométriques (les arbres applicatifs).

Prenons un exemple très simple. Considérons le mot «endort»; il est décomposable en une organisation applicative représentable par l'arbre:

$$\begin{array}{ccc}
 & \text{en-} & \text{-dor-} \\
 & \hline
 \text{-t} & & \text{en-(-dor-) = endor-} \\
 & \hline
 & & \text{(-t)(endor-) = endort}
 \end{array}$$

Prenons un deuxième exemple. Soit la phrase: *Tout garçon aime le football*. La décomposition applicative est donnée comme suit par l'arbre applicatif suivant:

			le	football
tout	garçon	aime	le football	
tout garçon		aime (le football)		
(tout garçon) (aime (le football))				

De nombreuses théories linguistiques, tant en morphologie qu'en syntaxe, ont adopté, quelquefois implicitement, ce point de vue: les unités linguistiques des langues sont des symboles opératoires. Dans la Grammaire de Harris (dans ses différentes versions, depuis 1968), les unités linguistiques sont des opérateurs ou des opérands, la décomposition d'une unité discursive (un discours) en opérateur ou opérande est essentielle dans la théorie. La Grammaire Applicative de S.K. Shaumyan (là aussi, dans ses différentes versions depuis 1963) considère explicitement les unités linguistiques (les sémions) comme des opérateurs ou des opérands organisés par l'opération d'application d'où le nom de «Grammaire Applicative». La syntaxe structurale de Tesnière repose implicitement et en partie seulement sur cette notion applicative. En effet, dans cette approche, le verbe, unité centrale de la phrase, peut être considéré comme un opérateur qui s'applique (globalement) aux opérands que sont les actants mais la notion de dépendance (par exemple, «l'article dépend du nom») n'exploite pas à fond la décomposition entre opérateur et opérande.

3. Grammaires catégorielles classiques

Les unités linguistiques ne sont pas homogènes, ni dans leurs morphologies, ni dans leurs fonctionnements. Il y a donc différents types d'unités linguistiques. Prenons le français par exemple. Une analyse, même très sommaire, fait reconnaître une différence entre les verbes et les noms, entre les noms et les adjectifs, entre les adjectifs et les adverbes, entre les prépositions et les conjonctions. Une des premières tâches du grammairien

revient donc à identifier les types des unités linguistiques manipulées: les types morphologiques, les types syntaxiques, les types sémantiques et, dans une approche cognitive du langage, les types cognitifs, ... Dans un certain nombre de cas, les types sont identifiées clairement par des classes d'unités linguistiques homogènes car ayant les mêmes propriétés formelles et un même fonctionnement dans les énoncés. Pour certaines unités cependant, la notion de type n'est pas toujours identifiable à une classe homogène d'unités, aussi n'assimilerons-nous pas *a priori* types et classes, la notion de type se prêtant à des abstractions opérationnelles et à des généralisations que n'autoriserait pas la notion de classe. Par exemple, la notion de type permet la construction de types, au moyen d'opérations précises – les constructeurs de types – à partir de types de base, puis d'établir des relations génétiques entre les types et parfois d'envisager un calcul inférentiel sur les types. Une entité typée est une entité à laquelle on a assigné un type; cela signifie que l'on connaît exactement: (i) quelles sont les opérations explicites que l'on peut faire sur cette entité; (ii) comment cette entité peut être représentée, dans certains cas même comment elle peut être physiquement implantée sur des supports matériels; (iii) comment cette entité peut être reliée à d'autres entités de type différent puisque connaissant son type, on connaît les relations génétiques que ce type entretient avec les autres types. Par exemple, le type **nat** des entiers naturels a un mode particulier d'implantation physique dans une machine informatique et ce mode diffère du type **réel** des réels. Si une entité est du type **bool** des entités booléennes, les opérateurs qui s'appliqueront à cette entité seront les opérateurs de négation, de conjonction, de disjonction, d'implication mais pas la soustraction arithmétique. Pour prendre une analogie, le type assigné correspond aux «dimensions» d'une expression du langage de la physique. Par exemple, dans l'expression linguistique du principe fondamental de la mécanique $F = ma$, les deux côtés du signe '=' doivent être de même type (principe d'isotypicalité), autrement dit la composition par multiplication d'une masse 'm' avec une accélération 'a' doit représenter une entité de même type qu'une force. En

prenant pour types de base: L (longueur), T (temps) et M (masse), le type d'une force est donc: $M \times (L/T^2)$.

En linguistique, les notions de type sémantique et syntaxique ont pris un sens technique et intuitif avec les Grammaires Catégorielles entrevues par Husserl. Ce formalisme a été ensuite développé en sémantique par Lesniewski, puis, en syntaxe, par Adjukiewicz, Curry, Bar Hillel pour être réexaminé et étendu par Lambek, Geatch, van Bethem, Bach et Partee, Morgaat, Steedman, ... Les types syntaxiques assignés aux unités et aux expressions linguistiques sont tous dérivés de types syntaxiques de base. Ces types syntaxiques (types de base et dérivés) représentent des catégories syntaxiques d'expressions linguistiques: noms, verbes intransitifs, verbes transitifs, adjectifs, ... Toutes les catégories syntaxiques en étant représentées par des types syntaxiques, sont analysées comme décomposables en catégories de base. On peut prendre selon les versions des Grammaires catégorielles, soit les deux catégories de base, «nom» et «phrase» ou soit les trois catégories de base, «syntagme nominal», «nom commun» et «phrase». Désignons ces dernières catégories de base par respectivement: 'N', 'NC' et 'P'. Les catégories dérivées sont construites à partir de deux constructeurs de types, les «divisions à droite et à gauche», désignées respectivement par '/' et '\'. On ajoute en général une autre constructeur, le «produit cartésien», désigné par 'x'.

L'ensemble $TYPE_{synt}$ de tous les *types syntaxiques dérivés* d'un ensemble de types syntaxiques de base $TYPE_{synt0}$ est défini récursivement comme suit:

- (i) Les types de base de $TYPE_{synt0}$ sont des types de $TYPE_{synt}$;
- (ii) Si: α et β sont des types de $TYPE_{synt}$
alors: ' α/β ', ' $\alpha\backslash\beta$ ', ' $\alpha \times \beta$ ' sont des types de $TYPE_{synt}$.

Les types ' α/β ', ' $\alpha\backslash\beta$ ' et ' $\alpha \times \beta$ ' se lisent respectivement « β sur α », « β sous α » et « α produit cartésien β ».

A l'opération d'application [**App**_t] sur les unités typées correspondent deux règles applicatives⁷ [**AB**>] et [**AB**<]:

7 [AB] pour rappeler Adjukiewicz-Bar Hillel.

$$\begin{array}{ccc}
 [X : \beta/\alpha] & [Y : \alpha] & \\
 [AB>] \xrightarrow{\quad} & & [Y : \alpha] \quad [X : \beta \backslash \alpha] \\
 & [X-Y : \beta] & \\
 & & [Y-X : \beta] \xleftarrow{\quad} \\
 & & [<AB]
 \end{array}$$

Dans la prémisses, les deux expressions X et Y ont des types assignés; dans la conclusion les expressions sont obtenues en juxtaposant, par l'opération notée '-', respectivement: l'opérateur X et l'opérande Y à sa droite⁸ au moyen de la règle $[AB>]$ et l'opérateur X et l'opérande Y à sa gauche au moyen de la règle $[AB<]$. Ces règles font apparaître l'interprétation fonctionnelle des types β/α et $\beta \backslash \alpha$.

Nous nous donnons également la règle d'assignation d'un type «produit» à un couple d'unités typées

$$\begin{array}{c}
 [X : \alpha], [Y : \beta] \\
 [x_t] \quad \frac{\quad}{[<X, Y> : \alpha \times \beta]}
 \end{array}$$

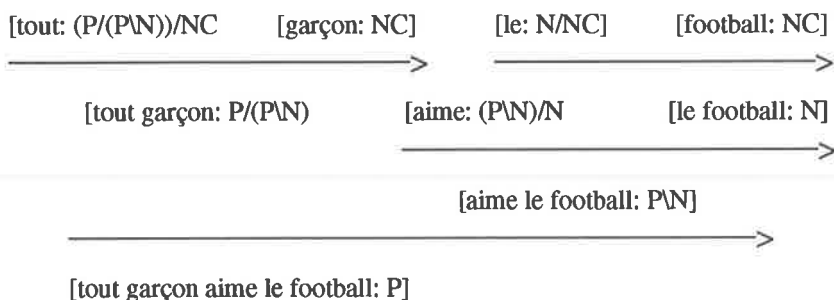
Cette règle nous invite à interpréter le type ' $\alpha \times \beta$ ' comme un «produit cartésien fini».

Prenons, à titre d'exemple, l'ensemble $TYPE_0 = \{N, NC, P\}$ comme types de base; nous en déduisons des types dérivés qui représentent les catégories syntaxiques les plus classiques:

8 Les deux expressions en conclusion de la règle ne sont pas des expressions applicatives mais des expressions concaténées; on a désigné par '-' le symbole de concaténation.

Catégories syntaxiques	Types syntaxiques
nom propre	N
phrase	P
nom commun	NC
verbe intransitif	P\N
verbe transitif	(P\N)/N
adjectif	NC/NC
article	N/NC
adverbe	(P\N)\(P\N)
quantificateur	(P/(P\N))/NC

Assignons des types aux unités linguistiques de base. Nous pouvons reprendre l'analyse applicative de la phrase *tout garçon aime le football*, en assignant maintenant des types à chaque unité linguistique (ici, dans l'exemple, des mots lexicaux), nous obtenons, en utilisant les règles $[AB>]$ et $[AB<]$, l'arbre suivant:



Les deux règles $[AB>]$ et $[AB<]$ sont analogues aux simplifications formelles du calcul sur les fractions ou encore aux simplifications formelles dans une vérification des «dimensions» d'une formule de physique. En se restreignant aux seuls types, les deux règles se ramènent formellement aux simplifications (à droite et à gauche) suivantes:

$$\frac{\alpha}{\beta} \quad \frac{\beta}{\alpha} \quad \frac{\alpha}{\beta} \quad \frac{\beta}{\alpha}$$

$$\frac{\alpha}{\beta} \quad \frac{\beta}{\alpha} \quad \frac{\alpha}{\beta} \quad \frac{\beta}{\alpha}$$

Cette remarque a conduit J. Lambek à introduire un préordre de réduction entre les types, notés ' $\Rightarrow$ ', pour structurer le système des types TYPE engendré à partir des types de base TYPE₀ par les constructeurs de types /, \, x et en faire un système inférentiel (par la relation $\Rightarrow$ entre types):

$$\langle \text{TYPE}_0, /, \backslash, x, \Rightarrow \rangle.$$

Le **Calcul de Lambek** est un système de déduction sur les types où l'on se donne des schémas d'axiome et des schémas de règles qui permettent de déduire un ensemble de relations significatives entre types (les théorèmes du système de Lambek). En particulier, en se donnant, entre autres, comme schéma d'axiome et schéma de règle de déduction:

$$\frac{\alpha \Rightarrow \alpha \quad \alpha \Rightarrow \beta, \beta \Rightarrow \gamma}{\alpha \Rightarrow \gamma}$$

on en déduit, en tant que théorèmes de déduction sur les types, les analogues des règles applicatives [AB>] et [AB<]:

$$\frac{(\alpha/\beta) \times \beta \Rightarrow \alpha \quad \beta \times (\alpha \backslash \beta) \Rightarrow \alpha.}{\alpha \Rightarrow \alpha.}$$

D'autres théorèmes sont également déduits comme, par exemple:

$$\frac{\gamma/(\alpha \times \beta) \Leftrightarrow \gamma/(\alpha/\beta) \quad \gamma/(\alpha \times \beta) \Leftrightarrow \gamma/(\beta/\alpha)}$$

qui ont l'interprétation donnée par le principe de «currification» (voir plus loin). Ces nouvelles relations entre types étendent le calcul initial de Adjuikiewicz et Bar-Hillel sur les types (dit sys-

tème **AB**) associé directement aux Grammaires Catégorielles classiques. Le Calcul de Lambek est donc plus riche que celui des catégories dans les Grammaires Catégorielles classiques.

Lorsque des types sont assignés aux unités linguistiques minimales, les Grammaires Catégorielles classiques substituent un calcul inférentiel sur les types à un calcul direct sur les unités linguistiques, de façon à vérifier que la formule de type qui spécifie l'énoncé est réductible au type caractéristique des énoncés: la formule du type de l'énoncé correspond au type P . A l'analyse catégorielle de la phrase *tout garçon aime une fille*, correspond donc une succession de réductions sur les types:

$$\begin{aligned}
 ((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC & \Rightarrow \\
 ((P/(P\backslash N)) \times ((P\backslash N)/N) \times (N/NC) \times NC & \Rightarrow \\
 (P/(P\backslash N)) \times ((P\backslash N)/N) \times N & \Rightarrow \\
 (P/(P\backslash N)) \times (P\backslash N) & \Rightarrow \\
 P. &
 \end{aligned}$$

Nous en déduisons le théorème: «le type initial est réductible au type P »:

$$((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC \Rightarrow P.$$

En inversant la relation de réduction sur les types, on obtient à partir du type de phrase P (et de l'axiome canoniquement associé: $P \Rightarrow P$), par expansions successives des types, une formule de type qui spécifie l'analyse syntaxique de toutes les phrases de même type. Ainsi, on obtient:

$$\begin{aligned}
 P & \Rightarrow P \\
 P & \Rightarrow (P/(P\backslash N)) \times (P\backslash N) \\
 P & \Rightarrow (P/(P\backslash N)) \times (((P\backslash N)/N) \times N) \\
 P & \Rightarrow (P/(P\backslash N)) \times (((P\backslash N)/N) \times ((N/NC) \times NC)) \\
 P & \Rightarrow (((P/(P\backslash N))/NC) \times NC) \times (((P\backslash N)/N) \times ((N/NC) \times NC)).
 \end{aligned}$$

Le type:

$$((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC$$

apparaît comme une *spécification syntaxique* d'une famille de phrases de même type c'est-à-dire organisées syntaxiquement de

la même façon: ces phrases sont toutes des instances du type donné.

Les notions d'application et de type permettent le déploiement des Grammaires Catégorielles. Le calcul inférentiel sur les types établit un lien entre l'analyse syntaxique applicative et la théorie de la preuve. En effet, effectuer une analyse applicative d'un énoncé, c'est construire explicitement, à partir d'une formule de type qui résulte d'une assignation de types aux unités minimales de l'énoncé, un arbre applicatif dont la racine est P, mais c'est aussi exhiber une preuve que la formule de type se réduit par la relation $\Rightarrow$ au type P. Or, de l'arbre applicatif, on extrait une déduction sur les types et, réciproquement, d'une déduction sur les types on peut construire un arbre applicatif. Les présentations usuelles de ces Grammaires ne se sont cependant pas complètement dégagées de la conception concaténationnelle. Pour intégrer complètement la conception applicative avec celle des types, il faut étendre les Grammaires Catégorielles en exploitant la signification profonde des types fonctionnels. C'est ce que nous avons entrepris avec les Grammaires Catégorielles Combinatoires et Applicatives.

4. Analogie de Curry-Howard

4.1. Currification

Les Grammaires Catégorielles classiques ou étendues (dans le Calcul de Lambek) substituent, comme nous venons de le voir, un calcul inférentiel sur les types syntaxiques à un calcul direct sur les unités linguistiques agencées dans une phrase ou dans un énoncé. Lorsqu'on effectue un calcul sur les unités linguistiques en *se laissant guider* par les réductions de types, on construit progressivement l'énoncé en juxtaposant les expressions linguistiques à chaque fois qu'on opère une réduction sur les types selon les règles $[AB>]$ ou $[AB<]$ ou des règles plus complexes du Calcul de Lambek. Remarquons bien que si l'expression obtenue est une expression présentée par la juxtaposition linéaire des unités linguistiques, cette expression est

construite suivant une structure applicative qui décrit «l'histoire de sa construction». Les règles applicatives $[AB>]$, $[AB<]$ nous amènent directement à interpréter les opérateurs ' $'$ ' et ' $\backslash$ ' comme des constructeurs de types fonctionnels et à interpréter l'opérateur ' $\times$ ' comme un constructeur d'un type »produit cartésien«. Ainsi, les types construits par les diviseurs *ont une signification fonctionnelle*: le type noté β/α exprime le type d'un opérateur qui construit à partir d'un opérande de type α une expression de type β ; il en est de même du type $\beta\alpha$. En «oubliant» maintenant les positions à droite ou à gauche de l'opérande par rapport à l'opérateur, il apparaît que les types syntaxiques de la forme α/β et $\alpha\beta$ sont des types fonctionnels dont les instances sont nécessairement des opérateurs, avec la correspondance suivante:

type fonctionnel	diviseur à gauche	diviseur à droite
$\alpha \rightarrow \beta$	β/α	$\beta\alpha$.

Définissons directement les types fonctionnels et les types «produit cartésien» de la même façon que nous avons défini les types syntaxiques de $TYPE_{\text{synt}}$ (voir 3.). Soit $TYPE_0$ un ensemble de types de base. Les types fonctionnels et les types «produit cartésien» sont construits au moyen des règles suivantes:

- (i) Les types de base de $TYPE_0$ sont des types de $TYPE$;
- (ii) Si α et β sont des types de $TYPE$ alors: ' $\alpha \rightarrow \beta$ ' est un type de $TYPE$;
- (iii) Si α et β sont des types de $TYPE$ alors ' $\alpha \times \beta$ ' est un type de $TYPE$.

On peut considérer des opérateurs unaires, binaires, ... n-aires qui construisent des expressions de type β ; ces opérateurs sont caractérisés par les types suivants:

type d'opérateurs unaires	$\alpha \rightarrow \beta$
type d'opérateurs binaires	$\alpha_1 \rightarrow (\alpha_2 \rightarrow \beta)$
...	
type d'opérateurs n-aires	$\alpha_1 \rightarrow (\alpha_2 \rightarrow (...(\alpha_n \rightarrow \beta) ...))$

Puisque l'on peut former le produit cartésien de types ' $\alpha_1 \times \alpha_2 \times \dots \times \alpha_n$ ', le type d'un opérateur n-aire doit être également de la forme ' $\alpha_1 \times \alpha_2 \times \dots \times \alpha_n \rightarrow \beta$ '. Un tel opérateur, disons X_n , construit une expression de type β à partir de la donnée d'un n-uple d'opérandes de types respectifs $\alpha_1, \dots, \alpha_n$. Plus précisément, nous avons en généralisant la règle d'application [App] au produit cartésien:

$$\begin{array}{c} [X_n : \alpha_1 \times \alpha_2 \times \dots \times \alpha_n \rightarrow \beta] \quad [Y^1, \dots, Y^n : \alpha_1 \times \alpha_2 \times \dots \times \alpha_n] \\ \hline [\text{App}_{tn}] \quad [X_n(Y^1 \dots Y^n) : \beta] \end{array}$$

Quels sont les rapports entre les deux types d'un opérateur n-aire, c'est-à-dire entre d'une part, le type suivant: $\alpha_1 \times \alpha_2 \times \dots \times \alpha_n \rightarrow \beta$ et d'autre part, le type: $\alpha_1 \rightarrow (\alpha_2 \rightarrow (\dots (\alpha_n \rightarrow \beta) \dots))$? Sont-ils identiques, bien que construits différemment? Le principe de «currification» (ou de «curriage») répond à cette question en établissant une bijection entre ces deux types:

$$[\text{CUR}_t] \quad \alpha_1 \times \alpha_2 \times \dots \times \alpha_n \rightarrow \beta = \alpha_1 \rightarrow (\alpha_2 \rightarrow (\dots (\alpha_n \rightarrow \beta) \dots)).$$

L'opération de «currification» d'un opérateur n-aire X_n , de type: $\alpha_1 \times \alpha_2 \times \dots \times \alpha_n \rightarrow \beta$, revient à lui associer biunivoquement l'opérateur unaire $X_1' = \text{Curr}(X_n)$, de type $\alpha \rightarrow \gamma$, où γ est le type fonctionnel $\alpha_2 \rightarrow (\dots (\alpha_n \rightarrow \beta) \dots)$. Le principe $[\text{CUR}_t]$ de currification indique que l'on peut considérer des opérateurs n-aires comme des opérateurs unaires et que le type produit cartésien $\alpha_1 \times \alpha_2 \times \dots \times \alpha_n$ n'est plus nécessaire pour différencier le type des différents opérateurs unaires, binaires, ..., n-aires.

Quelles sont les raisons mathématiques du principe de currification $[\text{CUR}_t]$? Remarquons que ce principe, dans le cas $n=2$, est un théorème du système de Lambek puisque nous avons (voir précédemment) la déduction entre types:

$$\alpha_1 \times \alpha_2 \rightarrow \beta \Leftrightarrow \alpha_1 \rightarrow (\alpha_2 \rightarrow \beta).$$

Lorsque l'on considère un type de base comme une sorte d'ensemble, les types fonctionnels représentent des ensembles de fonctions entre ensembles. Dans le cas de la Catégorie des ensembles *Ens*, A_1 , A_2 et B étant des ensembles et $\text{Ens}(X, Y)$

désignant l'ensemble de toutes les fonctions de X dans Y , nous avons la bijection suivante:

$$\mathbf{Ens}(A_1 \times A_2, B) = \mathbf{Ens}(A_1, \mathbf{Ens}(A_2, B))$$

qui à toute fonction f_2 du produit cartésien d'ensembles $A_1 \times A_2$ dans B associe canoniquement la fonction f_1 de A_1 dans l'ensemble $\mathbf{Ens}(A_2, B)$ des fonctions de A_2 dans B , c'est-à-dire que l'on a, dans l'ensemble B , l'égalité suivante:

$$f_2(a_1, a_2) = (f_1(a_1))(a_2) \quad \text{pour chaque couple } \langle a_1, a_2 \rangle \text{ de } A_1 \times A_2.$$

Cette propriété des ensembles se généralise aux produits finis cartésiens d'ensembles, elle est caractéristique des ensembles et plus généralement de toutes les Catégories cartésiennes fermées.

Du point de vue linguistique, l'opération de currification est importante puisqu'elle permet de ne considérer que des opérateurs unaires de type $\alpha \rightarrow \beta$ et des opérateurs zéro-aires considérés comme des opérandes absolues. Ainsi, en admettant la théorie de la valence, des verbes bivalent ou trivalent sont considérés comme des opérateurs unaires qui à partir d'un opérande construisent des opérateurs unaires ou binaires. Le type d'un opérateur donne donc toutes les informations sur son arité. En associant des types fonctionnels aux unités linguistiques à partir des catégories d'une Grammaire Catégorielle classique, nous pouvons considérer que les trois arbres applicatifs (A_x) (A_1) et (A_2) suivants construisent une expression applicative de type p :

$$[\text{admire} : n \times n \rightarrow p] \quad [\langle \text{Pierre}, \text{Marie} \rangle : n \times n]$$

[admire (Pierre, Marie): p]

Arbre applicatif (A_x)

[admire :n->(n ->p)] [Pierre :n]

[admire :n->(n ->p)] [Marie :n]

[admire Pierre :n ->p]

[Marie :n]

[admire Pierre :n ->p]

[Pierre :n]

[admire Marie Pierre : p]

[admire Marie Pierre : p]

Arbre applicatif (A₁)

Arbre applicatif (A₂)

Les résultats, qui sont de même type p, sont équivalents, c'est-à-dire:

admire (Pierre, Marie) = (admire Marie) Pierre.

Cependant, ces résultats ne sont pas construits de la même façon, ils n'ont pas la même histoire constitutive. Dans l'arbre applicatif (A_x), le prédicat verbal bivalent est appliqué globalement au couple des arguments <Pierre, Marie>, comme dans le calcul des prédicats classique. Dans l'histoire constitutive de l'arbre (A₁), il apparaît que le prédicat verbal, considéré comme un opérateur binaire, est appliqué dans un premier temps au terme qui fonctionne syntaxiquement comme «sujet», puis le résultat, considéré comme un opérateur unaire, est appliqué à son tour au terme qui fonctionne comme un «complément d'objet». Dans la constitution de l'arbre (A₂), la construction est effectuée en appliquant dans un premier temps le prédicat verbal au terme «objet» puis le résultat au terme «sujet», qui est ainsi caractérisé comme étant le dernier terme qui entre dans une relation prédicative, cette propriété servant à caractériser le terme «sujet». En explicitant l'histoire constitutive des résultats construits et en utilisant la λ -notation, nous avons:

$(\lambda \langle u_1, u_2 \rangle. \text{admire}_2(u_1, u_2)) (\langle \text{Pierre}, \text{Marie} \rangle) =$
 admire (Pierre, Marie);

$((\lambda x. (\lambda y. \text{admire}'_1(x, y))) (\text{Pierre})) (\text{Marie}) =$
 admire Marie Pierre;

$((\lambda y. (\lambda x. \text{admire}'_1(x, y))) (\text{Marie})) (\text{Pierre}) =$
 admire Marie Pierre.

Divers arguments linguistiques montreraient que la construction (A_2) est celle qui doit être retenue dans une théorie linguistique des opérations de prédication. En effet, entre autres arguments, on pourrait montrer qu'il y a une relation «plus serrée» entre le prédicat verbal et le terme «objet» qu'entre le prédicat verbal et le terme «sujet»: la construction (A_2) reflète parfaitement cette hiérarchisation des arguments qui entrent successivement dans la construction applicative d'une relation prédictive, hiérarchisation qui n'est absolument pas exprimée par le formalisme du calcul des prédicats.

4.2. Analogie de Curry-Howard entre types et propositions

H. B. Curry (1958) a remarqué que les types engendrés par des types de base à l'aide du constructeur des types fonctionnels étaient en bijection avec les propositions de la logique positive des propositions dont le seul connecteur est l'opérateur d'implication, désigné ici par ' $\rightarrow$ ': toute formule du Calcul sur les types qui est un théorème peut ainsi être traduite immédiatement par un théorème du Calcul positif des propositions et réciproquement. Curry a également remarqué que tous les types des combinateurs de la logique combinatoire, dont nous reparlerons au paragraphe suivant, sont des formules qui sont analogues à des théorèmes du calcul propositionnel intuitionniste. En effet, il suffit d'interpréter le constructeur de fonctionnalité par le connecteur d'implication et d'interpréter les types comme des propositions. Ainsi, les formules comme:

$$\begin{aligned} (B \rightarrow C) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C)) \\ (A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C)) \end{aligned}$$

s'interprètent aussi bien comme des types que comme des propositions; lorsque ce sont des types, la flèche ' $\rightarrow$ ' est le constructeur de fonctionnalité entre types; lorsque ce sont des propositions, la flèche ' $\rightarrow$ ' signifie l'implication entre propositions. Ces formules sont des théorèmes du Calcul propositionnel intuitionnisme et des théorèmes du Calcul sur les types. Dans une présentation par déduction naturelle dans le style de

Gentzen, on peut en donner des démonstrations en se servant des seules règles d'élimination et d'introduction de la flèche $\rightarrow$:

$$\begin{array}{c}
 \begin{array}{c} \alpha \rightarrow \beta \quad \alpha \\ \hline \beta \end{array} \\
 [e \rightarrow]
 \end{array}
 \qquad
 \begin{array}{c}
 [\alpha] \\
 \cdot \\
 \cdot \\
 \beta \\
 \hline
 [i \rightarrow] \quad \alpha \rightarrow \beta
 \end{array}$$

La preuve que $(B \rightarrow C) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$ est un théorème est donnée comme suit en «dédution naturelle»:

1.		B $\rightarrow$ C	hyp.
2.		A $\rightarrow$ B	hyp.
3.		A	hyp.
4.		A $\rightarrow$ B	rappel 2.
5.		B	[e $\rightarrow$], 3., 4.
6.		B $\rightarrow$ C	rappel 1.
7.		C	[e $\rightarrow$], 5., 6.
8.		A $\rightarrow$ C	[i $\rightarrow$], 3.-7.
9.		(A $\rightarrow$ B) $\rightarrow$ (A $\rightarrow$ C)	[i $\rightarrow$], 2.-8
10.		(B $\rightarrow$ C) $\rightarrow$ ((A $\rightarrow$ B) $\rightarrow$ (A $\rightarrow$ C))	[i $\rightarrow$], 1.-9.

La preuve que $(A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$ est un théorème est donnée ci-dessous en «dédution naturelle»:

1.		A $\rightarrow$ (B $\rightarrow$ C)	hyp.
2.		A $\rightarrow$ B	hyp.
3.		A	hyp.
4.		A $\rightarrow$ (B $\rightarrow$ C)	rappel 1.
5.		B $\rightarrow$ C	[e $\rightarrow$], 3., 4.
6.		A $\rightarrow$ B	rappel 2.
7.		B	[e $\rightarrow$], 3., 6.
8.		C	[e $\rightarrow$], 5., 7.
9.		A $\rightarrow$ C	[i $\rightarrow$], 3.-8.
10.		(A $\rightarrow$ B) $\rightarrow$ (A $\rightarrow$ C)	[i $\rightarrow$], 2.-9.
11.		(A $\rightarrow$ (B $\rightarrow$ C)) $\rightarrow$ ((A $\rightarrow$ B) $\rightarrow$ (A $\rightarrow$ C))	[i $\rightarrow$], 1.-10.

A titre d'exercice, montrons que le type syntaxique:

$$((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC$$

(ou son analogue propositionnel) est bien une expression de type P. Pour simplifier les notations, remplaçons les symboles qui représentent des catégories élémentaires P, N et NC par les symboles abstraits respectifs A, B et C interprétés comme des types ou comme des propositions. Substituons la flèche ' $\rightarrow$ ' aux constructeurs de division $/$ ou $\backslash$ et interprétons le produit ' $\times$ ' soit par le produit cartésien des types, soit par la conjonction de propositions. Le produit ' $\times$ ', tout comme la conjonction propositionnelle, est gouvernée par les deux règles d'élimination (associées aux projections) et d'introduction suivantes:

$$\begin{array}{ccc} \text{A} \times \text{B} & \text{A} \times \text{B} & \text{A, B} \\ [e_1-x] \frac{\quad}{\text{A}} & [e_2-x] \frac{\quad}{\text{B}} & [i-x] \frac{\quad}{\text{A} \times \text{B}} \end{array}$$

Après ces changements, le type syntaxique précédent devient:

$$(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C \times (B \rightarrow (B \rightarrow A))) \times (C \rightarrow B) \times C$$

Donnons la preuve que cette formule (type ou proposition) est une expression bien formée de type A.

- | | | |
|-----|---|------------------------------|
| 1. | $(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C \times (B \rightarrow (B \rightarrow A))) \times (C \rightarrow B) \times C$ | hyp. |
| 2. | $(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C$ | $[e-x]$, 1. |
| 3. | $(B \rightarrow (B \rightarrow A)) \times (C \rightarrow B) \times C$ | $[e-x]$, 1. |
| 4. | $(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C$ | rappel 2. |
| 5. | $(C \rightarrow ((B \rightarrow A) \rightarrow A)$ | $[e-x]$, 4. |
| 6. | C | $[e-x]$, 4. |
| 7. | $((B \rightarrow A) \rightarrow A)$ | $[e \rightarrow]$, 5., 6. |
| 8. | $(B \rightarrow (B \rightarrow A)) \times (C \rightarrow B) \times C$ | rappel 3. |
| 9. | $(B \rightarrow (B \rightarrow A))$ | $[e-x]$, 8. |
| 10. | $(C \rightarrow B)$ | $[e-x]$, 8. |
| 11. | C | $[e-x]$, 8. |
| 12. | B | $[e \rightarrow]$, 10., 11. |
| 13. | $B \rightarrow A$ | $[e \rightarrow]$, 9., 12. |
| 14. | A | $[e \rightarrow]$, 5., 13. |

On en déduit que:

$$(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C \times (B \rightarrow (B \rightarrow A))) \times (C \rightarrow B) \times C \Rightarrow A$$

est bien un théorème puisque $A \Rightarrow A$ est un schéma d'axiome (dans le Calcul sur les types ou, par isomorphisme, dans le Calcul propositionnel).

Le Calcul sur les types et le Calcul positif sur les propositions étant isomorphes, toute démonstration de théorèmes du Calcul des propositions se transporte immédiatement dans le Calcul sur les types. En particulier, comme l'a montré J. Lambek, le théorème de Gentzen d'élimination de la coupure (Cut), qui donne une méthode effective pour démontrer des théorèmes dans le Calcul des propositions, est applicable au Calcul sur les types.

4.3. Interprétation fonctionnelle guidée par la syntaxe

Les preuves de théorèmes donnent la possibilité de construire directement une interprétation sémantique fonctionnelle dans le λ -calcul. Aux règles d'élimination $[e \rightarrow]$ et d'introduction $[i \rightarrow]$ du symbole ' $\rightarrow$ ' du Calcul des types sont canoniquement associées des règles de construction des expressions applicatives typées. Plus précisément, nous avons les deux règles de construction par application $[e' \rightarrow]$ et par abstraction $[i' \rightarrow]$ suivantes:

$$\begin{array}{c}
 \begin{array}{ccc}
 X : \alpha \rightarrow \beta & Y : \alpha & [X : \alpha] \\
 [e' \rightarrow] \frac{}{XY : \beta} & & \cdot \\
 & & \cdot \\
 & & Y : \beta \\
 & [i' \rightarrow] \frac{}{\lambda X.Y : \alpha \rightarrow \beta}
 \end{array}
 \end{array}$$

La première règle exprime que si un opérateur X de type $\alpha \rightarrow \beta$ est appliqué à un opérande Y de type α alors le résultat XY est de type β . La seconde règle exprime que, avec l'hypothèse que X soit de type α , de l'expression Y de type β (dans laquelle apparaît une ou plusieurs occurrences de X), on peut abstraire une λ -expression $\lambda X.Y$, c'est-à-dire un opérateur, de type $\alpha \rightarrow \beta$.

Reprenons les exemples précédents. Les preuves qui ont été données permettent de construire des expressions applicatives (des λ -expressions) typées. Nous avons ainsi:

1.		X : B $\rightarrow$ C	hyp.
2.		Y : A $\rightarrow$ B	hyp.
3.		Z : A	hyp.
4.		Y : A $\rightarrow$ B	rappel 2.
5.		YZ : B	[e' $\rightarrow$], 3., 4.
6.		X : B $\rightarrow$ C	rappel 1.
7.		X(YZ) : C	[e' $\rightarrow$], 5., 6.
8.		$\lambda Z.X(YZ) : A \rightarrow C$	[i' $\rightarrow$], 3.-7.
9.		$\lambda Y.\lambda Z.X(YZ) : (A \rightarrow B) \rightarrow (A \rightarrow C)$	[i' $\rightarrow$], 2.-8.
10.		$\lambda X.\lambda Y.\lambda Z.X(YZ) : (B \rightarrow C) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$	[i' $\rightarrow$], 1.-9.

Nous venons de construire l'expression applicative (la λ -expression) $\lambda X.\lambda Y.\lambda Z.X(YZ)$ de type:

$$(B \rightarrow C) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$$

avec les variables abstraites X, Y et z de types respectifs: B $\rightarrow$ C, A $\rightarrow$ B, A. Cette expression applicative code directement la preuve que le type $(B \rightarrow C) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$ est un théorème.

De la même façon nous construisons l'expression applicative $\lambda X.\lambda Y.\lambda Z.XZ(YZ)$ à partir de la preuve que son type $((A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C)))$ est un théorème.

1.		X : A $\rightarrow$ (B $\rightarrow$ C)	hyp.
2.		Y : A $\rightarrow$ B	hyp.
3.		Z : A	hyp.
4.		X : A $\rightarrow$ (B $\rightarrow$ C)	rappel 1.
5.		XZ : B $\rightarrow$ C	[e' $\rightarrow$], 3., 4.
6.		Y : A $\rightarrow$ B	rappel 2.
7.		YZ : B	[e' $\rightarrow$], 3., 6.
8.		XZ(YZ) : C	[e' $\rightarrow$], 5., 7.
9.		$\lambda Z.XZ(YZ) : A \rightarrow C$:	[i' $\rightarrow$], 3.-8.
10.		$\lambda Y.\lambda Z.XZ(YZ) : (A \rightarrow B) \rightarrow (A \rightarrow C)$	[i' $\rightarrow$], 2.-9.
11.		$\lambda X.\lambda Y.\lambda Z.XZ(YZ) : ((A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C)))$	[i' $\rightarrow$], 1.-10.

Toujours à titre d'exercice, reprenons la preuve de:

$$(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C \times (B \rightarrow (B \rightarrow A))) \times (C \rightarrow B) \times C \Rightarrow A$$

Supposons que nous ayons l'assignation suivante des types à des expressions applicatives X, Y, Z, T et U:

$[X : (C \rightarrow ((B \rightarrow A) \rightarrow A))]; [Y : C]; [Z : (B \rightarrow (B \rightarrow A))]; [T : (C \rightarrow B)]; [U : C]$

Nous en déduisons la construction de l'expression applicative $'((XY) (Z(TU)))'$ à partir de la preuve que son type $'(C \rightarrow ((B \rightarrow A) \rightarrow A) \times C \times (B \rightarrow (B \rightarrow A)) \times (C \rightarrow B) \times C'$ est réductible à A.

- | | | | |
|--|--|----------|------------------------------|
| 1. $[X : (C \rightarrow ((B \rightarrow A) \rightarrow A))]$ | $[Y : C] \times [Z : (B \rightarrow (B \rightarrow A))]$ | $\times$ | |
| | $[T : (C \rightarrow B)] \times [U : C]$ | | hyp. |
| 2. $[X : (C \rightarrow ((B \rightarrow A) \rightarrow A))]$ | $\times [Y : C]$ | | $[e-x], 1.$ |
| 3. $[Z : (B \rightarrow (B \rightarrow A))]$ | $\times [T : (C \rightarrow B)] \times [U : C]$ | | $[e-x], 1.$ |
| 4. $[X : (C \rightarrow ((B \rightarrow A) \rightarrow A))]$ | $\times [Y : C]$ | | rappel 2. |
| 5. $[X : (C \rightarrow ((B \rightarrow A) \rightarrow A))]$ | | | $[e-x], 4.$ |
| 6. $[Y : C]$ | | | $[e-x], 4.$ |
| 7. $[XY : (B \rightarrow A) \rightarrow A]$ | | | $[e' \rightarrow], 5., 6.$ |
| 8. $[Z : (B \rightarrow (B \rightarrow A))]$ | $\times [T : (C \rightarrow B)] \times [U : C]$ | | rappel 3. |
| 9. $[Z : (B \rightarrow (B \rightarrow A))]$ | | | $[e-x], 8.$ |
| 10. $[T : (C \rightarrow B)]$ | | | $[e-x], 8.$ |
| 11. $[U : C]$ | | | $[e-x], 8.$ |
| 12. $[TU : B]$ | | | $[e' \rightarrow], 10., 11.$ |
| 13. $[Z(TU) : B \rightarrow A]$ | | | $[e' \rightarrow], 9., 12.$ |
| 14. $[(XY) (Z(TU))] : A]$ | | | $[e' \rightarrow], 5., 13.$ |

En revenant aux substitutions $[A := P]$, $[B := N]$ et $[C := NC]$ et en considérant que les expressions applicatives X, Y, Z, T et U sont respectivement tout', garçon', aime', le', football', la preuve sur les types nous permet de construire directement l'expression applicative:

$((\text{tout}' \text{ garçon}') (\text{aime}' (\text{le}' \text{ football'})))$

qui est sous-jacente à la phrase: *tout garçon aime le football*. Le type $'(NC \rightarrow ((N \rightarrow P) \rightarrow P) \times NC \times (N \rightarrow (N \rightarrow P)) \times (NC \rightarrow N) \times P'$ est une *spécification syntaxique* de cette phrase. La spécification syntaxique est insuffisante pour construire l'expression applicative sous-jacente, c'est-à-dire l'interprétation fonctionnelle de la phrase: il faut posséder une preuve que la spécification syntaxique est cohérente et réductible au type P des phrases. A par-

tir de cette preuve, on peut construire la représentation applicative sous-jacente à l'énoncé.

4.4. Théorie des types de Martin-Löf

L'analogie de Curry-Howard justifie les bases de l'approche intuitionniste de Martin-Löf qui identifie les types et les propositions: pour Martin-Löf, une proposition n'est pas le nom d'une «valeur de vérité» mais une proposition est un type dont les instances sont les preuves de cette proposition. Dans cette acception, une proposition A est vraie si et seulement si on peut exhiber une preuve de cette proposition A , autrement dit lorsque l'ensemble des preuves de la proposition A n'est pas vide. L'analogie entre types et propositions permet de donner également une signification fonctionnelle à l'implication propositionnelle. En effet, dans l'approche constructiviste de Martin-Löf, l'implication entre propositions, que nous avons désignée par: ' $A \rightarrow B$ ', s'interprète comme suit: la signification de la proposition ' $A \rightarrow B$ ' est donnée par une procédure fonctionnelle qui doit construire une preuve de la proposition B à partir d'une preuve de la proposition A . La théorie de Martin-Löf développe donc l'analogie de Curry-Howard en l'étendant aux autres connecteurs logiques ('&' et 'V'), à la négation (interprétée, comme dans l'approche intuitionniste de Heyting par l'absurdité) et aux quantificateurs d'existence et d'universalité. *Les types sont des spécifications de programmes.* A un type donné peut correspondre aucun programme (lorsque la spécification n'est pas correcte) ou un ou plusieurs programmes qui sont «équivalents en extension», puisqu'ils construisent le résultat ce qui est spécifié par le type, mais qui diffèrent «en intension» puisqu'ils construisent le résultat par des chemins différents. Les preuves que ces types sont des théorèmes du calcul inférentiel sur les types sont des instances de ces types; ces instances sont alors représentables par des λ -expressions, construites directement à partir de ces preuves; ces λ -expressions codent donc les preuves associées à un type, elles décrivent également les programmes dont le type est une spécification.

5. Composition des unités linguistiques

Les unités linguistiques, étant des opérateurs, sont composables entre elles. Les unités linguistiques étant, au moins pour certaines d'entre elles, des opérateurs typés, elles sont éventuellement composables entre elles, tout comme deux fonctions X et Y sont composables entre elles lorsque l'image de l'une est dans le domaine de définition de l'autre. Les combinateurs de la logique combinatoire étant des opérateurs abstraits, dont l'action intrinsèque est indépendante de tout domaine d'interprétation, ils servent à déterminer des modes de composition qui généralisent la composition entre fonctions. Un combinateur permet ainsi de construire des *opérateurs complexes* à partir d'opérateurs plus élémentaires. Chaque combinateur peut être représenté par une λ -expression close (c'est-à-dire sans variables libres) du λ -calcul de Church. Chaque combinateur peut aussi être inséré directement dans le système applicatif de la logique combinatoire de Curry au moyen de règles qui spécifient exactement son action sur les opérandes sur lesquelles il porte. Comme toutes les constantes logiques (flèche ou produit cartésien), la signification d'un combinateur est précisée par des règles d'élimination et d'introduction. Donnons quelques exemples de combinateurs.

5.1. Combinateurs élémentaires

Notation: Dans ce qui suit, les symboles $X, Y, Z, \dots, x, y, z, \dots$ désignent des expressions applicatives quelconques qui représentent des opérateurs. Comme un combinateur construit un opérateur complexe à partir d'autres opérateurs, les symboles X, Y, Z désignent les opérateurs qui sont constitutifs de l'opérateur complexe construit par un combinateur et les symboles $x, y, z, \dots$ désignent les opérandes de l'opérateur complexe.

Chaque combinateur est défini soit par une λ -expression close soit par une règle d'élimination et une règle d'introduction. Nous donnons pour les combinateurs élémentaires les

λ -expressions associées et les règles d'élimination et d'introduction.

Combinateur I. Le plus simple des combinateurs est le combinateur d'identité **I**.

$$\mathbf{I} =_{\text{def}} \lambda X. \lambda x. Xx$$

$$[\text{e-I}] \quad \frac{(\mathbf{IX})x}{Xx}$$

$$[\text{i-I}] \quad \frac{Xx}{(\mathbf{IX})x}$$

L'opérateur **IX** est construit à partir de l'opérateur **X**; l'action de **IX** sur un opérande x est la même que celle de **X**.

Combinateur B. Le combinateur **B** exprime un mode de composition qui s'interprète comme la simple composition de fonctions lorsque les opérateurs sont interprétés par des fonctions ensemblistes.

$$\mathbf{B} =_{\text{def}} \lambda X. \lambda Y. \lambda x. f(gx).$$

$$[\text{e-B}] \quad \frac{(\mathbf{BXY})x}{X(Yx)}$$

$$[\text{i-B}] \quad \frac{X(Yx)}{(\mathbf{BXY})x}$$

L'opérateur **BXY** est un opérateur complexe obtenu en composant les deux opérateurs **X** et **Y** entre eux. L'action de cet opérateur complexe **BXY** sur un opérande x est précisée par les règles d'élimination et d'introduction. Lorsque les opérateurs **X** et **Y** sont interprétés comme des fonctions ensemblistes, le combinateur **B** représente exactement la composition des fonctions, d'où: $X \circ Y = \mathbf{BXY}$.

Combinateur S. Le combinateur **S** est une généralisation de **B**, il compose deux opérateurs **X** et **Y**, d'où l'opérateur complexe **SXY**:

$$S =_{\text{def}} \lambda X. \lambda Y. \lambda x. Xx(Yx)$$

$$\begin{array}{ccc} & (SXY)_x & Xx(Yx) \\ [e-S] & \frac{\quad}{\quad} & [i-S] \frac{\quad}{\quad} \\ & Xx(Yx) & (SXY)_x \end{array}$$

L'opérateur complexe SXY , lorsqu'il a pour opérande x , donne le même résultat que l'action de X sur l'argument x et sur le résultat de l'opérateur Y sur x .

Combinateur W. Le combinateur de diagonalisation W duplique l'argument d'un opérateur binaire X .

$$W =_{\text{def}} \lambda X. \lambda x. Xxx$$

$$\begin{array}{ccc} & (WX)_x & Xxx \\ [e-W] & \frac{\quad}{\quad} & [i-W] \frac{\quad}{\quad} \\ & Xxx & (WX)_x \end{array}$$

L'opérateur WX est un opérateur complexe construit à partir de l'opérateur X ; l'action de W consiste à dupliquer un même argument x de X (processus de diagonalisation).

Combinateur C. Le combinateur C de conversion est défini comme suit:

$$C =_{\text{def}} \lambda X. \lambda x. \lambda y. Xyx$$

$$\begin{array}{ccc} & (CX)_{xy} & Xyx \\ [e-C] & \frac{\quad}{\quad} & [i-C] \frac{\quad}{\quad} \\ & Xyx & (CX)_{xy} \end{array}$$

L'opérateur (CX) est un opérateur complexe (le converse de l'opérateur X); l'action de CX est précisée par les règles: si X opère sur les opérandes x puis y dans cet ordre applicatif, alors l'opérateur (CX) opère sur les opérandes y puis x dans cet ordre applicatif, autrement dit (CX) opère sur les arguments permutés de X .

Combinateur C*. Le combinateur C* (dit de “lifting”) est défini comme suit:

$$C^* =_{\text{def}} \lambda Y. \lambda X. YX$$

$$\begin{array}{ccc} & (C^*X)Y & YX \\ [e-C^*] & \frac{}{} & [i-C^*] \frac{}{} \\ & YX & (C^*X)Y \end{array}$$

L'opérateur (C*X) est un opérateur complexe construit à partir de X; son rôle opératoire est tel que si X était opérande par rapport à l'opérateur Y alors l'opérateur Y devient opérande de l'opérateur (C*X).

Combinateur K. Le combinateur K (création d'argument fictif) est défini par:

$$K =_{\text{def}} \lambda X. \lambda x. X$$

$$\begin{array}{ccc} & (KX)x & X \\ [e-K] & \frac{}{} & [i-K] \frac{}{} \\ & X & (KX)x \end{array}$$

Combinateur Φ. Le combinateur Φ (intrication d'opérateurs agissant en parallèle) est défini par:

$$\Phi =_{\text{def}} \lambda X. \lambda Y. \lambda Z. \lambda u. X(Yu)(Zu)$$

$$\begin{array}{ccc} & \Phi XYZu & X(Yu)(Zu) \\ [e-\Phi] & \frac{}{} & [i-\Phi] \frac{}{} \\ & X(Yu)(Zu) & \Phi XYZu \end{array}$$

Chaque combinateur introduit une relation *β-réduction élémentaire* (et sa converse une *β-expansion*) entre l'expression du programme de combinaison spécifié par le combinateur et le résultat de l'exécution du programme, autrement dit entre l'expression *avec* combinateur et le résultat de l'action du combinateur *sans* ce combinateur. Nous avons par exemple, les β-réduc-

tions (β -expansions) élémentaires suivantes, associées aux combinateurs élémentaires précédents:

IX	$\beta \Rightarrow X$	X	$\Leftarrow \beta$	IX
BXYz	$\beta \Rightarrow X(Yz)$	X(Yz)	$\Leftarrow \beta$	BXYz
WXy	$\beta \Rightarrow Xyy$	Xyy	$\Leftarrow \beta$	WXy
CXyz	$\beta \Rightarrow Xzy$	Xzy	$\Leftarrow \beta$	CXyz
C*XY	$\beta \Rightarrow YX$	YX	$\Leftarrow \beta$	C*XY
SXYz	$\beta \Rightarrow Xz(Yz)$	Xz(Yz)	$\Leftarrow \beta$	SXYz
$\Phi XYZu$	$\beta \Rightarrow X(Yu)(Zu)$	$X(Yu)(Zu)$	$\Leftarrow \beta$	$\Phi XYZu$

La fermeture réflexive et transitive de la relation engendrée par les β -réductions (respectivement β -expansions) élémentaires constitue la relation de β -réduction (respectivement de β -expansion). Cette relation de β -réduction (et de β -expansion) structure l'ensemble des expressions applicatives avec combinateurs.

5.2. Propriétés des combinateurs

Les combinateurs ne sont pas tous indépendants. Par exemple, on a: $[I = BCC]$. Cette identité entre combinateurs signifie que «le converse du converse est la même chose que l'identité». En effet:

- | | |
|-------------------|--------------|
| 1. BCCXyx | hyp. |
| 2. C(CX)yx | $[e-B]$, 1. |
| 3. CXxy | $[e-C]$, 2. |
| 4. Xyx | $[e-C]$, 3. |
| 5. IXyx | $[i-I]$, 4. |

Il existe une infinité dénombrable de combinateurs puisqu'il est possible de les faire agir les uns sur les autres par l'application et donc de construire des combinateurs dérivés. Par

exemple, on peut définir par récurrence les combinateurs B^2 , plus généralement B^n , comme suit $[B^2 =_{\text{def}} \mathbf{BBB}]$, $[B^n =_{\text{def}} \mathbf{BBB}^{(n-1)}]$. On démontre que les actions de ces combinateurs sont définies par les réductions:

$$\begin{array}{ll} \mathbf{BXYZ}_1 & \Rightarrow X(YZ_1) \\ \mathbf{B^2XYZ}_1Z_2 & \Rightarrow X(YZ_1Z_2) \\ \mathbf{B^nXYZ}_1Z_2 \dots Z_n & \Rightarrow X(YZ_1Z_2 \dots Z_n) \end{array}$$

On démontre qu'il est possible de définir tous les combinateurs à partir de combinateurs de base. On peut prendre pour combinateurs de base les combinateurs $\mathbf{S}$ et $\mathbf{K}$, les autres combinateurs étant tous définissables en termes de $\mathbf{S}$ et $\mathbf{K}$. Ainsi, nous avons la définition de $\mathbf{I}$: $[\mathbf{I} =_{\text{def}} \mathbf{SKK}]$. En effet:

- | | |
|----------------------|-----------------------|
| 1. $\mathbf{SKKx}$ | hyp. |
| 2. $\mathbf{Kx(Kx)}$ | $[\mathbf{c-S}]$, 1. |
| 3. $\mathbf{x}$ | $[\mathbf{c-K}]$, 2. |
| 4. $\mathbf{Ix}$ | $[\mathbf{i-I}]$, 3. |

L'utilisation des combinateurs permet de *définir des propriétés intrinsèques d'opérateurs particuliers* sans faire appel à des variables. Ainsi, lorsque $\mathbf{X}$ est un opérateur binaire, exprimer qu'il est *intrinsèquement commutatif*, c'est affirmer que les actions de $\mathbf{X}$ et de son converse $\mathbf{CX}$ conduisent à des *résultats qui ont même signification*, en posant, *sans faire appel à des variables*, l'égalité intrinsèque relative à $\mathbf{X}$: $[\mathbf{CX} =_{\text{def}} \mathbf{X}]$. Donnons, à titre d'exemple, les formulations intrinsèques de la commutativité et de l'associativité de l'opérateur binaire, noté '+', défini aussi bien sur le domaine des entiers naturels que sur celui des nombres réels, en laissant au lecteur le soin d'en vérifier l'exactitude:

$$\begin{array}{l} [\mathbf{C+} = +] \\ [\mathbf{BC(BC(W(BB^2C)))+} = \mathbf{WB^2+}] \end{array}$$

Les combinateurs étant des opérateurs, ils sont typés non par des types constants mais par des *schémas de type*. Dans un contexte donné, le combinateur acquiert un type déterminé qui dépend directement du type assigné aux opérands.

5.3. Schémas de type des combinateurs

Calculons par exemple le schéma de type du combinateur **B**. Avec les trois hypothèses: $[X(Yzx) : \gamma]$, $[Yz : \beta]$ et $[z : \alpha]$, dont on peut s'abstraire, nous avons les deux arbres suivants:

$$\begin{array}{c}
 \begin{array}{c}
 [Y : \alpha \rightarrow \beta] \quad [z : \alpha] \text{ (hyp. 3)} \\
 \hline
 [Yz : \beta] \text{ (hyp. 2)} \\
 \hline
 [X(Yz) : \gamma] \text{ (hyp. 1)} \\
 \hline
 \text{=def } [BXYz : \gamma]
 \end{array} \\
 \hline
 \begin{array}{c}
 [BXY : \alpha \rightarrow \gamma] \quad [z : \alpha] \\
 \hline
 [BX : (\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma)] \quad [Y : \alpha \rightarrow \beta] \\
 \hline
 [B : (\beta \rightarrow \gamma) \rightarrow ((\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma))] \quad [X : \beta \rightarrow \gamma]
 \end{array}
 \end{array}$$

d'où le schéma de type assigné au combinateur **B**:

$$[B : (\beta \rightarrow \gamma) \rightarrow ((\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma))]$$

L'algorithme consiste à analyser la construction de l'expression qui définit l'action du combinateur, à savoir $X(Yz)$, en décomposant, du sommet vers les feuilles, cette expression en faisant apparaître l'opérateur et l'opérande (arbre du haut). On assigne un type (en utilisant la règle **[App_t]**) à partir de types assignés par hypothèse à certaines expressions. On analyse ensuite la construction du programme de construction défini par le combinateur, à savoir **BXY**, en assignant un type par abstraction, ce qui revient à se libérer progressivement des hypothèses introduites (arbre du bas).

Par la même méthode, calculons le schéma de type du combinateur **S**.

$$[X : \alpha \rightarrow (\beta \rightarrow \gamma)] \quad [z : \alpha] \text{ (hyp. 3)} \quad [Y : \alpha \rightarrow \beta] \quad [z : \alpha] \text{ (hyp. 3)}$$

$$[Xz : \beta \rightarrow \gamma]$$

$$[Yz : \beta] \text{ (hyp. 2)}$$

$$[Xz(Yx) : \gamma] \text{ (hyp. 1)}$$

$$\text{def}$$

$$[SXYz : \gamma]$$

$$[SXY : \alpha \rightarrow \gamma]$$

$$[z : \alpha]$$

$$[SX : (\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma)]$$

$$[Y : \alpha \rightarrow \beta]$$

$$[S : (\alpha \rightarrow (\beta \rightarrow \gamma)) \rightarrow ((\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma))] \quad [X : \alpha \rightarrow (\beta \rightarrow \gamma)]$$

Le schéma de type du combinateur **S** est donc :

$$[S : (\alpha \rightarrow (\beta \rightarrow \gamma)) \rightarrow ((\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma))]$$

D'après l'isomorphisme de currification $[CUR_t]$, nous avons l'équivalence entre les schémas de type pour respectivement le combinateur **B** et pour **S** :

$$((\beta \rightarrow \gamma) \times (\alpha \rightarrow \beta)) \rightarrow (\alpha \rightarrow \gamma) = (\beta \rightarrow \gamma) \rightarrow ((\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma))$$

$$((\alpha \times \beta \rightarrow \gamma) \times (\alpha \rightarrow \beta)) \rightarrow (\alpha \rightarrow \gamma) = (\alpha \rightarrow (\beta \rightarrow \gamma)) \rightarrow ((\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \gamma))$$

La lecture de ces schémas de type et des équivalences donne des informations sur la syntaxe des combinateur **B** et **S**. Le combinateur **B** est un opérateur binaire qui compose entre eux deux opérateurs de types respectifs $\beta \rightarrow \gamma$ et $\alpha \rightarrow \beta$ de façon à construire un opérateur complexe de type $\alpha \rightarrow \gamma$. Le combinateur **S** est un opérateur qui compose entre eux deux opérateurs, le premier est

binaire $\alpha \times \beta \rightarrow \gamma$, le second est unaire de type $\alpha \rightarrow \beta$; le résultat est un opérateur complexe de type $\alpha \rightarrow \gamma$.

Lorsqu'on interprète les types par des propositions, on remarque que les schémas de type de **B** et de **S** correspondent aux deux théorèmes, dont on a donné plus haut une démonstration, du Calcul propositionnel:

$$\begin{aligned} & (B \rightarrow C) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C)) \\ & (A \rightarrow (B \rightarrow C)) \rightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C)) \end{aligned}$$

On démontre que tout schéma de type assigné à un combinateur de la logique combinatoire est l'analogue d'un théorème du Calcul propositionnel.

5.4. Grammaire Catégorielle Combinatoire Applicative (GCCA)

Les combinateurs nous permettent de construire l'interprétation fonctionnelle d'un énoncé sous forme d'une représentation applicative normalisée sous-jacente à cet énoncé. En effet, les unités linguistiques étant des opérateurs, ils se voient assignés des types fonctionnels qui sont éventuellement composables entre eux puisque certains opérateurs, comme les fonctions, sont composables. Une extension des Grammaires Catégorielles classiques (avec le Calcul de Lambek notamment) a consisté à introduire des transformations sur les types (par exemple, l'opération de «type raising» (comme: $\alpha \Rightarrow \beta / (\beta \backslash \alpha)$) et des compositions entre types fonctionnels composables (comme: $(\gamma / \beta) \times (\beta / \alpha) \Rightarrow \gamma / \alpha$). Ces compositions entre les types d'unités linguistiques et les transformations de types sont justifiées lorsqu'on remarque que ces unités typées sont composées ou transformées *en tant qu'opérateurs*. La Grammaire Catégorielle Combinatoire et Applicative (GCCA) étend les Grammaires Catégorielles classiques en: 1°) exploitant systématiquement cette remarque ce qui revient à effectuer parallèlement des opérations sur les types qui guident les opérations sur les unités linguistiques; 2°) en donnant une stratégie heuristique, sous forme de méta-règles, qui permettent de trouver une preuve de réduction du

type syntaxique au type des phrases. En effet, nous avons vu que dès que l'on avait obtenu une preuve de ce que le type de la suite, avec types assignés, était réductible au type des phrases (vérification de la correction syntaxique ou de la cohérence de la spécification syntaxique), on pouvait en déduire immédiatement la construction de l'expression applicative associée à la suite analysée. Cependant, l'obtention de la preuve n'est pas immédiate: quelle doit être la procédure que l'on doit suivre pour trouver cette preuve? Les métarègles de la GCCA (thèse de I. Biskri) donnent des éléments de réponse à cette question. Par ailleurs, lorsqu'on procède à l'analyse d'un énoncé, on doit passer progressivement d'une *expression concaténée typée*, c'est-à-dire à la présentation «en surface» de l'énoncé, à une *expression applicative typée*, c'est-à-dire l'interprétation fonctionnelle de l'énoncé. Pour cela, on introduit «localement» dans un premier temps des combinateurs qui correspondent à des compositions de types ou à des transformations de types nécessaires pour que l'analyse syntaxique puisse conclure. Ensuite, dans un second temps, la «normalisation» de l'expression applicative ainsi construite conduit, par une succession d'élimination des combinateurs introduits «localement», à construire, par des β -réductions, une représentation applicative normalisée, appelée forme normale de l'énoncé analysé. On sait que, d'après le théorème de Church-Rosser, cette forme normale est *unique* et, d'après le théorème de normalisation, qu'elle *existe*: toute expression applicative bien typée a une forme normale.

Pour résumer, on peut dire que les règles d'inférence sur les types, les introductions «locales» de combinateurs et les β -réductions entre expressions applicatives constituent un «univers de calcul sur les expressions applicatives typées» dans lequel sont situées les preuves de réduction (réduction entre types; normalisation) que l'on cherche. Les métarègles sont utilisées pour trouver et construire, dans un certain nombre de cas, une preuve particulière de réduction sur les types, cette preuve est codée par une expression applicative que l'on normalisera ensuite pour y trouver la forme normale.

Pour illustrer notre propos, donnons un échantillon de règles extraites de la GCCA (voir la totalité des règles dans: Biskri,

Desclés, Biskri 1996). Aux règles de vérification syntaxique de la Grammaire Catégorielle Combinatoire (GCC) de M. Steedman, on associe canoniquement, dans la GCCA, des règles de construction sémantique qui introduisent «localement» des combinateurs.

règle de vérification syntaxique (GCC: Steedman)

$$\frac{(\beta/\alpha) \times \alpha}{\longrightarrow} \rightarrow [>]$$

β

$$\frac{(\alpha/\beta) \times (\beta/\gamma)}{\longrightarrow} \rightarrow [B]$$

α/γ

$$\frac{\alpha}{\longrightarrow} \rightarrow [T]$$

$\beta/(\beta\alpha)$

règle de construction sémantique (GCCA: Desclés-Biskri)

$$\frac{[X : (\beta/\alpha)] \times [Y : \alpha]}{\longrightarrow} \rightarrow [>]$$

$[XY : \beta]$

$$\frac{[X : (\alpha/\beta)] \times [Y : (\beta/\gamma)]}{\longrightarrow} \rightarrow [B]$$

$[BXY : \alpha/\gamma]$

$$\frac{[X : \alpha]}{\longrightarrow} \rightarrow [C*]$$

$[C*X : \beta/(\beta\alpha)]$

Les prémisses dans les règles de construction syntaxique sont des conjonctions d'expressions typées en relation biunivoque avec une juxtaposition d'expressions linguistiques (donc non applicatives) mais les conclusions sont des expressions applicatives typées et «locales».

Nous allons présenter un exemple d'analyse d'un énoncé dans le cadre des GCCA, en reprenant l'analyse syntaxique et sémantique de l'énoncé:

Tout garçon aime le football

que l'on comparera à l'analyse du même énoncé effectuée dans un précédent paragraphe (voir § 4). En assignant des types syntaxiques aux unités linguistiques élémentaires de cet énoncé, comme nous l'avons déjà effectué, nous avons obtenu une expression concaténée (ec) sous forme d'une juxtaposition

d'unités linguistiques (ici des mots) avec types assignés, l'opération de juxtaposition étant désignée par 'x' :

- (ec) [tout :((P/(P\N))/NC)] x [garçon :NC] x [aime :((P\N)/N):] x
[le :(N/NC)] x [football :NC]

Les types assignés dans (ec) sont les types constitutifs du type τ de (ec):

$$(\tau) \quad \tau = ((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC$$

Le calcul inférentiel sur les seuls types syntaxiques, avec les règles précédentes (inférence (I)), conduit à vérifier que l'expression est syntaxiquement «bien formée» puisque de type P:

$$\begin{array}{l} ((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC \\ \hline \rightarrow [>] \\ P/(P\backslash N) \\ \hline \rightarrow [B] \\ P/N \\ \hline \rightarrow [B] \\ P/NC \\ \hline \rightarrow [>] \\ P \end{array}$$

Ce calcul est la preuve que relation (R):

$$(R) \quad ((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC \Rightarrow P$$

est un théorème du système inférentiel des types. Des méta-règles (que nous ne donnons pas ici) ont pour but de déclencher des changements de types à certains endroits et de composer les types à d'autres endroits. Ces métarègles nous conduisent à construire, parmi toutes les preuves possibles, une preuve particulière en adoptant une stratégie «quasi-incrémentale», c'est-à-dire en procédant, autant que faire se peut, à une analyse de la suite textuelle «de gauche à droite». Comme on peut le voir, la preuve proposée est différente de celle qui a été donnée dans un précédent paragraphe, cette dernière procédant par un balayage entier de la suite textuelle avec retour en arrière (c'est-à-dire

avec «backtracking»). A cette *inférence sur les types*, les règles de construction sémantique de la GCCA associent canoniquement *un calcul sur les unités linguistiques typées*: à la règle d'application [$>$] correspond une simple application d'un opérateur à un opérande, à l'utilisation inférentielle de la règle [B] sur les types correspond une composition fonctionnelle des unités linguistiques, considérées comme des opérateurs, ce qui conduit à l'introduction «locale» du combinateur **B**: le combinateur note dans le calcul lui-même la composition fonctionnelle des unités linguistiques. Ce calcul, de même structure que le calcul inférentiel sur les types, est effectué à l'aide des règles de construction sémantique. Nous avons:

1. [tout : (P/(PN))/NC] x [garçon : NC] x [aime : (PN)/N:] x
[le : N/NC] x [NC : football] hyp.
2. [(tout garçon): P/(PN)] x [aime : ((PN)/N)] x
[le : (N/NC) : le] x [football : NC] [>], 1.
3. [B(tout garçon) (aime)] x [le : N/NC] x [football : NC] [B], 2.
4. [B(B(tout garçon) (aime)) (le) : N/NC] x [football : NC] [B], 3.
5. [B(B(tout garçon) (aime)) (le) (football): P] [>], 4.

Ce processus de calcul (C) construit *progressivement une expression applicative globale* (ea) de type P à partir de l'*expression concaténée typée* (ec) de type τ , en introduisant «localement» des combinateurs et en composant «localement» les unités par l'opération d'application:

(ea) **B**(B(tout garçon) (aime)) (le) (football)

Toute l'information pour construire l'expression (ea) est entièrement déterminée par: (i) l'assignation de types aux unités linguistiques élémentaires de (ec) et donc par le type τ , et par (ii) l'inférence (I) du type τ au type P. Comme l'expression typée (ec) encode à la fois les informations données par (i), la construction (C) de (ea) à partir de (ec) a été effectuée, au moyen des règles de la GCCA, en se laissant entièrement guider par l'inférence (I) sur les types.

L'information syntaxique sur l'énoncé est entièrement fournie par le type τ . *Ce type τ code les contraintes d'agencement syntaxique des unités linguistiques avec des types assignés.* Si le

type τ constitue une spécification syntaxique de l'expression applicative (ea), sa seule connaissance n'est pas suffisante pour fournir la méthode de construction de l'expression applicative (ea) associée. Dans le cadre de la GCC de Steedman, avec les seules règles syntaxiques comme $[>]$, $[B]$ et $[T]$, la construction de l'expression applicative sous-jacente à l'énoncé doit être effectuée dans un second temps en se donnant des règles d'interprétation dans le λ -calcul et des principes d'unification. En revanche, les règles constructives de la GCCA nous permettent de vérifier, par inférence, d'une part que le type τ est «bien formé» et réductible au type P des phrases, autrement dit, de vérifier que les contraintes d'agencement syntaxique codées par τ sont cohérentes et, si tel est le cas, d'autre part d'en déduire ensuite la construction de l'expression applicative (ea), qui donne l'interprétation fonctionnelle de l'énoncé. La construction (C) est entièrement guidée par la structure de l'inférence (I). Il apparaît alors que l'expression applicative (ea) est une certaine instance du type P, instance dont la construction effective est entièrement déterminée par (i) la spécification syntaxique τ et (ii) une certaine preuve que la relation inférentielle (R) ' $\tau \Rightarrow P$ ' est bien un théorème du système inférentiel sur les types. D'autres instances de type P avec une spécification syntaxique de type τ seraient construites de la même façon. Par exemple, les expressions applicatives:

(ea') $B(B(\text{toute fille}) (\text{aime})) (la) (\text{cuisine})$

(ea'') $B(B(\text{tout surveillant}) (\text{respecte})) (la) (\text{réglementation})$

sont des instances du type P avec une spécification syntaxique de type τ : les constructions de ces expressions sont analogues, seules les unités linguistiques avec une même assignation des types constitutifs de τ , changent.

Alors que l'expression (ec) n'est pas une expression applicative mais une *expression concaténée* typée – de type τ –, l'expression (ea) est une *expression applicative* typée – du type P obtenue par inférence à partir de la spécification syntaxique fournie par le type τ . Or, l'analyse de la structure applicative (ea) montre que cette expression applicative exprime *directement son propre mode de construction*, ce qui n'était pas le cas

de l'expression concaténée (ec). Ainsi, *la seule connaissance de la structure applicative d'une expression applicative de type P indique comment cette expression est construite*. Une telle expression applicative code donc biunivoquement sa propre construction.

Maintenant, étant donnée une expression applicative de type P avec des types syntaxiques assignés aux unités linguistiques qui la composent, on peut chercher à vérifier que la relation inférentielle entre le type syntaxique τ , déterminé par les seuls types assignés aux unités élémentaires, et le type P, est bien un théorème du système inférentiel sur les types; pour cela, on procède aux inférences élémentaires associées à chacune des étapes de la construction de l'expression applicative. Pour l'expression applicative (ea) qui nous sert d'exemple conducteur, nous avons:

1. $[B(B(\text{tout garçon}) (\text{aime})) (\text{le}) (\text{football}) : P]$ hyp.
2. $[B(B(\text{tout garçon}) (\text{aime})) (\text{le}) : P/NC] \times [\text{football} : NC]$ [$>$], 1.
3. $[B(\text{tout garçon}) (\text{aime}) : P/N] \times [\text{le} : N/NC] \times [\text{football} : NC]$ [B], 2.
4. $[(\text{tout garçon}) : P/(P \setminus N)] \times [\text{aime} : (P \setminus N)/N] \times [\text{le} : N/NC] \times [\text{football} : NC]$ [B], 3.
5. $[\text{tout} : (P/(P \setminus N))/NC] \times [\text{garçon} : NC] \times [\text{aime} : (P \setminus N)/N] \times [\text{le} : N/NC] \times [\text{football} : NC]$ [$>$], 4.

Nous en déduisons le théorème sur les types: $P \leq \tau$. L'expression applicative (ea) de type P peut donc bien être également considérée comme une instance de type τ . On remarquera, au passage, que les schémas de types des occurrences du combinateur **B** seraient instanciés en tenant compte des types des opérateurs sur lesquels il agit.

Revenons maintenant à l'expression applicative (ea):

(ea) **B**(**B**(tout garçon) (aime)) (le) (football)

Cette expression contient des combinateurs que l'on peut chercher à éliminer. Nous en déduisons la β -réduction (considérée comme un processus de normalisation) suivante entre expressions applicatives de même type:

Processus de normalisation:

- | | |
|--|-------------|
| 1. $[B(B(\text{tout garçon}) (\text{aime})) (\text{le}) (\text{football}): P]$ | hyp. |
| 2. $[(B(\text{tout garçon}) (\text{aime})) (\text{le football}) : P]$ | $[e-B], 1.$ |
| 3. $[(\text{tout garçon}) (\text{aime} (\text{le football})) : P]$ | $[e-B], 2.$ |

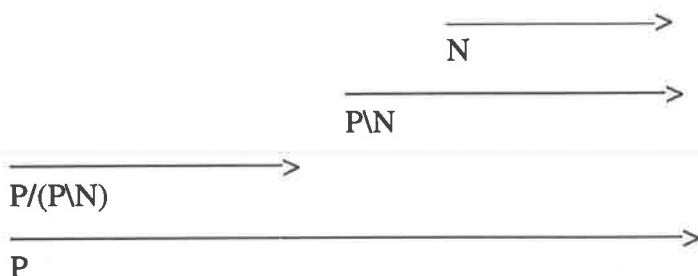
Le dernier pas est l'expression applicative:

(fn) (tout garçon) (aime (le football))

constitue de ce fait une «forme normale» (au sens technique de la logique combinatoire) de l'expression applicative (ea) car elle ne peut plus être réduite. Cette expression est aussi la forme normale sous-jacente à l'énoncé analysé *tout garçon aime le football*.

Remarquons que, pour les mêmes types assignés aux mêmes unités linguistiques élémentaires, la forme normale (fn) décrit un mode d'agencement applicatif différent de celui qui est exprimé par (ea); *cette forme normale code donc un autre mode de construction applicative* que celui qui est codé par (ea). Il s'ensuit que la forme normale (fn) correspond à une autre inférence entre les types; cette inférence est directement déduite de la construction applicative (fn). Nous avons:

$$((P/(P\backslash N))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC$$



On voit donc que les expressions applicatives (ea) et (fn) peuvent, toujours pour les mêmes assignations de types aux unités linguistiques élémentaires, être considérées comme deux instances différentes du même type τ . Elles expriment deux *programmes applicatifs distincts de même spécification syntaxique* τ . Cependant, ces deux programmes applicatifs ne sont pas

indépendants puisque le programme (ea) est réductible au programme (fn):

[B(B(tout garçon) (aime)) (le) (football): P]

$\beta \Rightarrow [(tout\ garçon) (aime\ (le\ football))]: P]$.

Le programme (fn) exprime la construction applicative «normale» sous-jacente à l'énoncé (ec). Il a été obtenu au moyen d'un «processus de normalisation» portant sur (ea). Les deux expressions applicatives (ea) et (fn) représentent deux interprétations fonctionnelles du même énoncé *tout garçon aime le football*. Puisque ces deux interprétations sont interliées par une β -réduction, elles expriment une *signification commune* dont un représentant canonique est justement la forme normale (fn). Ces programmes diffèrent uniquement par des modes de construction différents des mêmes unités linguistiques agencées. L'expression (ea) est un programme construit par une «stratégie incrémentale» à partir des informations contenues dans le type τ , c'est-à-dire en allant «de gauche à droite» et en composant le syntagme quantifié (*tout garçon*), considéré comme un opérateur, avec le verbe (*aime*) puis le résultat est de nouveau composé avec l'article (*le*), lui aussi opérateur, et l'expression ainsi obtenue est un opérateur qui est appliqué finalement à l'opérande nominal (*football*). Bien que l'expression (fn) soit de même type, elle exprime un programme de construction qui coïncide directement avec son interprétation fonctionnelle: dans un premier temps, un premier syntagme nominal (*le football*) est construit par application, puis ensuite le syntagme verbal (*aime le football*) est construit en appliquant le prédicat (*aime*) sur le syntagme nominal déjà construit (*le football*); le syntagme quantifié (*tout garçon*) est ensuite appliqué au syntagme verbal. Par conséquent, *ces deux programmes applicatifs* (ea) et (fn) construisent, par des moyens différents, *une même signification*. On dira que ces programmes constructifs diffèrent *en intension*, mais que leurs extensions sont équivalentes.

La construction sémantique de l'énoncé *tout garçon aime le football* est obtenue en «recollant» le processus de normalisation (pas 5 à 7) à la construction (pas 1 à 5) directement calquée sur l'inférence entre types: $\tau \Rightarrow P$:

1. $[\text{tout} : (P/(P\backslash N))/NC] \times [\text{garçon} : NC] \times [\text{aime} : (P\backslash N)/N] :] \times$
 $[le : N/NC] \times [NC : \text{football}]$ hyp.
 2. $[(\text{tout garçon}) : P/(P\backslash N)] \times [\text{aime} : ((P\backslash N)/N)] \times$
 $[le : (N/NC) : le] \times [\text{football} : NC]$ $[>], 1.$
 3. $[B(\text{tout garçon}) (\text{aime})] \times [le : N/NC] \times [\text{football} : NC]$ $[B], 2.$
 4. $[B(B(\text{tout garçon}) (\text{aime})) (le) : N/NC] \times [\text{football} : NC]$ $[B], 3.$
 5. $[B(B(\text{tout garçon}) (\text{aime})) (le) (\text{football}) : P]$ $[>], 4.$
- =====
5. $[B(B(\text{tout garçon}) (\text{aime})) (le) (\text{football}) : P]$ $[>], 4.$
 6. $[(B(\text{tout garçon}) (\text{aime})) (le \text{ football}) : P]$ $[e-B], 5.$
 7. $[(\text{tout garçon}) (\text{aime} (le \text{ football})) : P]$ $[e-B], 6.$

6. Opérateurs illatifs

Nous allons passer à une représentation en «logique illative» (branche de la logique combinatoire obtenue en adjoignant aux combinateurs des opérateurs qui jouent le rôle inférentiel des quantificateurs et des connecteurs dans la logique classique du premier ordre). Introduisons les opérateurs de quantification universelle non restreinte et restreinte Π_1 et Π_2 .

Soient deux expressions applicatives X et Y représentant deux opérateurs unaires, l'expression $\Pi_1 X$ (lire: «tout est X ») et $\Pi_2 XY$ (lire: «tout X est Y » ou «tout ce qui est X est Y ») sont deux expressions applicatives propositionnelles obtenues en appliquant respectivement le quantificateur Π_1 sur X ou le quantificateur Π_2 sur X , puis le résultat sur Y . L'opérateur E indiquant l'existence (dans un univers spécifié d'entités) d'opérandes de l'opérateur X sur lequel E porte (ce qui signifie que le domaine de définition de X n'est pas vide), les règles d'élimination des deux quantificateurs sont données comme suit:

$$\begin{array}{ccc}
 \Pi_1 X, EX & & \Pi_2 XY, Xx \\
 [e-\Pi_1] \frac{\quad}{Xx} & & [e-\Pi_2] \frac{\quad}{Yx}
 \end{array}$$

La première règle signifie que si X est un prédicat (unaire) et si x est un élément (quelconque) qui existe dans le domaine

(univers de discours) de définition de X , ce que l'on représente par EX , alors le prédicat X s'applique à x . La seconde règle signifie que si l'on a la proposition $\Pi_2 XY$ et si l'opérateur X s'applique à x , alors l'opérateur Y s'applique également à x .

Désignons par 'p' le type particulier des expressions propositionnelles. Supposons que X et Y soient des opérateurs de type fonctionnel $t \rightarrow p$ où t est le type des entités individuelles. Ces opérateurs sont donc des constructeurs de propositions. Les types de Π_1 et de Π_2 sont définis comme suit

$$\Pi_1: (t \rightarrow p) \rightarrow p$$

$$\Pi_2: (t \rightarrow p) \rightarrow ((t \rightarrow p) \rightarrow p).$$

Le quantificateur Π_1 est un opérateur qui s'appliquant à un prédicat de type $t \rightarrow p$ construit une proposition de type p . Le quantificateur Π_2 est un opérateur qui en s'appliquant à un prédicat de type $t \rightarrow p$ construit un autre opérateur de type $(t \rightarrow p) \rightarrow p$. Puisque, d'après le principe de currification, nous avons l'équivalence de types:

$$(t \rightarrow p) \times (t \rightarrow p) \rightarrow p = (t \rightarrow p) \rightarrow ((t \rightarrow p) \rightarrow p)$$

nous pouvons dire que le quantificateur Π_2 s'applique à un couple de prédicats pour construire une proposition.

L'expression applicative obtenue précédemment à partir de l'énoncé *tout garçon aime le football*:

$$(\text{tout garçon}) (\text{aime (le football)})$$

est représentée maintenant par une autre expression applicative en utilisant le quantificateur restreint Π_2 :

$$\Pi_2 (\text{garçon}) (\text{aime (le football)}).$$

Nous en déduisons l'inférence suivante entre expressions applicatives de même type P :

$$\Pi_2 (\text{garçon}) (\text{aime (le football)}), \text{garçon (Julien)}$$

$$\text{aime (le football) (Julien)}$$

En effet, nous avons:

- | | |
|---|----------------------|
| 1. Π_2 (garçon) (aime (le football)) | hyp. |
| 2. 2.1. garçon (Julien) | hyp. |
| 3. 2.2. Π_2 (garçon) (aime (le football)) | rep. 1. |
| 4. 2.3. (aime (le football)) (Julien) | $[e-\Pi_2]$, 3., 2. |

En français l'inférence:

«si tout garçon aime le football et si Julien est un garçon,
alors Julien aime le football»

exprime directement l'inférence précédente.

Etant donné un énoncé, dont la spécification syntaxique est exprimée par un type syntaxique τ . Nous avons un parallélisme étroit entre d'une part le choix d'une preuve syntaxique que le type τ se réduit au type P des phrases et d'autre part la construction de l'interprétation fonctionnelle qui est directement guidée par la preuve choisie. A propos de l'exemple traité *tout garçon aime le football*, trouver une preuve de la relation entre types:

$$((P/(PN))/NC) \times NC \times ((P\backslash N)/N) \times (N/NC) \times NC \Rightarrow P$$

conduit directement à trouver une construction d'une expression applicative qui peut elle-même être réduite à sa forme normale. Cette forme normale rend possible le remplacement des quantificateurs linguistiques par des quantificateurs illatifs. Nous avons ainsi:

$$\begin{aligned} [\text{tout}: ((P/(PN))/NC)] \times [\text{garçon}: NC] \times [\text{aime}: (PN)/N] \times [(N/NC): \text{le}] \times [NC: \text{football}] \\ \rightarrow [B(B(\text{tout garçon aime})(\text{le football})): P] \\ \beta \Rightarrow [(\text{tout garçon})(\text{aime}(\text{le football})): P] \\ = [\Pi_2(\text{garçon})(\text{aime}(\text{le football})): P]. \end{aligned}$$

Nous avons des règles analogues aux règles relatives aux quantificateurs universels pour les quantificateurs existentiels Σ_1 et Σ_2 . Soient X et Y deux prédicats unaires. Donnons les règles d'élimination:

$$\begin{array}{c}
 \Sigma_1 X, Xx \Rightarrow Z \\
 [e-\Sigma_1] \frac{\quad}{Z}
 \end{array}
 \qquad
 \begin{array}{c}
 \Sigma_2 XY, Xx \Rightarrow Z, Yx \Rightarrow Z \\
 [e-\Sigma_2] \frac{\quad}{Z}
 \end{array}$$

où 'x' est un élément non spécifié du domaine de définition de X (ou d'un sous domaine commun de définition de X et de Y).

Les règles d'introduction sont:

$$\begin{array}{c}
 Xx \\
 [i-\Sigma_1] \frac{\quad}{\Sigma_1 X}
 \end{array}
 \qquad
 \begin{array}{c}
 \&(Xx)(Yx) \\
 [i-\Sigma_2] \frac{\quad}{\Sigma_2 XY}
 \end{array}$$

Nous utiliserons plus tard ces règles.

Donnons un autre exemple d'énoncé: *Marie admire un certain garçon*. Nous avons l'assignation suivante de types syntaxiques:

[Marie :N] x [admire :(P\N)/N:] x [un:(P/(P\N))/NC] x [garçon :NC].

Nous en déduisons la construction de la forme normale avec un quantificateur existentiel:

1. [Marie :N] x [admire :(P\N)/N:] x [un.:(P/(P\N))/NC] x [garçon :NC]
2. [C*Marie: P/(P\N)] x [admire :(P\N)/N:] x [un :(P/(P\N))/NC] x [garçon :NC]
3. [B(C*Marie)(admire) :P/N] x [un.:(P/(P\N))/NC] x [garçon :N]
4. [B(C*Marie)(admire) :P/N] x [un garçon :(P/(P\N))]
5. [(un garçon) (B(C*Marie)(admire)) :P]
6. [(\Sigma_2 garçon) (B(C*Marie)(admire)) :P]

7. Prédicat complexe

En linguistique, les combinateurs sont les opérateurs qui permettent de construire des «prédicats complexes» à partir de prédicats plus élémentaires.

7.1. Prédicat réflexif

Cette notion de prédicat complexe a été entrevue par G. Geatch pour rendre compte, entre autres, des énoncés avec réflexifs. En effet, prenons l'exemple suivant:

Satan s'aime.

Quel est le rôle sémantique du réflexif *se*? Ce clitique a-t-il un simple rôle anaphorique avec une fonction syntaxique de «complément d'objet» ou bien assume-t-il un autre rôle? Lorsqu'on substitue (à la position près) *Satan* au clitique *se*, on obtient bien la paraphrase:

Satan s'aime -> *Satan_i aime Satan_i*.

On pourrait par conséquent penser que *se* est la trace, en position d'objet, d'une anaphore du sujet syntaxique. Or, deux difficultés surgissent: premièrement, on peut avoir l'anaphorique lui-même en position d'objet avec la co-présence de *se*:

Satan s'aime lui-même

ensuite, avec une quantification universelle, le réflexif *se* n'assume pas la fonction anaphorique de renvoi au sujet syntaxique puisque l'on a pas la paraphrase:

Chacun s'aime $\neq$ *chacun aime chacun*.

La solution entrevue par Geatch, à la suite de recherches logiques menées au moyen âge, consiste à construire le prédicat complexe «aimer soi-même» qui s'applique alors directement au terme sujet. Geatch ne propose pas une formalisation adéquate de cette notion de prédicat complexe. Selon nous, la logique combinatoire est un cadre formel tout désigné pour formaliser cette notion. Le prédicat «aimer soi-même» est construit à partir du prédicat plus élémentaire «aimer» à l'aide du combinateur *W*. Nous le définissons comme suit:

[«aimer-soi-même» =_{def} *W* aimer].

La phrase *Satan s'aime* est analysée par l'application du prédicat complexe «aimer-soi-même» au terme *Satan*, d'où la représentation applicative préfixée (l'opérateur précède l'opérande):

(«aimer-soi-même») (Satan).

La définition du prédicat complexe permet de résoudre les difficultés signalées. En effet:

- | | | |
|----|---|---------------|
| 1. | («aimer-soi-même») (Satan) | hyp. |
| 2. | [«aimer-soi-même» = _{def} W aimer] | def |
| 3. | (W aimer) (Satan) | rempl. 2., 3. |
| 4. | aimer (Satan) (Satan) | [e-W] |

d'où la paraphrase (techniquement, c'est une β -réduction entre expressions applicatives):

(«aime-soi-même») (Satan) $\beta \Rightarrow$ aime Satan Satan

qui correspond à l'inférence suivante: de *Satan s'aime* on peut déduire que *Satan aime Satan*, les deux occurrences de *Satan* ayant pour référence le même individu /Satan/.

Le passage de l'expression concaténée à la forme normale s'effectue comme suit:

- | | | |
|----|--|----------------|
| 1. | [Satan :N] x [se :((P\N)\N)/((P\N)/N)] x [aime :(P\N)/N] | hyp. |
| 2. | [C* Satan :(P/N)/((P\N)\N)] x [se :((P\N)\N)/((P\N)/N)] | |
| | x [aime :(P\N)/N] | [C*] |
| 3. | [B(C* Satan)(se) : (P/N)/((P\N)/N)] x [aime :(P\N)/N] | [B] |
| 4. | [B(C* Satan)(se) (aime) :P] | [>] |
| 5. | [(C* Satan) (se aime) :P] | [e-B] |
| 6. | [(se aime)(Satan) :P] | [e-C*] |
| 7. | [se aime = _{def} W aime] | def |
| 8. | (W aime) (Satan) | rempl., 7., 6. |
| 9. | aime Satan Satan | [e-W] |

Prenons maintenant la quantification universelle:

Chacun s'aime.

Cette phrase est analysée à l'aide de l'opérateur illatif Π_1 ; cet opérateur s'applique au prédicat complexe «aimer-soi-même»

pour construire une proposition, d'où les représentations applicatives équivalentes suivantes:

$$(\text{chacun})(\text{se aime}) = \Pi_1(\mathbf{W}(\text{aimer}))$$

Nous avons en effet le passage de la présentation de l'énoncé sous forme d'une expression concaténée typée à sa forme normale:

- | | |
|---|-----------------------|
| 1. $[\text{chacun} : P/(P\backslash N)] \times [[\text{se} : ((P\backslash N)/((P\backslash N)/N))] \times [\text{aime} : (P\backslash N)/N]$ | |
| 2. $[\mathbf{B}(\text{chacun})(\text{se}) : (P/((P\backslash N)/N))] \times [\text{aime} : (P\backslash N)/N]$ | $[\mathbf{B}]$, 1. |
| 3. $[\mathbf{B}(\text{chacun})(\text{se})(\text{aime}) : P]$ | $[>]$, 2. |
| 4. $[\text{chacun} =_{\text{def}} \Pi_1]$ | def. de chacun |
| 4. $[\mathbf{B}\Pi_1(\text{se})(\text{aime}) : P]$ | rempl. 4., 3. |
| 5. $[\Pi_1(\text{se aime}) : P]$ | $[e-\mathbf{B}]$, 4. |
| 6. $[(\text{se})(\text{aime}) =_{\text{def}} \mathbf{W} \text{ aime}]$ | def. se aime |
| 7. $\Pi_1(\mathbf{W} \text{ aimer})$ | rempl. 6., 5. |
| 8. $\mathbf{E}(\mathbf{W} \text{ aimer})$ | hyp. |
| 9. $(\mathbf{W} \text{ aime})(\text{Satan})$ | $[e-\Pi_1]$, 7., 8. |
| 10. $\text{aime}(\text{Satan})(\text{Satan})$ | $[e-\mathbf{W}]$, 9. |

ce qui correspond à la relation d'inférence en français:

si chacun s'aime, étant donné Satan, alors Satan s'aime et donc Satan aime Satan.

On voit que, dans le contexte de l'énoncé analysé, la *trace linguistique* du combinateur $\mathbf{W}$ de réflexivité est *se* et celle de *chacun* est le quantificateur Π_1 . Dans une quantification restreinte nous avons le même genre d'analyse. Prenons l'énoncé suivant:

Chaque homme s'aime

dont la représentation applicative construite est:

$$\Pi_2(\text{homme})(\mathbf{W} \text{ aimer}).$$

où «homme» et « $\mathbf{W}$ aimer» sont deux prédicats unaires. Nous avons comme dans le calcul précédent:

1. [chaque: (P/(P\N)/NC) x [homme: NC] x [se :((P\N)\N)/((P\N)/N)]
x [aime : (P\N)/N]
 2. [chaque homme : (P/(P\N))] x [se :((P\N)/((P\N)/N)] x [aime : (P\N)/N]
-
3. [chaque homme =_{def} Π_2 (homme)] def.
 4. [se aime =_{def} W aime] def.
-
5. Π_2 (homme) (W aimer)
 6. homme (Satan) hyp.
 7. (W aimer) (Satan) [e- Π_2]
 8. aimer (Satan) (Satan) [e-W]

ce qui correspond à l'inférence:

si chaque homme s'aime et si Satan est un homme alors
Satan s'aime et donc Satan aime Satan.

Avec la notion de prédicat complexe, on peut analyser des énoncés tels que:

Satan a pitié de lui-même
(a-pitié-de-soi)(Satan).

Seul Satan a pitié de lui-même
(est-le-seul-à-avoir-pitié-de-soi) (Satan).

Satan n'a pitié que de lui-même
(n'a-pitié-que-de) (Satan).

Seul, Satan n'a pitié que de lui-même
(est-le-seul-à-n'avoir-pitié-que-de-soi) (Satan).

Chacun n'a pitié que de soi-même
(Chacun)(n'a-pitié-que-de-soi).

Nous n'indiquons pas ici comment on construit les prédicats complexes mais les techniques de représentation à l'aide des combinateurs sont les mêmes que dans les exemples précédents.

7.2. Prédicat avec ellipse

On peut, par la même méthode, rendre ainsi compte de la signification de:

Pierre aime sa femme. Jacques aussi

qui est ambiguë. Cette suite textuelle peut signifier:

soit: *Pierre aime sa propre femme et Jacques aime sa propre femme,*

soit: *Pierre aime sa femme et Jacques aime la femme de Pierre.*

Dans la première interprétation, il faut analyser et représenter le prédicat complexe «aimer-sa-propre-femme» et le combinateur **W** intervient de nouveau dans sa construction. Montrons comment traiter la relation paraphrastique:

1 *Jean aime sa (propre) femme* $\beta \Rightarrow$ *Jean_i aime (la femme de Jean_i)*

Soit t le type des termes et p le type des propositions. Exprimons dans un premier temps l'opérateur unaire «la-mère-de» de type: $t \rightarrow t$. Nous avons la β -réduction suivante:

(est-la-mère-de T) $x \beta \Rightarrow$ la-mère-de T .

En effet, avec les assignations de types:

type (a-pour-mère) = $t \rightarrow (t \rightarrow p)$

type(la-mère-de) = $t \rightarrow t$

type(**K**(la -mère-de)) = $t \rightarrow (t \rightarrow p)$

nous avons le calcul suivant:

- | | | |
|----|---|--------------------|
| 1. | I (est-la-mère-de T) x | hyp. |
| 2. | BI (est-la-mère-de) $T x$ | [i- B], 1. |
| 3. | C (BI (est-la-mère-de)) $x T$ | [i- C], 2. |
| 4. | [a-pour-mère = _{def} C (BI (est-la-mère-de))] | def. |
| 5. | a-pour-mère $x T$ | rempl. 4., 3. |
| 6. | [a-pour-mère = _{def} K (la -mère-de)] | def. |
| 6. | K (la -mère-de) $x T$ | rempl. 6., 5. |
| 7. | la-mère-de T | [e- K], 6. |

Montrons maintenant comment on peut construire le prédicat complexe «aimer-sa-(propre)-mère». En procédant par β -expansions successives (c'est-à-dire par introduction successives de combinateurs), nous obtenons:

- | | |
|--|---------------------------------|
| 1. aime (la-mère-de Jean) Jean | hyp. |
| 2. B aime (la-mère-de) Jean Jean | [i-B], 1. |
| 3. W (B aime (la-mère-de)) Jean | [i-W], 2. |
| 4. BW (B aime) (la-mère-de) Jean | [i-B], 3. |
| 5. B (BW) B aime (la-mère-de) | [i-B], 4. |
| 6. [REFL ₂ = _{def} B (BW) B] | def de REFL ₂ |
| 7. REFL ₂ aime (la-mère-de) Jean | rempl., 5., 6. |
| 8. [aime-sa-(propre)-mère = _{def} REFL ₂ aime (la-mère-de)] | def. du prédicat |
| 9. aime-sa-(propre)-mère Jean | rempl., 7., 8. |

Nous pouvons analyser maintenant:

Chaque homme aime sa (propre) mère.

En effet, nous avons:

- | | |
|--|---------------------------------|
| 1. (chaque homme)(aime-sa-mère) | hyp. |
| 2. Π_2 (homme) (aime-sa-mère) | représent. de 1. |
| 3. 3.1. (homme)(Jean) | hyp. |
| 3.2. Π_2 (homme) (aime-sa-mère) | rappel, 2. |
| 3.3. (aime-sa-mère)(Jean) | [e- Π_2], 3.1., 3.2. |
| 3.4. [aime-sa-mère = _{def} REFL ₂ aime (la-mère-de)] | def. du prédicat |
| 3.5. REFL ₂ aime (la-mère-de)(Jean) | rempl., 3.4., 3.3. |
| 3.6. [REFL ₂ = _{def} B (BW) B] | def de REFL ₂ |
| 3.7. B (BW) B aime (la-mère-de) Jean | rempl., 3.6., 3.5. |
| 3.8. aime (la-mère-de Jean) Jean | [e-B(BW) B] |

c'est-à-dire que nous avons le raisonnement:

Si chaque homme aime sa (propre) mère et si Jean est un homme alors Jean aime sa (propre) mère et donc Jean aime la mère de Jean.

7.3. Prédicat avec anaphore

Donnons un autre exemple où la représentation de la signification peut être traitée par un prédicat complexe. Prenons l'énoncé:

J'ai téléphoné à un électricien, il viendra demain.

L'occurrence de l'anaphorique *il* n'est pas mise pour *un électricien* (selon la théorie du pronom paresseux) puisque en remplaçant le pronom *il* par l'antécédent *un électricien*, nous n'obtenons pas un énoncé de même signification:

J'ai téléphoné à un électricien, un électricien viendra demain.

La signification de l'énoncé avec l'anaphorique *il* peut être révélée par la paraphrase suivante:

J'ai téléphoné à un électricien et cet électricien auquel j'ai téléphoné viendra demain

ce que l'on peut également exprimer dans la logique des prédicats du premier ordre par une quantification existentielle d'une variable d'individu x :

$$[\exists x][(\text{électricien}(x) \ \& \ ((\text{ai-téléphoné-à}(je,x)) \ \& \ \text{vient-demain}(x)))]$$

Nous donnons une solution applicative en considérant le prédicat complexe unaire:

«ai-téléphoné-à-qui-vient-demain»

auquel s'applique l'opérateur qui correspond au syntagme quantifié existentiellement.

Σ_2 (électricien) («ai-téléphoné-à-qui-vient-demain»(je)).

Nous en déduisons, d'après la règle d'élimination du quantificateur existentiel Σ_2 :

1. Σ_2 (électricien) («ai-téléphoné-à-qui-vient-demain»(je))
2. $\&((\text{électricien}(x)) \ (\text{«ai-téléphoné-à-qui-vient-demain»}(je)) \ (x)) \quad [e-\Sigma_2]$

c'est-à-dire:

il existe un certain électricien et c'est à lui que j'ai téléphoné et ce dernier viendra demain.

7.4. Problème des «donkey sentences»

Les «donkey sentences» posent des problèmes de composition des unités linguistiques. En effet, les représentations d'un syntagme nominal avec article indéfini *un* ne sont pas identiques selon les contextes dans lesquelles ce syntagme s'insère. Pour nous en rendre compte, considérons l'énoncé:

(a) *Pedro possède un âne*

dont la représentation classique dans le calcul des prédicats est:

(a') $(\exists x) [\text{âne}(x) \ \& \ \text{possède}(\text{Pedro}, x)]$.

Le syntagme quantifié *un âne* se voit représenter à l'aide d'un quantificateur existentiel. Considérons maintenant l'énoncé suivant:

(b) *Si Pedro possède un âne, alors il [= Pedro] est heureux.*

Le connecteur *si ... alors* est étroitement associé, comme on le sait, à une quantification universelle restreinte. La représentation de cet énoncé dans le calcul des prédicats est donc exprimée par une quantification universelle restreinte:

(b') $(\forall x)[\text{âne}(x) \ \& \ \text{possède}(\text{Pedro}, x) \rightarrow \text{est heureux}(\text{Pedro})]$.

Cependant, dans le calcul des prédicats, nous avons l'équivalence entre cette représentation (b'), avec quantificateur universel, et la représentation suivante (b''), avec quantificateur existentiel:

(b'') $\{(\exists x)[\text{âne}(x) \ \& \ \text{possède}(\text{Pedro}, x)] \rightarrow \text{est heureux}(\text{Pedro})\}$

la portée des opérateurs de quantification étant différente dans (b') et (b''). On voit donc que la composition des significations dans la représentation sémantique est conservée puisque le syntagme nominal *un âne* est toujours représenté (à une équivalence

d'écritures près) à l'aide d'une quantification existentielle dans les deux énoncés (a) et (b). Prenons maintenant un autre contexte avec l'énoncé:

(c) *Si Pedro possède un âne, alors il le [= cet âne] bat.*

Remarquons immédiatement que la «théorie du pronom paresseux» ne fonctionne pas ici puisque si on remplace l'anaphorique *le* par son antécédent *un âne*, on obtient alors un énoncé (d) qui ne conserve plus la signification initiale:

(d) *Si Pedro possède un âne, alors il bat un âne.*

La signification de l'énoncé analysé *Si Pedro possède un âne alors il le bat* est rendue transparente par la glose suivante:

(c') Si Pedro possède un âne il bat l'âne que Pedro possède

que nous pouvons représenter dans le calcul des prédicats par une quantification universelle restreinte:

(c'') $(\forall x)[\text{âne}(x) \ \& \ \text{possède}(\text{Pedro}, x) \rightarrow \text{bat}(\text{Pedro}, x)]$.

Cette représentation n'est cependant plus équivalente à la représentation avec quantification existentielle et changement de portée de la quantification:

(e) $\{(\exists x)[\text{âne}(x) \ \& \ \text{possède}(\text{Pedro}, x)]\} \rightarrow \text{bat}(\text{Pedro}, y)$.

En effet, l'occurrence de *y* est libre et n'est pas liée par la quantification existentielle, cette dernière représentation n'exprime donc plus la signification de (c).

Ces trois énoncés (a) *Pedro possède un âne*; (b) *Si Pedro possède un âne, il est heureux* et (c) *Si Pedro possède un âne, il le bat* font apparaître des représentations différentes pour le syntagme nominal *un âne* selon les contextes où ce syntagme apparaît; ce dernier est tantôt représenté par une quantification existentielle, tantôt représenté par une quantification universelle. Il s'ensuit que *ou bien* les langues ne sont pas compositionnelles au sens suivant: la syntaxe doit guider directement la construction sémantique, *ou bien* que la composition est étroitement

dépendante du formalisme utilisé. Dans ce dernier cas, le calcul des prédicats n'est pas un formalisme qui laisse transparaître la composition du français.

Plusieurs théories tentent d'apporter une solution à ce problème avec plus ou moins d'élégance.. Citons par exemple les tentatives dans le cadre de la DRT de Kamp ou dans l'interprétation dynamique qui prolonge les analyses linguistiques de la Grammaire Universelle de R. Montague. Une solution formelle élégante a été apportée par A. Ranta (1994) qui fait appel à la théorie intuitionniste des types de Martin-Löf. La logique combinatoire illative permet de donner une autre solution formelle en construisant, à l'aide de combineurs et du quantificateur universel illatif Π_2 , le prédicat complexe:

«dès-que-on₁-possède-âne-on₁-le-bat».

Pour construire ce prédicat complexe, nous allons analyser auparavant l'énoncé suivant:

(f) *Si Pedro possède un âne alors Juan le vend.*

Nous partons de la forme normale glosée et représentée dans le formalisme applicatif par respectivement:

«si x est un âne et si Pedro possède cet x alors Juan vend ce même x»

-> (& (âne x) (possède x Pedro)) (vend x Juan).

Par des β -expansions successives qui introduisent des combineurs et le quantificateur Π_2 , nous construisons le prédicat complexe «dès-que-on₁-possède-un-âne-(quelconque)-on₂-le-vend» de façon à avoir une représentation applicative «sans variables liées». Nous avons:

1. $\rightarrow (\& (\hat{\text{âne}} x) (\text{possède } x \text{ Pedro})) (\text{vend } x \text{ Juan})$ hyp.
2. $\rightarrow (\& (\hat{\text{âne}} x) (\text{Cpossède Pedro } x)) (\text{Cvend Juan } x)$ [i-C], 1.
3. $\rightarrow (\Phi \& \hat{\text{âne}} (\text{Cpossède Pedro } x)) (\text{Cvend Juan } x)$ [i- Φ], 2.
4. $\Phi \rightarrow (\Phi \& \hat{\text{âne}} (\text{Cpossède Pedro})) (\text{Cvend Juan } x)$ [i- Φ], 3.
5. $\Pi_2 (\Phi \& \hat{\text{âne}} (\text{Cpossède Pedro})) (\text{Cvend Juan})$ [i- Π_2], 4.
6. $\Pi_2 (\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède Pedro})) (\text{Cvend Juan})$ [i-B], 5.
7. $\mathbf{B}\Pi_2 (\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède})) \text{ Pedro } (\text{Cvend Juan})$ [i-B], 6.
8. $\mathbf{C}(\mathbf{B}\Pi_2 (\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède}))) (\text{Cvend Juan}) \text{ Pedro}$ [i-C], 7.
9. $\mathbf{B}(\mathbf{C}(\mathbf{B}\Pi_2 (\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède})))) (\text{Cvend Juan}) \text{ Pedro}$ [i-B], 8.
10. 10.1. $\mathbf{B}(\mathbf{C}(\mathbf{B}\Pi_2 (\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède}))))$ sous expression
- 10.2. $\mathbf{BBC}(\mathbf{B}\Pi_2 (\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède})))$ [i-B], 10.1
- 10.3. $\mathbf{B}(\mathbf{BBC})(\mathbf{B}\Pi_2)(\mathbf{B}(\Phi \& \hat{\text{âne}}) (\text{Cpossède}))$ [i-B], 10.2
- 10.4. $\mathbf{B}^2(\mathbf{B}(\mathbf{BBC})(\mathbf{B}\Pi_2)) \mathbf{B} (\Phi \& \hat{\text{âne}}) (\text{Cpossède})$ [i- $\mathbf{B}^2$], 10.3
11. $\mathbf{B}^2(\mathbf{B}(\mathbf{BBC})(\mathbf{B}\Pi_2)) \mathbf{B} (\Phi \& \hat{\text{âne}}) (\text{Cpossède})$
(Cvend) Juan Pedro rempl., 10.3, 9.
12. [tout-âne-que-l'on₁-possède, on₂-le-vend =_{def}
 $\mathbf{B}^2(\mathbf{B}(\mathbf{BBC})(\mathbf{B}\Pi_2)) \mathbf{B} (\Phi \& \hat{\text{âne}}) (\text{Cpossède}) (\text{Cvend})]$ def.
13. (tout-âne-que-l'on₁-possède, on₂-le-vend)
(Juan) (Pedro) rempl. 12., 11.

Le prédicat complexe «dès-que-on₁-possède-âne-on₂-le-vend» est donc défini par une expression applicative de la logique combinatoire illative:

dès-que-l'on₁-possède-un-âne-(quelconque), on₂-le-vend
 = tout-âne-que-l'on₁-possède, on₂-le-vend
 =_{def} $\mathbf{B}^2 (\mathbf{B}(\mathbf{BBC})(\mathbf{B}\Pi_2)) \mathbf{B} (\Phi \& \hat{\text{âne}}) (\text{Cpossède}) (\text{Cvend})$.

Par la même méthode, exprimons maintenant le prédicat complexe:

«dès-que-l'on₁-possède-âne, on₁-le-bat».

Nous avons:

1. $\rightarrow$ ($\&$ (âne x) (possède x Pedro)) (bat x Pedro) hyp.
2. $\Phi \rightarrow$ ($\Phi \&$ âne (Cpossède Pedro)) (Cbat Pedro) x [i- Φ], 1.
3. $\Phi \rightarrow$ (**B** ($\Phi \&$ âne) (Cpossède) Pedro) (Cbat Pedro) x [i-**B**], 2.
4. Π_2 (**B** ($\Phi \&$ âne) (Cpossède) Pedro) (Cbat Pedro) [i- Π_2], 3.
5. $\Phi \Pi_2$ (**B** ($\Phi \&$ âne) (Cpossède)) (Cbat) Pedro [i- Φ], 4.
6. [dès-que-on₁-possède-âne-on₁-le-bat =_{def} $\Phi \Pi_2$ (**B** ($\Phi \&$ âne) (Cpossède)) (Cbat)] def.
7. (dès-que-on₁-possède-âne-on₁-le-bat) (Pedro) rempl., 6., 5.

Le prédicat complexe «dès-que-on₁-possède-âne-on₁-le-bat» est donc défini par une expression applicative de la logique combinatoire illative:

tout-âne-que-l'on₁-possède, on₁-le-bat =_{def}
 $\Phi \Pi_2$ (**B**($\Phi \&$ âne) (Cpossède)) (Cbat).

Nous avons donc représenté l'énoncé *Si Pedro possède un âne alors il le bat* par l'application du prédicat complexe «dès-que-on₁-possède-âne-on₁-le-bat» au terme «Pedro» sans faire appel à des variables liées:

(dès-que-l'on₁-possède-âne, on₁-le-bat)(Pedro).

8. Conclusion

Nous avons montré comment *le type syntaxique d'un énoncé spécifie l'organisation syntaxique de l'expression* (énoncé). A l'énoncé correspond un programme applicatif dont *la spécification syntaxique est donnée par son type*; ce programme décrit comment les unités linguistiques agissent en tant qu'opérateurs sur les autres pour construire dynamiquement une signification. Cependant, plusieurs programmes applicatifs peuvent correspondre à un même énoncé, *la construction d'un programme particulier dépendant étroitement* (d'après l'isomorphisme de Curry-Howard) *de la réduction inférentielle du type syntaxique*, c'est-à-dire de la stratégie adoptée pour réduire le type syntaxique au type des phrases. Le programme normalisé, que l'on

déduit du programme applicatif particulier construit à partir de la réduction inférentielle du type syntaxique, décrit la représentation applicative sous-jacente à l'énoncé. Il existe une stratégie (donnée par les métarègles de la Grammaire Catégorielle Combinatoire Applicative) pour: 1° construire un programme applicatif correspondant à l'énoncé examiné à partir de type syntaxique de cet énoncé (utilisation de règles inférentielles sur les types, guidées par des métarègles); 2° en déduire le programme applicatif normal sous-jacent à l'énoncé, en normalisant, par des β -réductions, le programme déjà construit; ce programme normal représente l'interprétation sémantique fonctionnelle de l'énoncé analysé.

La discussion sur la composition (ou la non composition) des significations associées aux unités linguistiques dans les langues doit être reprise dans le nouveau cadre formel plus large de la logique combinatoire. Plus généralement, c'est la quantification des termes nominaux dans les langues naturelles qui doit être «revisitée», notamment en intégrant aux opérations de quantification, de détermination et de prédication les problèmes relatifs aux représentants objectaux, plus ou moins typiques et atypiques, d'un concept; la logique classique des prédicats est incapable de prendre en compte ces problèmes et donc de les formaliser alors que la psychologie et la linguistique cognitives en ont reconnu la pertinence dans chaque processus de catégorisation. Il apparaîtrait alors que le cadre post-frégéen, bien que très élégant et relativement adéquat à la description de la plupart des énoncés mathématiques, reste trop étroit et très insuffisant pour l'expression formelle des significations entre des énoncés qui sont paraphrastiquement très proches mais que la linguistique cherche à différencier explicitement pour mieux étudier les opérations langagières qui sont constitutives des énoncés. Les analyses, effectuées par Aristote, Ockham et les Messieurs de Port-Royal sont sans doute plus respectueuses des structures linguistiques que les analyses post-frégéennes de Montague et de Dowty, comme en témoigne le débat ouvert par Fred Sommers. Elles devraient être réexaminées, sans doute avec fruit, à l'aide du formalisme applicatif «sans variables» de la logique combi-

natoire typée de Curry (inconnu, bien entendu, des auteurs anciens).

Institut des Sciences Humaines Appliquées
Université de Paris-Sorbonne
 96, boulevard Raspail
 F 75006 PARIS

Références bibliographiques

- AJDUKIEWICZ K. (1935). Die syntaktische Konnexität. *Studia Philosophica* 1, 1-27.
- BACH E. (1988). Categorical grammars as theories of language. In: Oehrle (1988), 17-34.
- BAR HILLEL Y. (1953). A quasi-arithmetical notation for syntactic description. *Language* 29, 47-58.
- BISKRI I. (1996). Logique combinatoire et linguistique: grammaire catégorielle combinatoire applicative. *Mathématiques, Informatique et Sciences Humaines* 32, 39-68.
- BISKRI I., DESCLÈS J.-P. (1996). Du phénotype au génotype: la grammaire catégorielle combinatoire applicative. *Actes de TALN*. Marseille, 87-96.
- CURRY H.-B. (1930). Grundlagen der kombinatorischen Logik (Inauguraldissertation). *Amer. J. Math.*, 509-536, 789-834.
- CURRY H.-B., FEYS R. (1958). *Combinatory Logic*. Amsterdam: North Holland, vol. 1.
- CURRY H.-B., HINDLEY J.R. & SELDIN J.-P. (1992). *Combinatory Logic*. Amsterdam: North Holland, vol. 2.
- DESCLÈS J.-P. (1980). *Opérateurs / opérations: méthodes intrinsèques en informatique fondamentale. Applications aux bases de données relationnelles et à la linguistique*. Paris: Thèse de doctorat de l'Université René Descartes.
- DESCLÈS J.-P. (1990). *Langages applicatifs, langues naturelles et cognition*. Paris: Hermès.
- DESCLÈS J.-P. (1991). La double négation dans l'Unum Argumentum analysé à l'aide de la logique combinatoire. In:

- La négation. Le rôle de la négation dans l'argumentation et la raisonnement.* Université de Neuchâtel: Travaux du Centre de Recherches Sémiologiques n° 59, 33-74
- FITCH F. (1974). *Elements of Combinatory Logic*. Yale: University Press.
- FREY L. (1967). Langues logiques et processus intellectuels. In: *Modèles et formalisation du comportement*. Paris: Editions du CNRS, 327-345.
- GEACH P.T. (1962). *Reference and Generality*. Ithaca: Cornell Univ. Press.
- GEACH P.T. (1971). A program for syntac. *Synthese* 22, 3-17.
- GIRARD J.Y. (1995). Linear logic: its syntax and semantics. In: J.Y. Girard et al. (eds), *Advances in Linear Logic*. Cambridge: Cambridge University Press, 1-42.
- HARRIS Z.S. (1982). *A Grammar of English on Mathematical Principles*. New York: Wiley.
- HOWARD W.A. (1980). The formulae-as-types notion of construction. In: J.-P. Seldin & J.R. Hedin (eds), *To H.B. Curry: Essays on Combinatory Logic, Lambda Calculus and Formalism*. London: Academic Press, 479-490.
- HUSSERL E. (1913). *Logische Untersuchungen*. Halle: M. Niemeyer.
- LAMBEK J. (1958). The mathematics of sentence structure. *American Mathematical Monthly* 65, 154-165.
- LAMBEK J. (1961). On the calculus syntactic types. *Proceedings of Symposia in Applied Mathematics* vol. XII. Providence: America Mathematical Society, 166-178.
- LAMBEK J. (1972). Deductive systems and categories. III Cartesian closed categories, intuitionist propositional calculus and combinatory logic. *Lectures Notes in Mathematics* 274, 57-82.
- LAMBEK J., SCOTT P.J. (1986). *An Introduction to Higher Order Categorical Logic*. Cambridge: Cambridge University Press.
- LESNIEWSKI S. (1989). *Sur les fondements de la mathématique. Fragments (Discussions préalables, méréologie, ontologie)*. Trad. de G. Kalinowski. Paris: Hermès.

- OEHRLE R.T. et al. (eds) (1988). *Categorial Grammars and Natural Languages Structures*. Dordrecht: Reidel.
- MARTIN-LÖF P. (1984). *Intuitionistic Type Theory*. Naples: Bibliopolis.
- MOORTGAT M. (1988). *Categorial Investigation. Logical and Linguistic Aspects of the Lambek Calculus*. Dordrecht, Providence RI: Foris Pub.
- PARTEE B.H., MEULEN A. & WALL R.E. (1993). *Mathematical Methods in Linguistics*. Dordrecht: Kluwer Academic Pub.
- SHAUMYAN S.K. (1977). *Appplcational Grammar as a Semantic Theory of Natural Language*. Edinburgh: University Press.
- SHAUMYAN S.K. (1987). *A Semiotic Theory of Natural Language*. Bloomington: Indiana University Press.
- STEEDMAN M. (1988). Combinators and grammars. In: Oherle (1988), 207-263.
- STEEDMAN M. (1989). *Work in Progress: Combinators and Grammars in Natural Language Understanding*. Tuscon University: Summer Institute of Linguistic.
- STEEDMAN M. (1989). Constituency and coordination in a combinatory grammar. In: M. Baltin & J. Kroch, *Alternative Conceptions of Phrase Structure*. Chicago: University Press, 201-231.
- VAN BENTHEM J. (1991). *Language in Action: Categories, Lambdas and Dynamic Logic*. Amsterdam: North Holland.

LA LOGIQUE DÉVELOPPEMENTALE

Denis MIÉVILLE

Je vous écris du bout du monde. Vous n'imaginez pas tout ce qu'il y a dans le ciel, il faut l'avoir vu pour le croire. Ainsi, tenez, les ... mais je ne vais pas vous dire leur nom tout de suite.

Henri Michaux

Préambule

Notre intérêt va vers l'étude d'une logique développementale, une logique avec laquelle il est possible de marquer les étapes du développement cognitif. Il s'agit donc d'une logique qui, bien que ne contenant pas de foncteurs de vérité liés à la temporalité, est à même de révéler la chronologie de son développement. Cette logique peut être caractérisée globalement de la manière suivante:

- 1) elle est une logique vérifonctionnelle et extensionnelle;
- 2) elle se distingue des logiques classiques par l'ensemble particulièrement important des foncteurs logiques qu'elle peut contenir;
- 3) elle est conçue de manière à être progressivement étendue pour donner accès à de nouveaux foncteurs.

Cette caractérisation doit être affinée et justifiée. Dans un premier temps, nous nous y emploierons en proposant un ensemble de considérations qui sont de nature à éclairer cette perspective développementale. Ensuite nous exposerons les fondements d'une telle logique. Enfin, nous en évoquerons quelques conséquences.

Quelques considérations à propos de la logique développementale

Première considération:

La logique classique du premier ordre n'est pas le tout de la logique extensionnelle. Il s'agit, certes d'une banalité, mais une banalité qu'il est nécessaire parfois de rappeler. En effet, la logique classique du premier ordre est avant toute chose la logique mathématique. Elle a été élaborée en fonction d'un dessein bien particulier, très explicitement lié à la construction des fondements de l'arithmétique. Elle a donc été développée de manière à disposer des objets conceptuels indispensables à sa finalité. Une conséquence de cette démarche est que la logique qui en résulte contient un nombre relativement modeste de catégories syntaxico-sémantiques et, par voie de conséquence, elle fait usage d'un nombre restreint d'opérations logiques.

Deuxième considération:

Pour aborder logiquement les démarches raisonnées des sciences déductives, et cela dans leur dimension extensionnelle et vérifonctionnelle, les catégories syntaxico-sémantiques des noms [N] et des propositions [S] s'imposent comme les catégories fondamentales. La catégorie des propositions est indispensable pour aborder ce qui relève des valeurs de vérité, et la catégorie des noms est indissociablement liée aux objets qui peuplent les univers sur lesquels les raisonnements logiques sont appliqués.

Troisième considération:

Toute catégorie syntaxico-sémantique différente des catégories S et N nécessaires à l'édifice logique est définissable en termes des deux catégories de base. Ainsi, dans la

logique classique des propositions, en plus des deux catégories fondamentales S et N, deux autres catégories sont nécessaires. Il y a d'une part la catégorie des foncteurs formateurs de propositions à un opérande propositionnel [S/S], à laquelle appartient l'opération de négation. Il y a d'autre part la catégorie des foncteurs formateurs de propositions à deux opérandes propositionnels [S/SS], dont la conjonction et la disjonction font notamment partie.

Suite aux considérations qui précèdent, il est dès lors légitime de se poser la question du nombre des catégories syntaxico-sémantiques que la logique doit contenir, et donc du nombre des foncteurs qui leur sont associés, pour remplir son rôle par rapport à l'ensemble du savoir déductif. Cette question n'est pas nouvelle et en 1931, Tarski se la posait déjà.

Le langage d'un système complet de logique devrait contenir en tant que tel – en acte ou en puissance – toutes les catégories sémantiques possibles apparaissant dans le langage des sciences déductives. Cette circonstance confère justement à ce langage un caractère universel dans un certain sens et est l'un des facteurs auxquels la logique doit son importance fondamentale pour l'ensemble du savoir déductif. La variété des catégories sémantiques dans tels ou tels systèmes fragmentaires de logique ou dans d'autres sciences déductives peut être considérablement limitée – tant au point de vue de leur nombre que de leur ordre (Tarski 1974, vol. 1: 219, 1ère éd. 1931).

Comme l'évoque Tarski, la variété des catégories sémantiques peut être limitée dans certains systèmes – fragmentaires, ajoute-t-il – de logique. Le plus bel exemple de cette limitation est encore la logique classique du premier ordre qui ne connaît que les variétés suivantes: N, S, S/S, S/SS, et la famille des quantificateurs qui opèrent sur des fonctions propositionnelles à arguments nominaux pour former des expressions de la catégorie des propositions, S/fonction propositionnelle.

Il est donc admis qu'il existe plusieurs langages liés à divers systèmes de logique; Tarski offre même une classification de ceux-ci:

...nous pouvons distinguer quatre *espèces de langages*: (1) les langages dont toutes les variables appartiennent à une même catégorie sémantique; (2) les langages où le nombre de catégories contenant des variables est plus grand que un mais fini; (3) les langages où les variables appartiennent à un nombre infini de catégories différentes, l'ordre de celles-ci ne dépassant cependant pas un nombre naturel n déterminé à l'avance; enfin (4) les langages contenant des variables d'un ordre quelconque si élevé qu'il soit (Tarski 1974, vol. 1: 219-220).

Considérant que le langage d'un système complet de la logique doit contenir des variables d'un ordre quelconque, si élevé qu'il soit, il est indispensable de s'interroger sur la forme que prendrait la manifestation d'un tel projet. Avant d'avancer quelques éléments de réponse, nous poserons ce qui suit:

Quatrième considération:

Le nombre des catégories syntaxico-sémantiques d'un système complet de logique est infini. En effet, la pensée est capable de concevoir une infinité dénombrable de catégories à partir de celles de base que sont les catégories des noms et des propositions. La définition inductive suivante est de nature à révéler progressivement toutes ces catégories:

- i. N et S sont des catégories.
- ii. Si $C, C_1, C_2, \dots, C_n$ sont des catégories, alors $C/C_1C_2\dots C_n$ est une catégorie. Il s'agit de la catégorie des foncteurs formateurs d'expressions de la catégorie C à n arguments dont le premier est de la catégorie $C_1, \dots$, le $n^{\text{ième}}$ de la catégorie C_n .
- iii. Rien n'est une catégorie sinon par ce qui précède.

Cinquième considération:

S'il est concevable d'envisager, potentiellement parlant, toutes les possibilités issues des deux catégories de base, le nombre effectif des opérations logiques connues, et

donc des catégories qui leur correspondent dans le langage des sciences déductives est à un moment donné fini.

Dans l'impossibilité de disposer d'une théorie logique contenant en acte la totalité des opérations logiques, il faut concevoir un système logique qui soit en mesure d'être progressivement enrichi de nouvelles catégories et de nouveaux foncteurs. Les considérations précédentes nous engagent alors à aborder les questions suivantes:

Si l'ensemble des opérations logiques liées au langage des sciences déductives n'est pas connu de manière exhaustive, une fois pour toutes, mais seulement de manière progressive, en fonction des objets problématiques que la logique parvient à résoudre:

- 1 – De quelle manière inscrire progressivement la signification de nouveaux foncteurs?
- 2 – A partir de quelle base en termes de significations primitives fonder de nouveaux foncteurs?
- 3 – De quel type de système faut-il disposer pour qu'il puisse effectivement être façonné de manière progressive?
- 4 – Quelles sont les conséquences liées à une telle manière de faire?
- 5 – Une sémantique purement distributive est-elle encore adéquate pour représenter des opérations logiques d'une complexité croissante?

Nous répondrons tout d'abord à la troisième question. Disposer d'un système logique capable d'intégrer progressivement des idées nouvelles nous engage à réfléchir d'une part sur le choix de la base axiomatique sur laquelle fonder l'édifice logique en devenir et d'autre part, sur la nature d'une procédure rendant effective une telle progression d'idées logiques nouvelles. Il est donc nécessaire, dans un premier temps, de justifier un choix possible de significations primitives initiales à partir desquelles on peut accéder à de nouvelles significations. Construire des significations nouvelles à partir d'autres significations préalablement inscrites présuppose une activité définitoire. Il est dès lors intéressant de s'intéresser aux conditions qui règlent toute

bonne définition explicite. Celles-ci sont connues depuis fort longtemps et nous ne nous y attarderons guère (à ce sujet, cf. Miéville 1994). Pour notre propos, nous ne retiendrons pour l'instant qu'une condition: une bonne définition établit une relation d'équivalence entre une expression propositionnelle *A* qui contient le terme constant nouveau ou le foncteur nouveau à définir, et une autre expression propositionnelle *B*, qui présente dans son organisation ce qui sert à définir. Or, deux entités propositionnelles ont entre elles une relation d'équivalence si et seulement si, en leur appliquant l'opération de biconditionnelle, on obtient une expression logiquement valide:

$$A \longleftrightarrow B \text{ si et seulement si } \vdash A \equiv B.$$

Cette information est importante à deux niveaux: d'une part elle fournit un candidat particulièrement intéressant, la biconditionnelle, pour remplir le rôle d'une des significations primitives à inscrire dans la base axiomatique du système en devenir et d'autre part, elle offre une indication fondamentale quant à la nature des définitions dans le cadre d'un système qui est enrichi progressivement d'idées logiques nouvelles. Traditionnellement, toute définition explicite n'est qu'une commodité linguistique. En effet, bien que subjectivement parlant on lui accorde une finalité cognitive, elle n'apparaît formellement que comme une abréviation. Et c'est bien sa fonction effective dans les systèmes classiques, qui introduisent toute définition à l'aide d'un signe métalinguistique,

$$A = \text{df. } B.$$

N'accorder qu'un statut abrégatif à la définition n'est guère satisfaisant intellectuellement parlant et surtout inutile par rapport au projet de l'élaboration d'un système qui devrait toujours pouvoir s'enrichir d'idées logiques nouvelles. Russell, déjà, à l'aube de notre siècle, en était conscient.

it is a curious paradox, puzzling to the symbolic mind, that definitions, theoretically, are nothing but statements of symbolic abbreviations, irrelevant to the reasoning and inserted only for practical convenience, while yet in the development of a subject, they always require a very

large amount of thought, and often embody some of the greatest achievements of analysis (Russell 1956: 63, 1ère éd. 1903).

Afin de briser le caractère pratique lié à la dimension abrégative de la définition et lui substituer le rôle quasi créatif qui doit lui revenir, il est indispensable de penser l'activité définitoire comme une règle d'inférence qui introduit des expressions-définitions comme des théorèmes du système. A ce prix-là, il y aurait rupture avec le rôle purement et effectivement abrégatif des définitions. Une telle règle devra être conçue de manière à déterminer les critères d'admission d'expressions-définitions nouvelles en fonction des significations préalablement inscrites dans le système. Ces expressions-définitions apparaîtraient dans le système sous la forme d'expressions biconditionnelles ayant statut de théorème:

|— $A \equiv B$.

Sur la base de ces remarques liées à la réalisation d'une logique développementale, il convient donc d'une part de déterminer les axiomes nécessaires afin de fonder au moins la signification de l'opération de biconditionnelle et d'autre part, d'explicitier les conditions inférentielles permettant d'introduire dans le système en devenir de nouvelles expressions biconditionnelles porteuses de concepts logiques nouveaux, expressions ayant le statut de théorème. Mais avant d'aborder cet aspect des choses, il est nécessaire de franchir un obstacle important. En effet, cette manière d'envisager le développement d'une théorie logique ne présuppose en aucune manière un ordre d'accès aux catégories et aux opérations logiques, ni une limite exhaustive visant à les circonscrire. A cette indétermination sémantique liée à l'ensemble des opérations logiques effectivement connues correspond une indétermination syntaxique. Ne connaissant pas à l'avance ce que l'on va définir, il n'est tout simplement pas possible de proposer une liste de symboles préalablement fixés en fonction d'un projet sémantique catégoriellement et fonctionnellement fermé. Un système développemental ne saurait donc posséder une des caractéristiques essentielles de tout système formel au sens classique du terme: celui d'être «*donné tout entier*»,

en une seule fois» (Chazal 1995: 73). Au temps initial d'une logique développementale, seules les significations initiales indispensables sont présentes; toute autre signification est l'objet d'une définition interne au système, définition qui ne soutient aucun lien avec un objet syntaxique préalablement inscrit. Hormis ce qui est proposé dans la base axiomatique, nul autre symbole n'est explicité *a priori*. La sélection de nouveaux symboles supportant de nouvelles significations, et leur reconnaissance catégorielle et fonctionnelle, doivent donc passer par une procédure d'identification qui s'appuie sur d'autres critères que ceux d'une forme qui renvoie univoquement à une signification. La solution passe par une détermination contextuelle des symboles.

Une détermination contextuelle n'est pas une nouveauté. Qui-conque s'est interrogé sur la catégorie grammaticale de la marque conjonctive en français, par exemple, a forcément été sensible au fait que, prise isolément, cette conjonction n'appartient pas à une catégorie univoquement déterminée. Pour attribuer à telle ou telle conjonction «et» sa catégorie, il est nécessaire d'étudier la nature des objets linguistiques sur lesquels elle porte. Ainsi, dans la proposition «Pierre lit le journal et Jean rit» la conjonction est de la catégorie des foncteurs formateurs de propositions à deux arguments propositionnels parce que cette conjonction-là porte sur deux entités propositionnelles: «Pierre lit le journal» d'une part, et «Jean rit» de l'autre. Alors que dans l'exemple suivant, «Le drapeau suisse est rouge et blanc», la conjonction est de la catégorie des foncteurs formateurs de noms à deux arguments nominaux parce que cette conjonction-ci porte sur deux entités nominales: «rouge» d'une part, et «blanc» de l'autre. Ainsi, pour obvier à l'indétermination sémantique, et donc syntaxique, propre aux logiques développementales, la manière choisie est justement cette détermination contextuelle dont on esquissera ici le principe. Nous poserons dorénavant que tout foncteur (constant ou variable) d'une quelconque catégorie précédera une paire de parenthèses symétriques – pour l'instant, peu importe sa forme –, parenthèses qui enserrent un nombre fini d'arguments. Ainsi, dans les inscriptions suivantes:

*(ay%) et +[xcvj]

* et + sont deux foncteurs. Le premier agit sur trois arguments, le deuxième sur quatre. Ainsi, et nous le préciserons un peu plus loin, la catégorie d'un foncteur est déterminable en fonction de la forme des parenthèses et du nombre des arguments qui lui sont associés.

Mais comment réaliser si soit «*», soit «+» a le statut de variable ou de constante? La question se pose tout naturellement dans la mesure où nous ne disposons pas ici de listes préalables de symboles, indiquant s'ils fonctionneront comme variables ou comme constantes. La réponse passe encore une fois par un mode de détermination contextuelle. Précisons les choses en imposant ce qui suit:

- 1) Toute expression bien formée du système en devenir doit être une expression quantifiée.
- 2) Une expression quantifiée est une expression complexe composée de deux sous-parties concaténées; la première est une expression parenthésée, non vide dont les parenthèses sont équiiformes aux parenthèses suivantes: [et], et la deuxième, non vide également, dont les parenthèses sont équiiformes aux parenthèses suivantes: [et]. On appelle généralement la première partie un quantificateur, et la deuxième, un sous-quantificateur.

Dans ces conditions, une expression bien formée a la forme globale suivante:

[... $\epsilon\beta\alpha$...][... $\alpha\gamma\mu\epsilon^4$...], peu importe pour le moment la signification et la forme des éléments que cette expression contient.

Il est dès lors possible de définir, dans une expression, quels symboles ont statut de variable ou statut de constante:

- 3) Un symbole dans un sous-quantificateur possède le statut de variable si et seulement si
 - il n'est pas une parenthèse symétrique,
 - il est différent des symboles choisis pour définir la partie quantification et la partie sous-quantification d'une expression bien formée,

- il possède un symbole équiforme dans le quantificateur.

Ainsi, sans disposer d'une liste de symboles préalablement déterminés sémantiquement parlant, il est possible de déterminer le statut de chaque symbole de manière contextuelle.

Nous pouvons donc ainsi attribuer aux symboles e et α dans [... $\alpha v \mu e 4$...] le statut de variable parce qu'ils possèdent l'un et l'autre un signe équiforme dans le quantificateur correspondant au sous-quantificateur dans lequel ils sont inscrits: [... $e \beta \alpha$...].

Les symboles v , μ et 4 dans [... $\alpha v \mu e 4$...] ont le statut de constante parce qu'ils ne sont pas des parenthèses symétriques, ils ne sont pas des symboles équiformes aux parenthèses utilisées pour les quantificateurs et les sous-quantificateurs, et parce qu'ils n'ont pas de signes équiformes dans le quantificateur [... $e \beta \alpha$...] qui correspondent au sous-quantificateur dans lequel ils sont inscrits.

Nous sommes en présence d'un certain nombre d'éléments déterminants pour nous diriger vers le projet qui nous intéresse:

- nous savons distinguer ce qui peut avoir statut de constante ou de variable,
- nous savons reconnaître ce qui peut jouer le rôle de foncteur et la forme globale d'une expression de type fonctionnel,
- par ailleurs nous savons que le mode de reconnaissance de la catégorie d'un symbole de foncteur passera par la forme des parenthèses symétriques de l'expression parenthésée qui le suit, ainsi que du nombre des arguments que cette expression contient,
- enfin, la procédure définitoire permettant de donner accès à de nouvelles significations, voire catégories, est directement conçue sur l'existence de l'opération logique de biconditionnelle.

Il est dès lors possible de faire un pas de plus en proposant les significations primitives et les contextes primitifs qui seront à la base de notre système en devenir. La signification primitive sera celle de la biconditionnelle, inscrite dans une base axiomatique et qui fixe, par décision initiale, le premier contexte. Cette décision initiale est d'autant plus importante qu'elle donne le

code de détermination de la catégorie syntaxique de la première signification du système, celle qu'il possède à l'état initial. A ce temps-là du système, rien d'autre, syntaxiquement parlant, n'existe.

Par décision, nous imposerons dorénavant que tout symbole, différent d'un symbole de parenthèse ou différent d'un symbole de délimitateur de quantification ou de sous-quantification, sera de la catégorie des foncteurs formateurs de propositions à deux arguments propositionnels [S/SS] si et seulement si il précède une paire de parenthèses symétriques équiiformes à (, respectivement), et que cette expression parenthésée contient exactement deux arguments.

Ainsi, dans l'expression suivante, $[... \beta\alpha ...] [\dots \alpha(v w) \dots]$, je reconnais que le α est un foncteur variable de la catégorie des foncteurs formateurs de propositions à deux arguments propositionnels, non pas parce qu'il est inscrit avec la forme α , mais parce qu'il précède une expression parenthésée dont les parenthèses symétriques sont équiiformes à (--) et que cette expression possède deux arguments. Il s'ensuit, et c'est un corollaire, que toute entité inscrite dans une expression parenthésée, dont les parenthèses sont équiiformes à (--), et qui contient deux unités, est une expression de la catégorie des propositions. Sur la base de cette réalité ainsi décrite, et en nous basant sur le fait que nous voulons disposer de la signification de la biconditionnelle comme première signification primitive, il est indispensable d'inscrire une base axiomatique pour fixer les choses. Nous emprunterons celle que propose S. Lesniewski, qui fut le premier, à l'aube de ce siècle, à formuler les bases d'une logique développementale. Les axiomes sont exprimés dans le cadre d'un système d'écriture préfixée, parenthésée, et à détermination de type contextuel:

$$A1: \quad [pqr] \vdash (\equiv (\equiv (pr) \equiv (qp)) \equiv (rq)) \vdash$$

$$A2: \quad [pqr] \vdash (\equiv (p \equiv (qr)) \equiv (\equiv (pq) r)) \vdash$$

$$A3: \quad [pg] \vdash (([f] \vdash (g(pp) \equiv ([r] \vdash (f(rr)g(pp)))) \vdash$$

$$[r] \vdash (f(rr)g(\equiv (p [q] \vdash [q] p))) \vdash [q] \vdash g(pq) \vdash$$

En regard de cette écriture rigoureuse qui respecte les conditions strictes liées au projet d'une logique développementale, présentons une transcription bâtarde, mais plus conforme à notre manière habituelle de procéder:

$$A1: (\forall pqr)((p \equiv r) \equiv (q \equiv p)) \equiv (r \equiv q))$$

$$A2: (\forall pqr)((p \equiv (q \equiv r)) \equiv ((p \equiv q) \equiv r))$$

$$A3: (\forall pg) [(\forall f)(g(pp) \equiv ((\forall r)(f(rr) \equiv g(pp)) \equiv (\forall r)(f(rr) \equiv g((p \equiv (\forall q)(q)p)))) \equiv (\forall q)(g(pq)))]$$

Au temps initial de l'existence de notre système, il n'existe que trois expressions logiques, les trois axiomes, et ceux-ci en constituent donc les trois premières thèses. A l'état initial, le système dispose donc de deux catégories syntaxiques: celle des propositions, S, présente à travers l'inscription de symboles variables et celle des foncteurs formateurs de propositions à deux arguments propositionnels, S/SS, catégorie inscrite à travers le symbole du foncteur constant $\equiv$, et de variables. Le système proposé est donc un système d'ordre supérieur. Rappelons que la détermination de la catégorie du symbole $\equiv$ est de la catégorie S/SS non parce qu'il a cette forme précise, mais parce qu'il apparaît dans un contexte particulier dans lequel il est un signe qui précède l'expression parenthésée (pq), et pour aucune autre raison formelle. Les axiomes 1 et 2 expriment certaines des propriétés essentielles qui caractérisent la biconditionnelle; le troisième contient d'une part le principe de bivalence pour la biconditionnelle et les foncteurs variables de la catégorie S/SS et d'autre part, quelques formes du principe d'extensionnalité pour les propositions.

L'avenir de ce système est d'une certaine manière entièrement indéterminé dans la mesure où son expansion progressive peut prendre de multiples chemins. Cette liberté, liée à l'ensemble infini des possibilités de développement, entraîne une indétermination syntaxique. Nous montrerons de quelle manière, grâce à la règle d'inférence définitoire, il est possible de développer dans le temps et l'espace le cheminement logique que l'on veut concevoir, en marquant le temps de la progression cognitive qui lui est liée.

Où il est question de la définition

Une règle inférentielle de définition doit, pour remplir le contrat d'une logique développementale, permettre d'inscrire des théorèmes dans le système. Il s'agit de théorèmes particuliers dans la mesure où une expression définitoire doit être de préférence une expression biconditionnelle; nous avons précédemment expliqué les raisons de ce choix. Ainsi, la règle de définition permettra d'introduire des expressions complexes de la forme très schématique suivante:

$$[v_1 v_2 \dots v_n] \vdash \equiv (f \langle v_1 v_2 \dots v_n \rangle E_{v_1 v_2 \dots v_n}) \quad]$$

dont l'expression bâtarde correspondante est la suivante:

$$(\forall pqr)(f \langle v_1 v_2 \dots v_n \rangle \equiv E_{v_1 v_2 \dots v_n})$$

dans laquelle le symbole **f** représente un symbole de foncteur constant nouveau par rapport à la catégorie à laquelle il appartient, foncteur portant sur **n** variables; l'expression $f \langle v_1 v_2 \dots v_n \rangle$ est le *definiendum* de la définition, et l'expression $E_{v_1 v_2 \dots v_n}$, le *definiens*. Cette dernière expression doit être formée avec ce que le système a préalablement construit et contenir les variables $v_1, v_2, \dots, v_n$. On peut remarquer que, conformément à ce qui avait été annoncé, le *definiendum* et le *definiens* sont bien inscrits dans une expression biconditionnelle:

$$\equiv (f \langle v_1 v_2 \dots v_n \rangle E_{v_1 v_2 \dots v_n})$$

Dans l'explicitation de la mise en oeuvre de la règle inférentielle définitoire il sera nécessaire de spécifier que:

- i) Si le nouveau foncteur constant **f**, destiné à être introduit, est d'une catégorie syntaxico-sémantique précédemment introduite, il est nécessaire de respecter formellement le contexte qui caractérise cette catégorie préalablement introduite. Ainsi, si l'on désire introduire une nouvelle constante de la

catégorie S/SS, il est indispensable de respecter le contexte préalablement choisi, à savoir (--).

- ii) Si le nouveau foncteur constant *f*, destiné à être introduit, est d'une catégorie que le système ne connaît pas encore, il est nécessaire de choisir un contexte qui lui attribue une nouvelle identité catégorielle et cela, en prenant la précaution de choisir un parenthésage qui évite toute confusion. A titre d'exemple, prenons la situation où il est nécessaire d'introduire un foncteur constant binaire, formateur de propositions à deux arguments de la catégorie des foncteurs formateurs de propositions à deux arguments propositionnels, S/(S/SS)(S/SS). Dans ce cas, il est exclu de faire usage des parenthèses équiiformes à (et). En effet, ce nouveau foncteur étant binaire lui aussi, le choix des parenthèses (et) rendrait totalement ambiguë la catégorie du nouveau foncteur puisque cette paire de parenthèses a été préalablement utilisée pour identifier dans un autre contexte binaire, une autre catégorie. A part ce cas, toute autre paire de parenthèses différentes de (et) et d'autres paires de parenthèses liées à des contextes catégoriels binaires préalablement inscrits peut être choisie.

Par ailleurs, l'expression de la thèse-définition doit respecter entre le *definiens* et le *definiendum* les conditions fondamentales liées à toute définition explicite. Nous les rappelons ici:

1. Etablir une relation d'équivalence entre un *definiendum* (ce qui porte le défini), et un *definiens* (ce qui permet de définir).
2. Construire le *definiens* sur ce qui a été préalablement posé, voire défini.
3. Inscrire un terme constant nouveau unique dans le *definiendum*. Celui-ci ne doit être équiiforme à aucun autre terme préalablement posé ou défini de la même catégorie syntaxico-sémantique.
4. Prendre garde à ce que toute variable du *definiens* possède une occurrence dans le *definiendum*. Ne pas respecter cette condition conduit à la possibilité de constructions contradictoires.

5. Prendre garde à ce que toute variable du *definiendum* possède une occurrence dans le *definiens*. Il s'agit d'une condition de pure esthétique!
6. Veiller à ce qu'aucun signe du *definiendum* ne soit répété. Une infraction à cette condition peut poser des problèmes dans la mise en oeuvre de la règle de substitution.

La procédure définitoire proposée par S. Lesniewski, se fondant sur la constante logique de biconditionnelle, respecte ces conditions.

Nous avons beaucoup insisté, dans l'esprit tout au moins, sur la règle d'inférence définitoire proposée par Lesniewski. Celle-ci permet, à partir d'une base syntaxico-sémantique modeste, de développer progressivement l'ensemble des catégories logiques – et donc des constantes qui lui correspondent – concevables à partir de la catégorie initiale et primitive des propositions, S. Cependant un logicien ne saurait se contenter d'enrichir progressivement la bibliothèque potentiellement infinie des opérations logiques propositionnelles. Il doit encore être capable de les mettre en oeuvre, afin d'atteindre l'ensemble des théorèmes logiques de son système, quel qu'en soit l'état de développement. Une logique développementale doit donc disposer de règles d'inférence supplémentaires: les règles de détachement de la biconditionnelle, de substitution et d'extensionnalité, pour ne citer que les plus importantes. La logique esquissée ici les possède bien entendu, mais notre propos est d'insister avant tout sur la définition temporelle liée au développement syntaxico-sémantique d'une telle logique. Nous ne ferons ainsi qu'évoquer l'existence de ces autres règles d'inférence. En exposant les bases de cette logique des propositions élargie, la protothétique, S. Lesniewski offre une méthodologie permettant un accès à une logique issue de la catégorie des propositions la plus généreuse qui soit.

Pour caractériser l'esprit d'une logique développementale dans laquelle la progression temporelle est dominante, nous avons antérieurement consacré notre propos à la conception purement propositionnelle de la logique. Elle n'est enfin qu'un aspect, certes fondamental, mais incomplet de l'édifice logique.

Si nous avons procédé de cette manière c'est pour insister sur ce qui nous importe ici: la dynamique liée au temps. Mais ce qui relève des activités liées aux valeurs de vérité est aussi déterminant; il s'agit de ce qui participe au discours logique portant sur les problèmes de la référence et des structures relationnelles. Dans cette perspective, la catégorie syntaxico-sémantique des noms est fondamentale et, ajoutée à celle des propositions, elle permet, par combinatoire, d'épuiser l'ensemble des catégories syntaxico-sémantiques qu'une logique complète et idéale devrait posséder. Ici encore, nous empruntons à Lesniewski un système élargi des prédicats d'ordre supérieur, habité du dynamisme développemental qui marque la temporalité. Ce système, que nous allons maintenant aborder, porte le nom d'ontologie.

Un aperçu de l'ontologie de Lesniewski

Proposons une approche présémantique pour aborder l'ontologie de Lesniewski qui constitue un calcul d'ordre supérieur des termes nominaux. Une telle approche est constituée par une appréhension naïve et naturelle que nous avons du monde et de notre manière d'en parler. Nous croyons à l'existence matérielle ou non matérielle de choses; Socrate, cette page, la Lune appartiennent à notre réalité. Nous savons raisonner avec de tels objets et, pour le faire, nous leur associons des noms. Ces noms sont considérés comme des *noms individuels*, des noms qui dénotent des choses considérées comme des entités. Mais on sait qu'il existe des noms d'une autre nature; il y a des *noms généraux*, Nicolas Bourbaki – ce mathématicien polycéphale – est l'un d'entre eux. Il existe également un autre type de noms; les logiciens le savent bien, eux qui n'ont eu de cesse de raisonner sur le thème de Pégase, de l'actuel roi de France, et autre cercle-carré. Il s'agit des *noms vides*, c'est-à-dire, de noms qui ne dénotent aucun objet. Enfin, il existe des objets sans nom et, pour cette raison, nous ne saurions en parler. Lesniewski va admettre cette variété sous la catégorie des noms.

Lorsque nous parlons, lorsque nous raisonnons, nous ne cessons d'utiliser – dans les langues indo-européennes en tous les

cas – la copule *est*. Cette copule joue un rôle logique considérable, même si elle n'apparaît pas explicitement dans le calcul logique des prédicats, où elle se voit amalgamée d'une certaine manière aux propriétés et aux prédicats. Lesniewski s'y intéresse directement, et ceci pour des raisons d'évidence toute pratique. En effet, lorsqu'il développe les linéaments de sa théorie des ensembles collectifs, il les expose dans une langue naturelle. Axiomes et théorèmes apparaissent comme des propositions particulières dans lesquelles la copule *jest* – l'équivalent polonais du *est* – abonde. Cette copule est associée aux noms qu'il utilise pour organiser les objets de sa théorie. Il s'intéresse donc à la signification de cette copule lorsqu'elle articule des noms. Cet intérêt aboutit à l'explicitation d'une théorie des termes capable de représenter un calcul des noms.

Le génie de Lesniewski est associé à une exigence de rigueur peu commune, ainsi qu'à la conscience qu'une langue formelle doit hériter, dans la mesure de ce qui est possible, de toutes les richesses que nous offre la pensée en discours, et notamment de son pouvoir de créativité. Cette attitude explique en partie son refus de travailler avec les systèmes de la tradition russellienne. Il va donc offrir un système logique qui fixe de manière axiomatique une signification de la copule. Celui-ci, à l'image de tous les systèmes de Lesniewski, est d'une nature conceptuelle et développementale. Il peut donc être également développé de manière progressive, sur la base de ce qui a été préalablement posé ou défini. Cette dynamique est conduite par une directive de définition à caractère ontologique qui permet d'introduire des thèses-définition internes au système. Une telle liberté définitoire permet de représenter progressivement des idées nouvelles dans le système. Elle est rendue possible, rappelons-le, par le fait que toute catégorie syntaxico-sémantique est déterminée de manière contextuelle, autrement dit qu'elle ne dépend pas d'un ensemble prédéterminé de symboles.

Sur la base de ces idées, Lesniewski établira de manière univoque la signification de la copule qu'il utilise dans les discours qui règlent ses déductions logiques. Il en présentera une première version formelle en 1930 sous la forme d'un unique axiome contenant un seul terme primitif, l'épsilon ϵ . Il ne s'agit

en aucun cas du symbole d'appartenance de la théorie classique des ensembles. Ce terme apparaît dans des propositions dites *singulières* dont la forme est la suivante $a \varepsilon b$, et qui peut se lire de la manière présémantique suivante:

$a \varepsilon b$: a est le (ou un des) b .

Les termes a et b représentent des objets formels de la catégorie syntaxico-sémantique des noms. L'épsilon ε est donc un foncteur formateur de propositions à partir de deux arguments de la catégorie des noms, ce que nous désignerons par l'équation formelle suivante: S/NN.

L'évaluation d'une proposition singulière de la forme $a \varepsilon b$ est le vrai si et seulement si toutes les conditions suivantes sont réalisées:

- 1) Le terme a ne représente pas un nom sans dénotation;
- 2) le terme a représente un nom individuel. Ce nom ne peut pas dénoter plus d'un individu;
- 3) si un terme est associé à un nom qui possède la même dénotation que celui associé à a , alors il est en correspondance avec les objets – ou l'objet – dont le nom est associé au terme b .

Cette formulation n'est pas très élégante. La première clause stipule l'existence de a , la deuxième inscrit l'unicité de a et enfin la troisième clause explicite un principe de convergence: tout ce qui pourrait être a est aussi un des b .

Cette signification de l'épsilon de Lesniewski s'exprime au travers de la réalisation formelle suivante dans laquelle $[a]$ peut être lu «quel que soit a » et $[\exists b]$, «il y a b ».

Axiome:

$$\begin{aligned}
 [ab][a \varepsilon b] &\equiv [\exists c][c \varepsilon a] \wedge && \text{(existence)} \\
 [dc][c \varepsilon a \wedge d \varepsilon a] &\supset d \varepsilon c] \wedge && \text{(unicité)} \\
 [c][c \varepsilon a \supset c \varepsilon b] &&& \text{(convergence)}
 \end{aligned}$$

En utilisant cette signification de la copule et en choisissant comme domaine sémantique celui des connaissances naïves

communément partagées, il est possible dès lors d'évaluer les propositions suivantes:

Aristote est un philosophe de l'Antiquité – comme une proposition vraie.

Jean-Paul II est un mathématicien célèbre – comme une proposition fausse. En effet, bien que *Jean-Paul II* existe et soit unique, il n'est pas le cas qu'il soit un mathématicien.

Pégase est un cheval ailé – comme une proposition fausse, parce que Pégase n'existe pas, Pégase ne dénote aucun objet.

L'homme est mortel – comme une proposition fausse, parce que l'homme dans ce contexte est un nom général, il dénote plus d'un objet. En fait, il s'agit de la forme contractée d'une universelle affirmative qui peut s'écrire ainsi:

$$[c][c \in a \supset c \in b]$$

et qui serait vraie.

Nous avons choisi de représenter la quantification d'une autre manière que celle généralement utilisée. Il ne s'agit pas d'une coquetterie, mais d'une nécessité dans la mesure où dans les théories de Lesniewski la quantification ne possède pas le caractère existentiel implicite des logiques classiques, ni celui explicite des logiques libres standard. Dans la perspective d'une interprétation sur un domaine sémantique, elle ne saurait donc être objectuelle. Répétons-le, dans cette théorie, existence et quantification restent deux notions distinctes.

Comme la protothétique, l'ontologie de Lesniewski est une théorie logique et, comme telle, elle contient des directives inférentielles. Elles sont au nombre de sept: une directive de détachement, une de substitution, une directive opérant sur la quantification, deux directives d'extensionnalité et deux directives de définition. Pour des raisons déjà évoquées, nous avons insisté ici uniquement sur les directives de définition.

Sur la base de ces directives et en utilisant les informations que contient l'axiome, à savoir les quatre catégories syntaxico-sémantiques primitives (S, N, S/NN, S/SS), ainsi que les constantes associées à certaines d'entre elles (ϵ , $\equiv$, $\wedge$, $\supset$), il est

possible d'inscrire progressivement de nouvelles constantes quelle qu'en soit la catégorie syntaxico-sémantique.

Nous nous sommes intéressés à définir une négation nominale de la catégorie N/N (Miéville 1991), une opération logique de subordination de la catégorie N/SN, ainsi que quelques opérations d'ordre supérieur (Miéville 1993).

Cette générosité syntaxico-sémantique a une conséquence particulièrement riche. En effet, elle fait reculer les limites élémentaires des descriptions extensionnelles classiques. Le fait de disposer de constantes des catégories syntaxiques suivantes, par exemple, N/N ou N/S, incite à nous inquiéter de la subtilité de la description du modèle sémantique. Et dans cette perspective, la volonté de disposer d'une sémantique capable de représenter l'organisation des diverses relations que soutiennent les parties au tout s'impose de manière quasi naturelle. Une fois encore, nous avons emprunté à S. Lesniewski la définition de la classe collective pour l'interpréter (Lesniewski 1916; Miéville 1984). En effet, cette acception collective de la notion de classe permet de rendre compte de l'aspect pluridimensionnel des objets. Elle offre la possibilité de réunir en un tout, une entité des plus complexes sans pour autant admettre n'importe quoi comme ingrédient. On comprendra donc notre intérêt pour ce modèle. De plus, on peut parler du même objet de discours sous des aspects différents, et là encore le recours au modèle collectif élaboré par Lesniewski se révèle d'une grande utilité. Explicitons-en maintenant les propriétés fondatrices.

La classe collective

Dans un premier temps, illustrons notre propos afin de mieux rendre compte de l'interprétation collective. Abordons la classe des carrés de la figure I. De manière distributive, cette classe est constituée de six et seulement six éléments: les carrés ACEG, BDFH, ABIH, HIFG, BCDI, IDEF.

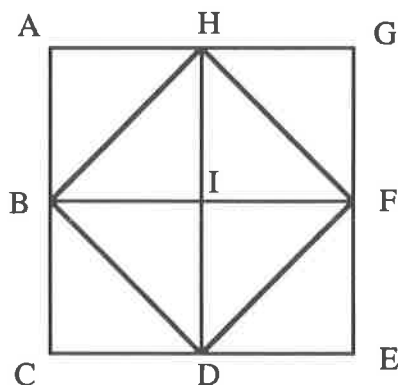


Figure I

Tous les éléments qui la composent sont de même nature: ils tombent tous sous la propriété caractéristique de la classe, autrement dit sous le concept qui l'engendre: être un carré de cette figure géométrique.

La lecture collective de cette classe comportera également ces six éléments, mais bien d'autres choses encore, en effet, les éléments suivants en font notamment partie:

- Les triangles ABH, BIH, HIF, HFG, BCD, BDI, IDF, DEF, BFH, BDF, BDH et HDF.
- Le polygone ABIDEG, entre autres polygones.
- Les segments AB, AC, BC, tout segment contenu dans AC, les points A, B, C, D, tout point contenu dans CD, toute com-

position imaginable de points et (ou) de segments, et (ou) d'autres éléments qui y sont inscrits, appartiennent encore à la classe collective des carrés de la figure I.

- La cardinalité de la classe distributive des carrés de la figure I est de six, celle de la classe collective possède la puissance du continu.
- Par ailleurs, la classe collective des carrés de la figure I est le même objet que la classe des triangles de la figure I. On a donc accès aux mêmes éléments par des modes d'accès conceptuellement très différents.

Lesniewski a défini le concept de classe collective dans le cadre d'une théorie axiomatique du rapport des parties au tout. La relation d'ingrédience qu'il propose est une relation réflexive, antisymétrique et transitive. Le modèle qui domine sa construction, bien que pluriextensionnel, reste extensionnel. Il est probable que sa perception de l'espace géométrique ait guidé ses réflexions. Mais il n'est pas inutile de présenter la définition de la classe collective de manière plus précise. Nous allons nous y employer.

Lorsque Lesniewski s'interroge sur la notion de classe, il refuse d'accepter qu'elle ne soit considérée qu'à l'image d'une commodité linguistique ou symbolique. Il refuse donc l'acception qu'en donne Russell (Whitehead & Russell 1910, I: 75):

The symbols for classes like those for descriptions, are, in our system, incomplete symbols: their "uses" are defined, but they themselves are not assumed to mean anything at all. That is to say, the uses of such symbols are defined such that, when "definiens" is substituted for the "definiendum", there no longer remains any symbol which could be supposed to represent a class. Thus classes, so far as we introduce them, are merely symbolic or linguistic conveniences, not genuine objects as their members are if they are individuals.

Pour Lesniewski, bien au contraire, une classe est une «chose». De manière intuitive, il trouve naturel de considérer qu'un tas de sable existe de manière analogue avec les grains de sable qui le composent. Ses réflexions le conduisent tout à la fois à donner une acception collective au terme de «classe» et à

exposer une théorie générale des ensembles dans laquelle la classe collective est définie: la méréologie (Lesniewski 1916).

C'est la base axiomatique de cette théorie que j'exposerai maintenant.

Axiome I

Quel que soit A , quel que soit B , si A est part de B , alors B n'est pas part de A .

Axiome II

Quel que soit A , quel que soit B , quel que soit C , si A est part de B et B part de C , alors A est part de C .

Définition I

Quel que soit A , quel que soit B , A est ingrédient de B si et seulement si, A existe et (A est identique à B ou A est part de B).

Définition II

Quel que soit A , quel que soit a , A est la classe des a si et seulement si,

* A existe

*il existe un B , B est un des a

*quel que soit B , si B est un des a alors B est un ingrédient de A

*quel que soit B , si B est un ingrédient de A alors il existe C et D tels que: C est un des a , D est ingrédient de C et D est ingrédient de B .

Axiome III

Quel que soit A , quel que soit B et quel que soit a , si A est la classe des a et B la classe des a , alors A est identique à B .

Axiome IV

Quel que soit A , quel que soit a , si A est un des a alors il existe B , B est la classe des a .

Ce qui précède appelle quelques remarques:

- Le seul terme primitif de cette base axiomatique est *part de*. La signification de ce terme est une relation d'ordre irréflexive, asymétrique et transitive.

- Le terme défini, *ingrédient de*, établit une relation d'ordre partiel.
- Si A est une classe de a , alors A est ingrédient de lui-même.
- Une même classe collective A peut être engendrée par des éléments générateurs a et b différents, sans pour autant qu'il y ait un rapport de l'ordre du même entre eux, et cette faculté d'offrir une multiplicité de révélateurs des ingrédients d'une même classe collective nous importe directement.
- Dans la perspective collective, la classe vide n'existe pas.
- La méréologie ne propose pas seulement les significations primitives nécessaires à la définition du concept de la classe collective. Cette théorie s'ouvre également sur la définition de tout un ensemble d'opérations possibles, opérations qui permettent un calcul avec et sur les entités collectives.

Epilogue

Dans les pages précédentes, nous avons exposé les éléments fondamentaux caractérisant une logique développementale qui conserve les principes de bivalence et d'extensionnalité. Nous avons également présenté les significations primitives permettant de fonder une telle logique. De plus, nous avons décrit les moyens inférentiels autorisant toute expansion en termes de constantes nouvelles et donc de catégories nouvelles.

Dans la forme comme dans le fond, nous ne sommes pas en présence d'une théorie logique unique qui aurait la propriété d'être développée progressivement. En fait, nous disposons d'une base axiomatique modeste en termes de significations primitives à partir de laquelle une multitude de systèmes logiques peuvent être élaborés. Cette très grande ouverture entraîne incontestablement au moins deux conséquences épistémologiques importantes: la première est liée, et on l'a largement éprouvée comme telle, à une nouvelle manière de concevoir la formalisation. La liberté expansive du système implique une indétermination de la syntaxe qui entraîne, à son tour, une détermination contextuelle de la signification des inscriptions d'une expression. La deuxième conséquence est liée au mouvement

progressif dans l'espace et le temps qui participe à l'élaboration de tout système. Ce temps est celui de l'expansion de la connaissance opératoire du système en développement. Il marque les étapes des modifications fortes du système; celles qui s'inscrivent par le biais de définitions créatives. Une définition est qualifiée de *créative* lorsqu'elle donne accès à des théorèmes qui ne contiennent pas le terme nouveau défini, mais que l'on ne saurait déduire sans cette signification nouvelle.

L'expression du temps que ces systèmes fixent n'est pas celle liée aux opérations de vérité. Il s'agit davantage de l'expression du temps associée à l'idée d'expansion et de croissance; cette croissance est celle qui rattache les compétences nouvelles à celles des stades antérieurs. Ce mouvement génétique est à même de pouvoir représenter une filiation de structures. Ainsi, une exploitation possible d'un système développemental est la possibilité de mettre en oeuvre la directive de définition de manière à marquer l'accès progressif au groupe des quatre transformations INRC, groupe qui spécifie le stade des opérations hypothético-déductives tel que Piaget et son école l'ont révélé (Piaget 1949). Des travaux sont en cours dans cette perspective.

Ainsi, un engagement dans la logique développementale ouvre sur de nouvelles aventures intellectuelles; il s'agit d'aventures où le temps lié au développement de la connaissance logique prédomine. Elles nous entraînent à enrichir constamment l'édifice des opérations logiques. Elles nous conduisent également à investiguer toujours davantage la structure de la sémantique formelle.

*Séminaire de logique
Université de Neuchâtel
Espace Louis-Agassiz 1
2000 NEUCHÂTEL*

Références

- CARNAP R. (1949). *The Logical Syntax of Language*. London: Routledge & Kegan.
- CHAZAL G. (1995). *Le miroir automate*. Champ Vallon: Seyssel.
- LESNIEWSKI S. (1916). Podstawy ogólnej teorii mnogości, I. (Les fondements d'une théorie générale des ensembles). *Prace polskiego kola naukowego w Moskwie. Sekcyz matematycznoprzyrodnicza*, 2. (Traduction anglaise anonyme, 1976).
- LUSCHEI E.C. (1962). *The Logical System of Lesniewski*. Amsterdam: North Holland.
- MIÉVILLE D. (1984). *Un développement des systèmes logiques de Stanislaw Lesniewski. Protothétique, ontologie, méréologie*. Berne: Lang.
- MIÉVILLE D. (1991). *La négation, une étude logique*. Université de Neuchâtel: Centre de Recherches Sémiologiques, (Travaux de logique, n° 6).
- MIÉVILLE D. (éd.) (1993). *Relations formelles et relations non formelles*. Université de Neuchâtel: Travaux du Centre de Recherches Sémiologiques, n° 61.
- MIÉVILLE D. (éd.) (1997). *Temps, logique et langage*. Actes du symposium tenu lors du colloque international «Penser le temps». Neuchâtel, 8-10 octobre 1996. Université de Neuchâtel: Travaux du Centre de Recherches Sémiologiques, n° 65.
- MIÉVILLE D. & VERNANT D. (sous la dir.) (1995). *Stanislaw Lesniewski aujourd'hui*. Grenoble: Département de philosophie / Université de Neuchâtel: Centre de Recherches Sémiologiques.
- PIAGET J. (1949). *Traité de logique. Essai de logique opératoire*. Paris: Colin.
- RUSSELL B. (1903). *The Principles of Mathematics*. London: Allen & Unwin.

- SURMA S.J., SRZEDNICK J.T, BARNETT D.I. (eds) (1992).
S. Lesniewski. *Collected Works*. Dordrecht: Kluwer
Academic Pub.
- TARSKI A. (1974). *Logique, sémantique, métamathématique*.
Paris: Colin
- WITEHEAD A.N. & RUSSELL B. (1910). *Principia Mathematica*.
Cambridge: Cambridge University Press.

PARADOXICAL SIMPLIFICATION

Helmut LINNEWEBER-LAMMERSKITTEN

In the following paper¹ I will construct a family of alternative propositional modal systems in which the paradoxical "from contradictories" formulas of simplification

$p \& \sim p \rightarrow p$ and

$p \& \sim p \rightarrow \sim p$

are not derivable. I will start with a short characterisation of what in my opinion are the *main points* in alternative logics in general (Section 1) and sketch a pragmatic point of view concerning inferential situations (Section 2) that is important as a background for the reconstruction of the pretheoretic notion of inferential validity. In Section 3 I will explain, why from this point of view the "from contradictories" formulas of simplification and, as a consequence, the general formulas of simplification should be avoided. In Section 4 I will formulate the requirements an alternative system without general principles of simplification ("WGPS-system") should fulfil. Section 5 to 9 contain the construction and description of a family ("SIM1") of such systems: starting with a weak axiomatic system SIM1₀ (Section 5), which satisfies all of the requirements of Section 4, the members of the family are obtained by a stepwise addition of axioms and thus form a sequence ordered by inclusion (of the sets of theses). The system MS1.0 (Section 10) determined by a system of 4-valued matrices can be regarded as an upper bound and thus provides a guideline for the construction of the SIM1 systems.

1 The paper summarises some results of my Habilschrift "Systems without general principles of simplification", which will appear in 1998.
Travaux de logique, 11, 1997.

1. Alternative logics

There are three different striking points in which some alternative logicians like C. I. Lewis, W. Ackermann, A. R. Anderson and N. B. Belnap e.a. are dissatisfied with the conception and the results of the classical propositional calculus of the *Principia Mathematica*² and the best way to mark these points is by making a distinction between the pretheoretic understanding of our propositions and inferences on the one hand and the formal reconstruction of these propositions and inferences by a logical theory on the other hand³. The first point concerns the reconstruction of *hypothetical propositions* as expressed by indicative "if p, then q" statements for instance "If Peter will marry Jacqueline, then he will have a lot of trouble with his father". For the pretheoretic *truth conditions* for such propositions are *more rigorous* than those of the corresponding reconstruction by means of the material implication connective, i.e. there are hypotheticals that are false in a pretheoretic sense, while the corresponding theoretical reconstruction is true. The second point concerns certain *inferences containing "if p, then q" structures* within the premises or the conclusion. The difference of the truth conditions between statements in which the pretheoretic "if p, then q" structure is used and their reconstruction by means of material implication has consequences with respect to the validity of such inferences. For inferences that are pretheoretically invalid may be valid according to the reconstruction, if the replacement of the "if p, then q" statement by material implication results in a weakening of the conclusion or in a strengthening of the premises; inferences, on the other hand, that are pretheoretically valid, may become theoretically invalid, if the replacement has the

2 I. e. the propositional calculus developed in the first volume of the *Principia Mathematica* (cf. Whitehead and Russell 1927).

3 For more details and a presentation of alternative logics that differs in some respects from my own presentation, cf. Haack 1996, Routley and Plumwood and Meyer and Brady 1982 (with a critique of Haack's classification of alternative logics 57ff.), Read 1988, Anderson and Belnap 1990, Anderson and Belnap and Dunn 1992.

effect of a strengthening with respect to the conclusion or of a weakening with respect to the premises⁴. The third point concerns the reconstruction of the pretheoretic *notions of inference and/or inferential validity themselves*. Again, the pretheoretic notions are *more rigorous* than the corresponding theoretic ones, i.e. there are inferential structures that are invalid in a pretheoretic sense (or rejected as not representing an inference in a proper sense), while the corresponding theoretic reconstruction of the pretheoretic notions of inference and inferential validity classes them as valid.

Ad 1: Though there may be some ordinary "if p, then q" statements in the pretheoretic sphere whose truth conditions coincide exactly with the truth conditions of material implication – i.e. "if p, then q" statements that are true, if the antecedent p is false or the consequent q is true, and that are false only, if the antecedent p is true and the consequent is false – this does not hold generally. For while the contradictory of such an "if p, then q" statement is the conjunction of the antecedent and the negation of the consequent: "p and not q", there are other "if p, then q" statements where the relation to the "p and not q" statement is only a relation of contrariety and not of contradiction, i.e. it is possible that both statements are false. Even if the truth conditions of this second type of "if p, then q" statements are not sufficiently clear, the difference in the relation to the corresponding conjunction statements shows at least, that statements of this kind cannot be satisfyingly reconstructed by material implication.

Ad 2: The most famous examples of such inferences are the so called "paradoxes of material implication":

$$\begin{array}{rcl} \sim p & & \text{"ex falso quodlibet" paradox} \\ \hline p \rightarrow q \end{array}$$

4 For some examples of inferences that are pretheoretically valid, but invalid according to the classical reconstruction, cf. McCall (1966). In the following I will concentrate on alternative reconstructions that lead to "regular" systems (cf. Definition 3), i.e. systems that do not allow inferences which would be invalid in the classical reconstruction.

$$q$$

"quodlibet ad verum" paradox

$$p \rightarrow q$$

but there are also other puzzling examples (due to the fact, that the negation of a material implication statement provides a stronger premise than the negation of the pretheoretic "if p, then q" statement):

$$\sim(p \supset q)$$

$$p$$

$$\sim(p \supset q)$$

$$p \supset \sim q$$

$$\sim(p \supset q)$$

$$q \supset p$$

Ad 3: According to the classical reconstruction of validity, an inference is valid just in case it is impossible for its premises to be true and its conclusion false. But this notion has two strange "limit cases"⁵: first, if it is already impossible for the premises to be true together, then an inference composed of these premises and an arbitrary conclusion is valid. Hence for instance

$$p$$

$$\sim p$$

$$q$$

as well as

5 Provided the term "impossible" is used in such a way that the condition "it is impossible for its premises to be true and its conclusion false" is already satisfied if only one of the conditions "it is impossible for its premises to be true" or "it is impossible for its conclusion to be false" holds.

$$\begin{array}{c}
 p \\
 \sim p \\
 \hline
 \sim q
 \end{array}$$

is valid according to the classical theory. But according to our pretheoretic understanding such constellations of propositions do not count as valid inferences⁶: it would be a bad idea to start with a contradictory pair of premises *in order to get a new insight* q , (since, beside other things, the 'expected insight' q is annihilated by $\sim q$), and it would be a bad idea as well to refer to a contradictory pair of premises *in order to give a reason for* a proposition q , (since, beside other things, the same reason could be given for the contradictory $\sim q$). Secondly if it is already impossible for the conclusion to be false, then an inference composed of this conclusion and some arbitrary premises is valid. Thus according to the classical theory the following pair

$$\begin{array}{c}
 p \\
 \hline
 \sim(q \& \sim q)
 \end{array}$$

and

$$\begin{array}{c}
 \sim p \\
 \hline
 \sim(q \& \sim q)
 \end{array}$$

is valid. But again: according to our pretheoretic understanding such constellations of propositions do not count as valid inferences (at least not by themselves): it would neither be a good idea to start with an arbitrary premise p in order to *derive* the tautological statement $\sim(q \& \sim q)$ *as a new insight*, nor would it be a good idea to refer to an arbitrary premise p *in order to give a reason for* a tautological statement.

There are three extreme strategies to settle the conflict between the pretheoretic understanding and the reconstruction given by the classical theory. Either we try a) to keep the classical theory as well as the pretheoretic understanding

6 Though they may, for other reasons, be tolerated as limit cases.

untouched – then we have to modify the relation between these spheres, i.e. we have to give a reinterpretation of what is reconstructed by the reconstruction. Or we try b) to keep the classical theory as well as the aims of the reconstruction without any restriction – then we have to modify our pretheoretic understanding. Or c) we try to keep our pretheoretic understanding and the aims of the construction – then we have to modify the classical theory.

Classical logicians typically use strategy a) to solve the truth condition problem of "if p, then q" statements (problem 1): the connective of material implication is no longer interpreted as an adequate reconstruction of ordinary "if p, then q" statements *in general*, but as a reconstruction of a special type of such statements, namely of "if p, then q" statements that are (in a pretheoretic sense) equivalent to "not (p and not-q)" statements. Concerning the validity problem (problem 3) they typically use strategy b): a pretheoretic understanding of validity that is in trouble with the just mentioned inferences has to be mistaken, since it is thought that the only *logical* condition that has to be satisfied by a valid inference is that it is impossible for its premises to be true and its conclusion false. Concerning the problem of the "paradoxes of material implication" and other (at least *prima facie*) strange inferences they typically apply a combination of strategy a) and b): since, on the one hand, the reconstruction is restricted to a special subclass of "if p, then q" statements, namely those that satisfy the truth conditions of material implication, and, on the other hand, the classical reconstruction of the notion of inferential validity is regarded as adequate, these inferences are not paradoxical at all: the validity of these inferences can be proved and the putative paradox or strangeness is thought of as resulting by a misunderstanding due either to the neglect of the intended restriction or to a mistaken idea of what "inferential validity" means (or to a combination of both).

Alternative logicians on the other hand would argue, that by restricting the notion of material implication to the reconstruction of negated conjunction statements (problem 1) the classical theory fails to give an adequate reconstruction of

ordinary "if p, then q" statements, though they will willingly admit, that the notion of material implication provides an adequate reconstruction of statements that are (in a pretheoretic sense) equivalent to negated conjunction statements. An adequate reconstruction of ordinary "if p, then q" statements, however, would require a notion of implication that is stronger than that of material implication, so, for instance, C. I. Lewis proposed "strict implication" (Lewis 1959), Ackermann "strong implication" (Ackermann 1956), Anderson, Belnap e. a. different sorts of "entailment" (Anderson 1990), McCall "connexive implication" (McCall 1966) etc.

Concerning the notion of inferential validity (problem 3) C. I. Lewis and W. Ackermann (and other contemporary alternative logicians) did not hold the same opinion, though both rejected material implication as inadequate for the reconstruction of "if p, then q" statements as well as for the definition of inferential validity. While Lewis was convinced and tried to prove that the "from contradictories" inferences must be part of every logical calculus (what he in fact showed was that the formulas corresponding to these inferences must be theorems in every logical calculus in which certain other principles are derivable⁷), Ackermann rejected these inferences as invalid, though, of course, he would have conceded that they satisfy the *classical* criterion of validity. That means, that according to Ackermann the classical criterion formulates only a necessary, but not a sufficient condition, i.e. that we need not only a stronger notion of implication for the reconstruction of "if p, then q" statements, but also a stronger notion of inferential validity for the reconstruction of our pretheoretic understanding of inferential validity.

Since alternative logicians regard material implication as inadequate for the reconstruction of "if p, then q" statements, they reject inferences like the paradoxes of material implication and similar inferences (problem 2) as invalid, if the material implication is replaced by a stronger implicative connective. Moreover, from the point of view of an alternative logician "if p, then q" statements have to be *reconstructed* and the notion of

7 Cf. below footnote 22.

implication has to be *determined* in such a way that these inferences turn out to be invalid. Typically, the determination of valid inferences (and the determination of invalid inferences) in alternative logics is not intended as the result of applying previously defined truth conditions and a previously defined notion of inferential validity, but as the first step in the direction of such definitions: the truth conditions for "if p, then q" statements and the notion of inferential validity has to be defined in such a way that the inferential requirements are fulfilled.

The problem of the paradoxes of material implication (problem 2) requires a solution either by a change with regard to the reconstruction of "if p, then q" statements (problem 1) or with regard to the reconstruction of the notion of inferential validity (problem 3) or by a change of both reconstructions. Indeed, it is sufficient for the solution of problem 2 to solve only one of the other two problems, i.e. it is sufficient either to replace the classical reconstruction of "if p, then q" statements or that of the notion of inferential validity. Thus, for instance, in modal systems like K, D, S4, S5⁸ etc. a notion of strict implication can be introduced by the definition: $p \rightarrow q := \sim M(p \supset q)$, which avoids the unwanted parts of the truth conditions (problem 1) as well as the paradoxes (problem 2), while it leaves the notion of inferential validity untouched (problem 3). On the other hand, one can keep the reconstruction of the "if p, then q" statements by material implication (problem 1) and solve problem 2 and 3 by imposing some restrictions on the classical reconstruction of inferential validity⁹.

But though it is possible to solve problem 1 and problem 3 separately, it is not only aesthetically more elegant but also philosophically preferable to solve both problems at once, namely by choosing an alternative kind of implication that is a) *stronger* than material implication and that b) can, at the same time, be used for the reconstruction of "If p, then q" statements and as main part within the reconstruction of the notion of

8 Cf. Hughes & Cresswell 1968, 51 et seq.

9 For instance, restrictions similar to those formulated by the von Wright/Geach/Smiley condition (cf. Anderson & Belnap 1990, 215 et seq.).

inferential validity. For though classical propositional logic fails in both cases, its main idea – to reconstruct the notion of inferential validity as a special case of material implication – is philosophically most interesting. Let us briefly re-examine how this idea works. The minimal requirements for the *pretheoretic* truth of "if p , then q " statements (that it is not the case that p is true and q false) and the minimal requirements for the *pretheoretic* notion of inferential validity (that it is for logical reasons not the case that the premises are true and the conclusion false) differ only with respect to the addition "for logical reasons". The classical approach, which reduces the truth conditions and the notion of validity exactly to these minimal requirements, is able to use material implication for the reconstruction of the notion of inferential validity by means of the so called *law of conditionalization*:

*Law of conditionalization (classical version)*¹⁰: An inference from premises $x_1, \dots, x_n$ to a conclusion y is *valid*, iff the corresponding¹¹ wff $x_1 \& \dots \& x_n \supset y$ is a thesis in K_c (where K_c is the classical calculus of propositional logic and the implication symbol " $\supset$ " means material implication.)

If we keep the formal connection expressed by the law of conditionalization, but replace the classical propositional calculus by some alternative propositional calculus and the connective of material implication by a stronger kind of implicational connective, we arrive at the following:

Law of conditionalization (alternative version): An inference from premises $x_1, \dots, x_n$ to a conclusion y is

-
- 10 Though the term "conditionalization" and a law similar to the principle formulated here are often used in the context of the so-called "deduction theorem" my main point here is, that it functions as a *bridge* principle between the sphere of inferential situations (cf. 12) and the sphere of propositional logic and not between two different formal techniques within the second sphere (cf. also Linneweber-Lammerskitten 1997).
- 11 Note that there are a lot of problems involved in connection with the formalization of inferences expressed in ordinary language. Some of these problems belong to logic especially the question of what is meant by "form", others belong to the semantics of the ordinary language, in which the inference is formulated. In the following text I silently presuppose that these problems can be solved in an ordinary inferential situation, i.e. that a speaker who claims that he infers something from something (else) is able to give a formal reconstruction of his inference in terms of the formal language.

valid, iff the corresponding wff $x_1 \& \dots \& x_n \rightarrow y$ is a thesis in K_a (where K_a is an alternative calculus of propositional logic and the implication symbol " $\rightarrow$ " means an alternative type of implication.)

It is essential to distinguish between the different calculi as such and the different versions of the laws of conditionalization. The replacement of the classical calculus by some alternative calculus does not by itself lead to an alternative notion of inference. It is the law of conditionalization that determines whether the notion of inferential validity is alternative or classical.

In its classical version the law of conditionalization *reduces* the notion of inferential validity to a special case of material implication and it is thus understandable that classical logicians tend to use strategy b) with respect to problem 3. The alternative version of the law of conditionalization in contrast cannot be used for such a reduction, since the truth conditions of an alternative kind of implication are not truth-functional. We can in a first step only give a sufficient condition for the falsehood and a necessary (but negative) condition for the truth of such statements. But we can use the law of conditionalization in the other direction and formulate in a second step restrictions to possible truth conditions: if, for instance, the transitivity formula

$$*1 \quad (p \rightarrow q) \& (q \rightarrow r) \rightarrow (p \rightarrow r)$$

is a thesis of an alternative calculus K (where " $\rightarrow$ " stands for the alternative implicative connective of K), then the truth conditions for a simple implication statement $p \rightarrow q$ must be such that (*1) is a true proposition whatever propositions we choose for the symbols p , q and r . But though it is true that, if (*1) is a thesis of K , this indicates via the law of conditionalization that the corresponding inference is valid according to K , the relation that is more important goes in the other direction: because the inference allowing the transition from implicative premises $p \rightarrow q$ and $q \rightarrow r$ to an implicative conclusion $p \rightarrow r$ is regarded as valid (whatever the exact reconstruction of "if p , then q " statements by the implication " $p \rightarrow q$ " may be), K has been constructed in such a way, that (*1) becomes a thesis of K and this in turn

restricts the possible reconstructions of the "if p, then q" statements.

If we compare the different usage that is made of the law of conditionalization in the classical and in the alternative case, we can say, that classical logic starts with the reconstruction of "if p, then q" statements and uses the law to reconstruct the notion of inferential validity, while alternative logic, typically, starts with the reconstruction of the notion of inferential validity (by singling out certain *inferences* of the classical approach as invalid and others as fundamental) and uses the law to single out certain complex implicative *propositions* as formally true and others as not formally true, and thereby determines the formal truth conditions of the simple implicative statements.

2. Non-classical aspects of inference

The *classical* reconstruction of the notion of inferential validity abstracts from most of the pretheoretic aspects of an inferential situation and concentrates on one single point: that inference is a transition from a set of premises to a conclusion and that the transition is valid iff the logical form of the propositions does not admit inferences from true premises to a false conclusion. However, this point is only one of many other aspects relevant for ordinary inferential situations, or for special subgroups of such situations¹². Many of these other aspects are logically irrelevant, but there are some non-logical aspects that depend on logical principles and there are also some *logical aspects* beyond the requirement of truth preservation. The most important of the last group, in my opinion, is (premise-) *relevance*: (at least one) of the premises should be relevant for the conclusion, i.e. there should be a logical connection between the (structure of the) premises and the (structure of the) conclusion. However this may be interpreted in detail, a necessary condition in the case of *propositional* logic is surely

12 By (the "idealtypus" of) an inferential situation I mean, roughly speaking, a situation, in which some rational beings with different beliefs mutually try to give reasons for their own position and mutually try to understand, what exactly the position of the others consists in and how it is grounded.

that (in the formalization of the inference) the premises and the conclusion must have at least one propositional variable in common, for otherwise this connection could not be expressed by means of propositional logic.

Though many of my considerations have been influenced by works on relevant logic (esp. Anderson 1990 and 1992) there is a non-logical aspect which, in my opinion, deserves some interest, namely that the *theoretical* question of inferential validity has *practical imports* concerning *rationality* and *obligation* in inferential situations. Thus for instance we would say that it is *rational* for someone who accepts some premises $x_1, \dots, x_n$ to pass to the conclusion y , provided the *inference is valid*; or we might ask "Am I *obliged* to pass from the premises $x_1, \dots, x_n$, which I granted for the sake of argument, to a certain conclusion y , or not?", "Am I *allowed* to pass from the premises $x_1, \dots, x_n$, which express my belief at t_0 , to the conclusion y , or not?"¹³, etc. Though inferential validity is a theoretical matter it *takes its importance* from the fact, that it is relevant for our intellectual *actions*, for what we do in conversation with others as well as for what we do in silent reflection; in turn the practical validity of these actions (in concreto) is dependent on the theoretical validity of the inferences at stake.

In an inferential situation we may be obliged to accept "for the sake of argument" premises to which we normally would not assent – but we are not obliged to accept such premises beyond necessity. We can ask for what kind of inference the premises are needed – and if we do not accept these inferences as valid, we are not obliged to accept the premises, since there is no (accepted) argument for the sake of which the acceptance of the premises could be demanded. If we do not accept certain inferences, we are not allowed to use them and we are not obliged to accept their use by others (we are not obliged to accept the conclusion temporarily, since we are not obliged to accept the premises).

13 Obligation in inferential situation was an important subject in the logic of the fourteenth and fifteenth century (cf. for instance Paul of Venice 1988; there are also treatises on obligation by Walter Burley, Roger Swyneshed, Albert of Saxony (1360-1390) and others) – though it seems that the interest in obligation was less directed to moral duty than to strategies for disputation.

Though the question what kind of obligation and what kind of rights one has in an inferential situation is not a logical question, its answer presupposes some logical investigation, since for the reconstruction of these obligations and rights in detail a logical reconstruction of the notion of inferential validity is needed. Of course, the latter can only be expected to succeed with respect to the non-logical aim, if it takes non-classical aspects of inferential situations serious. But this does not mean that the reconstruction thereby becomes a non-logical (psychological, sociological etc.) enterprise, for it is not the pretheoretic notion of persuasion, of reasoning as psychological act, of arguing as a social event etc. that is at stake here, but it is the pretheoretic notion of inferential validity, that is to be reconstructed.

3. What is wrong with simplification?

In *Principia Mathematica* B. Russell and A. N Whitehead refer to the following three formulas as "principles of simplification"¹⁴ which are theorems of the PM calculus:

$$*2 \quad q \rightarrow (p \rightarrow q)$$

$$*3 \quad (p \& q) \rightarrow p$$

$$*4 \quad (p \& q) \rightarrow q$$

The first expresses one of the well-known paradoxes of material implication that corresponds to the "verum sequitur ad quodlibet"-rule of the scholastic logicians¹⁵, the other two

14 In the notation used by Russell and Whitehead these formulas read (cf. Whitehead and Russell 1927, XI):

* 2·02. $\vdash : q \supset . p \supset q$ Pp. ["Simp"]

* 3·26. $\vdash : p \cdot q \supset p$ Pp. ["Simp"]

* 3·27. $\vdash : p \cdot q \supset q$ Pp. ["Simp"]

15 It should be noted that this rule formulated for instance in Buridanus' *De consequentiis*: "Et est notandum quod (...) omnis uera ad omnem aliam sequitur etiam consequentia ut nunc." (Buridanus 1976, 32) is part of a logical theory that differs in some respect from the theories we are accustomed with today. Buridan's rule does *not* mean, that we can pass from a true *premise* q to a true *conclusion* $p \rightarrow q$. It is *not* the rule corresponding to the syllogistic scheme:

formulas are sometimes called formulas of Petrus Hispanus¹⁶ - I will use the name "paradox of material implication" for (*2) and reserve the expression "*simplification*" for expressions of the conjunctive form (*3) and (*4) - the latter will be called "*general formulas of simplification*", if contrasted with its substitution instances for instance with the "*from contradictories formulas of simplification*"

*5 $p \& \sim p \rightarrow p$ and

*6 $p \& \sim p \rightarrow \sim p$

Since the paradox concerning formula (*2) has been one of the main reasons for the development of the systems of strict implication and has got a detailed discussion by C. I. Lewis and his school, I will concentrate on the other two formulas¹⁷ (*3) and (*4).

Prima facie there is nothing paradoxical with the general formulas of simplification (*3) and (*4)¹⁸. But they have paradoxical consequences in the sense that they can be used together with other theses and rules in order to derive paradoxical formulas like the paradox of material implication (*2). But before I can give a short survey of such hidden traps of paradoxes I must first make precise what is meant by the phrase

q

$p \rightarrow q$

rather it corresponds to the syllogistic scheme:

p

ææ

q

which is not valid with respect to every proposition in every situation, but only with respect to a situation in which the sentence q is actually true. Thus for instance the inference

Socrates runs

A white (pale) man runs

is valid in all situations in which a white man runs.

- 16 This is not quite correct, since they are not mentioned by Petrus Hispanus himself, but appear in a later commentary by a certain Versorius.

- 17 Despite the fact that formula (*2) is paradoxical in itself, the arguments for the renunciation of (*3) and (*4) will provide another reason for the rejection of (*2). For as long as the importation formula

$(p \rightarrow (q \rightarrow r)) \rightarrow ((p \& q) \rightarrow r)$

is a thesis, (*3) could be derived from (*2).

- 18 Nelson (1930) held a more radical view.

"formula x has certain (paradoxical) consequences y (as long as certain other rules and theses hold)", for a paradoxical import of a formula depends on the system in which it is a thesis.

Definition 1: An axiomatic system of propositional logic (based on a language L) is a *DCSE-system*, iff the rules of detachment, conjunction, substitution and the equivalence rule¹⁹ hold (i.e. if these rules are part of the axiom-basis or introducible).

Definition 2: Let $x, y, z_1, \dots, z_n$ be formulas of a language L . A formula x has the consequence y , if $z_1, \dots, z_n$ are theses, iff y is a thesis in every DCSE-system (based on the language L) in which $z_1 \dots z_n$ are theses.

Using these definitions we can give some examples of the paradoxical import resulting from the general formulas of simplification:

Theorem 1: The general formula of simplification (*3) has the paradoxical consequence

$$*7 \quad p \rightarrow (q \rightarrow p)$$

("verum sequitur ad quodlibet" paradox), if the formula of exportation

$$*8 \quad ((p \& q) \rightarrow r) \rightarrow (p \rightarrow (q \rightarrow r))$$

is a thesis²⁰.

Theorem 2: The general formula of simplification (*3) has the paradoxical consequence

$$*9 \quad (p \& \sim p) \rightarrow q$$

("from contradictories" paradox), if the formula of antilogism

$$*10 \quad ((p \& q) \rightarrow r) \rightarrow ((p \& \sim r) \rightarrow \sim q) \text{ and the equivalence formula}$$

¹⁹ For an explicit formulation of these rules cf. Section 5.

²⁰ *Proof:* a) $((p \& q) \rightarrow p) \rightarrow (p \rightarrow (q \rightarrow p))$ (*8) [p/r]
 b) $((p \& q) \rightarrow p)$ (*2)
 c) $p \rightarrow (q \rightarrow p)$ (a) (b) X modus ponens

*11 $\sim\sim q \leftrightarrow q$
are theses²¹.

Theorem 3: The general formulas of simplification (*3) and (*4) have the paradoxical consequence

*12 $(p \& \sim p) \rightarrow q$
("from contradictories" paradox), if the following formulas

*13 $p \rightarrow (pvq)$

*14 $((p \rightarrow q) \& (q \rightarrow r)) \rightarrow (p \rightarrow r)$

*15 $((p \rightarrow q) \& (p \rightarrow r)) \rightarrow (p \rightarrow (q \& r))$

*16 $((pvq) \& \sim p) \rightarrow q$

are theses²².

These paradoxical consequences are not "absolute consequences", but consequences under the condition that other formulas are theses. This means that the paradoxes can be avoided by renouncing one of the additional theses. Thus Lewis and Langford in their systems of strict implication renounce the *formula of exportation*, but accept the "*from contradictories*" paradox; Ackermann in his system of strong implication

- 21 *Proof:* a) $((p \& \sim q) \rightarrow p) \rightarrow ((p \& \sim p) \rightarrow \sim\sim q)$ (*10) $[\sim q/q \ p/r]$
b) $((p \& \sim q) \rightarrow p)$ (*2) $[\sim q/q]$
c) $(p \& \sim p) \rightarrow \sim\sim q$ (a) (b) X modus ponens
d) $(p \& \sim p) \rightarrow q$ (c) (*11) X equivalence

q.e.d.

- 22 The ideas of the following proof are due to Lewis and Langford (1959, 250f.) – The main ideas of the proof were already formulated by Buridanus in the 14th century (cf. Buridanus 1976, 37, 169–181).

Proof:

- a) $(p \& \sim p) \rightarrow p$ (*2) $[\sim p/q]$
b) $p \rightarrow (pvq)$ (*13)
c) $((p \& \sim p) \rightarrow p) \& (p \rightarrow (pvq))$ (a) (b) X conjunction
d) $((p \& \sim p) \rightarrow p) \& (p \rightarrow (pvq)) \rightarrow ((p \& \sim p) \rightarrow (pvq))$ (*14) $[(p \& \sim p)/p \ p/q \ (pvq)/r]$
e) $(p \& \sim p) \rightarrow (p \vee q)$ (c) (d) X modus ponens
f) $(p \& \sim p) \rightarrow \sim p$ (*4) $[\sim p/q]$
g) $((p \& \sim p) \rightarrow (pvq)) \& ((p \& \sim p) \rightarrow \sim p)$ (e) (f) X conjunction
h) $((p \& \sim p) \rightarrow (pvq)) \& ((p \& \sim p) \rightarrow \sim p) \rightarrow ((p \& \sim p) \rightarrow ((pvq) \& \sim p))$
(*15) $[(p \& \sim p)/p \ (pvq)/q \ \sim p/r]$
i) $(p \& \sim p) \rightarrow ((pvq) \& \sim p)$ (g) (h) X modus ponens
j) $((pvq) \& \sim p) \rightarrow q$ (*16)
k) $((p \& \sim p) \rightarrow ((pvq) \& \sim p)) \& (((pvq) \& \sim p) \rightarrow q)$ (i) (j) X conjunction
l) $((p \& \sim p) \rightarrow ((pvq) \& \sim p)) \& (((pvq) \& \sim p) \rightarrow q) \rightarrow ((p \& \sim p) \rightarrow q)$
(*14) $[(p \& \sim p)/p \ ((pvq) \& \sim p)/q \ q/r]$
m) $(p \& \sim p) \rightarrow q$ (k) (l) X modus ponens

q.e.d.

renounces the *formula of antilogism*, the *formula of disjunctive syllogism* etc.²³.

But beside these paradoxical consequences conditional on some other theses, the formulas (*3) and (*4) have also *immediate* consequences, which are in a certain sense paradoxical:

$$*17 \quad (p \& \sim p) \rightarrow p$$

$$*18 \quad (p \& \sim p) \rightarrow \sim p$$

Since these formulas are substitution instances of the general formulas, they can only be avoided by renouncing the general formulas of simplification themselves²⁴.

In which sense are (*17) and (*18) paradoxical? Obviously they are not only substitution instances of the general formulas of simplification, but also substitution instances of the "*from contradictories*" paradox

$$*19 \quad (p \& \sim p) \rightarrow q$$

This is not by itself a sufficient argument, since substitution instances of a paradoxical formula need not be paradoxical themselves, for instance

$$*20 \quad (p \& \sim p) \rightarrow (p \& \sim p)$$

is a *non-paradoxical* substitution instance of (*19). But it may be helpful to have a look at (*19) in order to get a better understanding of (*17) and (*18). I think there are three main points in which our pretheoretic notion of inference collides²⁵ with (*19): The *first* point is "missing relevance": we expect that

23 For a detailed survey cf. Linneweber-Lammerskitten 1994.

24 There is still another alternative, namely to give up the calculus rule of substitution: by giving up this rule we could keep (*3) and (*4) as theses and could get rid of (*17) and (*18). But disregarding the technical problems, this solution would undermine the idea that the theses in a logical calculus should not only guarantee the validity of their "material" substitution instances outside the calculus (i.e. the concrete inferences used in everyday life reasoning), but should also guarantee the theoremhood of their "formal" substitution instances within the calculus.

25 Of course, no thesis of the propositional calculus of PM is by itself a paradox: the sort of paradoxes I want to discuss here arise, since by the acceptance of PM as a reasonable systematisation of our intuitions concerning inferences, we are forced to accept inferences that contradict our previous intuitions or some of their consequences.

the conclusions of our inferences have something to do with the premises they are drawn from, but (*19) generates a host of examples of inferences in which the premises have nothing to do with the conclusion. The *second* point is that we expect that whenever there is a valid inference from premises x to a conclusion y then there cannot be also a valid inference from x to the negation of y , but by accepting (*19) we also accept its substitution instance

$$*21 \quad (p \& \sim p) \rightarrow \sim q$$

hence a formula whose consequent term is the negation of that in (*19). The *third* point is, that we expect – quite contrary to the school tradition of scholastic logic – that from a plain contradiction "as little as possible" follows²⁶.

Among these three points of paradox concerning the formula (*19), the first vanishes if we pass from (*19) to its substitution instances (*17) or (*18). For p or $\sim p$ indeed have something to do with the premise $p \& \sim p$: the former are constituents of the latter; (*17) and (*18) seem to be expressions of an ordinary analysis with respect to a complex proposition. I will criticise this view later, but I agree that the first objection against (*19) is not an objection against (*17) or (*18). The second and the third argument against (*19), however, can be applied to (*17) and (*18).

According to the second, the two formulas provide an example of inferences that lead from the very same premises to contradicting conclusions. Notice that this does not mean that the conclusions themselves are contradictories – p and $\sim p$ may well be contingent propositions – the point is that (*17) and (*18) as principles of inference give rise to two different patterns of inference with the same premise-structure, but with conclusions that cannot be true together. Why do we normally not want to have inference patterns with that property? I think it is because the first thing we expect from inferences is progression – progression concerning our own knowledge or

26 Of course we cannot avoid that the position of a contradiction (cf. *20) follows from itself (at least as long as we accept the identity formula of implication $p \rightarrow p$ as a thesis). Another formula we cannot avoid (if we keep $p \rightarrow Mp$) is $p \& \sim p \rightarrow M(p \& \sim p)$.

progression concerning the understanding of others or their understanding of us. But the acceptance of inference patterns that lead to contradicting conclusions frustrates the hope for such progress, since a) one such conclusion can at once be abolished by the second pattern and b) a complex proposition x cannot be used as an account for an opinion q , if it could also be used as an account for the opposite opinion. Now it may be objected, that – granted that we normally do not want to have such inference pattern – in this special case we should have them as instrument to detect and abolish contradictions, i.e. we should have them since otherwise we could not have a *reductio ad absurdum* instrument. But this is not true: the systems I want to propose provide the necessary *reductio ad absurdum* facilities.

The third argument aims at the point, that in discourse we should not be forced beyond necessity to accept (provisionally, for the sake of argument) a plain contradiction ("Imagine $p \& \sim p$ " – this is rather hard to imagine). But if we agree to (*17) and (*18) as principles for inferences in discourse, we are *obliged* to accept premises like "It does and it does not rain", whenever someone wants us to accept it, provided his underlying intention is to pass to one of its contingent parts. I think it would be reasonable to restrict such an obligation to accept (and the corresponding authorisation to demand the acceptance of) a plain contradiction like $q \& \sim q$ beyond necessity. A first step was made by W. Ackermann in abolishing the general formula of "*from contradictories*" paradox (*19) a second should be made by abolishing its substitution instances (*17) and (*18), the general formulas of "*from contradictories*" simplification.

So much to the formulas (*17) and (*18) as substitution instances of the paradoxical formula (*19). The defender of (*17) and (*18) will probably prefer the view that they are (admittedly odd) substitution instances of the general formulas of simplification (*3) and (*4) and thus *principles of inference* since the general formulas are. Now first of all it is true that *if* the general formulas of simplification are principles of inference then (according to our general ideas concerning inference) indeed the formulas (*17) and (*18) are also principles of

inference (but on the other hand: if we say that we do not accept the formulas (*17) and (*18) as principles of inference then the general formulas cannot be principles either). Is there an independent reason (concerning our normal understanding of inferences) for saying that the general formulas (and hence all of their formal substitution instances) are principles of inference? I think the most tempting reason is to say that all sorts of formulas of simplification are (in a very special sense of the word) "analytical" principles, since they allow us to pass from a complex whole to each of its parts. But should it be allowed? And is it really the case that "It rains" is part of that what someone means who is apt to hold "It rains and it does not rain". I am not sure whether this is so, but even if it were, the normal reaction to someone performing such an assertion on a rainy day, in my opinion, is to wonder what he is aiming at and not to analyse this assertion in the usual polite way: "I agree with you to the extent that it rains but I cannot see how you could come to the opinion that it does not" for he, being asked, might answer: "For the very same reason: it is a simple consequence of my more general view that it rains and it does not rain." I think that the information the discourse partner gets from someone who asserts " $p \& \sim p$ ", (and also the information the latter gets about his beliefs) is in a certain respect greater than the usual information in an assertion of the form $x \& y$. The surplus of information concerns the discovery that there is something seriously wrong with the set of beliefs oneself or the other has, since it contains a contradiction. This discovery is far more important than the "content" of such a complex proposition – its analysis would rather cloud this insight, than focus it. Note that in a normal discourse the discovery of a plain contradiction will cause an interruption and close that part of discussion in order to provide some basis for a *reductio ad absurdum* argumentation, revisions, reflections etc.

If we look at these two roots, the "*from contradictories*" paradox (*19) on the one hand and the general formulas of simplification (*3) and (*4) on the other hand, the first root exhibits that something is wrong with the "*from contradictories*" formulas of simplification (*17) and (*18), since it is (like every

principle of inference with a contradiction as premise) useless in the sense that it cannot be used to pass from true premises to a true conclusion, and, for that and the other reasons indicated, is in conflict with our pretheoretic notion of inference. At first, the second root drives us in the opposite direction: it seems that these formulas, since they can be understood as special cases of the general principles of simplification, content special principles of "analysis" that are tolerable. My point is that if we arrive at $p \& \sim p$ it is not analysis that is needed, but a new beginning, and that therefore the second root does not really give a reason in favour of these principles – the conflict with our pretheoretic notion of inference remains.

This is the reason why I think it worthwhile to study systems that try to give up (*17) and (*18). But in doing so, we should be modest and should not postulate that every substitution instance of (*17) and (*18) is not a thesis, for this, probably, would lead us to very weak systems. I think it is already an advance to construct and analyse systems in which the general formulas of "*from contradictories*" simplification are not derivable²⁷.

The family of systems I want to propose are neither the only nor the first systems in which the general formulas of simplification and some of their substitution instances are rejected²⁸, but they are motivated by different reasons and thus lead to different results.

4. Minimal requirements for a system without general principles of simplification

In order to determine which properties an alternative propositional system without general principles of simplification (short: "WGPS-system"), in my opinion, should at least have, it

27 In the family of SIM1 systems I found the following substitution instances of the general "from contradictories" formulas of simplification:

$Mp \& \sim Mp \rightarrow Mp$	but not:	$Mp \& \sim Mp \rightarrow \sim Mp$
$p \& \sim Lp \rightarrow \sim Lp$	but not:	$Lp \& \sim Lp \rightarrow Lp$
$(p \& \sim p) \& \sim (p \& \sim p) \rightarrow (p \& \sim p)$	but not:	$(p \& \sim p) \& \sim (p \& \sim p) \rightarrow \sim (p \& \sim p)$

28 Systems in which the general principles of simplification are not derivable have been proposed for instance by Nelson 1930, McCall 1966, Angell 1962, Cheng 199), et al.

is essential to distinguish by definition two different relations of inclusion that may or may not hold between the classical and the alternative WGPS-system:

Definition 3: Let S be an alternative WGPS-system with conjunction, negation, implication and necessity as primitives and T be a reformulation of the PM-calculus with conjunction, negation and material implication as primitives. Let

$\phi: \text{WFF}(S) \longrightarrow \text{WFF}(T)$ be a mapping such that

$$\phi(p) = p \text{ for every atomic formula } p$$

$$\phi(\sim x) = \sim \phi(x)$$

$$\phi(x \& y) = \phi(x) \& \phi(y)$$

$$\phi(x \rightarrow y) = \phi(x) \supset \phi(y)$$

$$\phi(Lx) = \phi(x)$$

then S is a *regular system* iff for every thesis x of S $\phi(x)$ is a thesis of T ²⁹.

In a certain sense a *regular WGPS-system is included* in the *PM-calculus*, namely if the necessity operator is ignored and the implication operator is identified with material implication. In quite another sense the PM-calculus *is included* in a WGPS-system, if the PM-calculus is a *subsystem* of a WGPS-system (cf. Definition 4), namely in the sense that every formula of the PM-calculus can also be derived in the alternative system (but with an important difference concerning the law of conditionalization: a material implication thesis in a WGPS-system does not count as a principle of inference.

Definition 4: Let S be a WGPS-system with conjunction, negation, implication and necessity as primitives and let T be a reformulation of the PM-calculus with conjunction and negation as primitives. Let

$\psi: \text{WFF}(T) \longrightarrow \text{WFF}(S)$ be a mapping such that

$$\psi(p) = p \text{ for every atomic formula } p$$

29 Cf. Hughes and Cresswell 1968, 267.

$$\psi(\sim x) = \sim \psi(x)$$

$$\psi(x \& y) = \psi(x) \& \psi(y)$$

then *PM* is a subsystem of *S* (or '*S* is a supersystem of *PM*') iff for every thesis *x* of *T*, $\psi(x)$ is a thesis of *S*.

The following definition concerning weakest and strongest formulas generalises the idea expressed by the paradoxical formulas (*30) and (*31).

Definition 5: A wff *x* of a propositional system *S* is called a *weakest formula*, iff it is implied by every other term; *x* is called a *strongest formula*, iff it implies every other term.³⁰

By using these definitions we can formulate the following requirements:

(i) None of the above mentioned paradoxes should be derivable, i.e. a WGPS-system should not have the "from contradictories" formulas of simplification

$$*22 \quad p \& \sim p \rightarrow p \quad \text{and} \quad *23 \quad p \& \sim p \rightarrow \sim p$$

as theses and therefore cannot have the general formulas of simplification

$$*24 \quad (p \& q) \rightarrow p \quad *25 \quad (p \& q) \rightarrow q$$

Also, it should not have any of the paradoxes of material or strict implication,

$$*26 \quad \sim p \rightarrow (p \rightarrow q) \quad *27 \quad q \rightarrow (p \rightarrow q)$$

$$*28 \quad \sim Mp \rightarrow (p \rightarrow q) \quad *29 \quad Lq \rightarrow (p \rightarrow q)$$

nor the "from contradictories" paradox or its conversion

$$*30 \quad (p \& \sim p) \rightarrow q \quad *31 \quad p \rightarrow (qv \sim q)$$

(ii) A WGPS-system should not contain any *weakest* or *strongest* formula, for in the pretheoretic sense there is no premise from which every arbitrarily chosen proposition follows as its conclusion, and no conclusion that can be drawn from every arbitrarily chosen premise.

30 These definitions have been introduced by T. Sugihara 1955.

(iii) It should be a *regular* system, for it should not lead to inferences that must be regarded as invalid from a classical point of view, i.e. it should, concerning the notion of inferential validity, be more restrictive than, but still compatible with the classical system.

(iv) It should be a *supersystem* of the PM-calculus, for it should respect the truth functional relations between the truth-functional connectives.

(v) It should guarantee the following equivalences as theses:

- | | |
|---|---|
| *32 $p \leftrightarrow p$ | *33 $p \leftrightarrow \sim(\sim p)$ |
| *34 $(p \& q) \leftrightarrow (q \& p)$ | *35 $(p \vee q) \leftrightarrow (q \vee p)$ |
| *36 $(p \& (q \& r)) \leftrightarrow ((p \& q) \& r)$ | |
| *37 $(p \vee (q \vee r)) \leftrightarrow ((p \vee q) \vee r)$ | |
| *38 $(p \& p) \leftrightarrow p$ | *39 $(p \vee p) \leftrightarrow p$ |
| *40 $(p \rightarrow \sim q) \leftrightarrow (q \rightarrow \sim p)$ | *41 $(p \rightarrow q) \leftrightarrow (\sim q \rightarrow \sim p)$ |
| *42 $Mp \leftrightarrow \sim L \sim p$ | *43 $Lp \leftrightarrow \sim M \sim p$ |
| *44 $(p \rightarrow q) \leftrightarrow L(p \rightarrow q)$ | |

(vi) It should guarantee the following implications as theses:

- | | |
|---|------------------------|
| *45 $((p \rightarrow q) \& (q \rightarrow r)) \rightarrow (p \rightarrow r)$ | |
| *46 $((q \rightarrow r) \& (p \rightarrow q)) \rightarrow (p \rightarrow r)$ | |
| *47 $((p \rightarrow q) \& (p \rightarrow r)) \rightarrow (p \rightarrow (q \& r))$ | |
| *48 $((p \rightarrow r) \& (q \rightarrow r)) \rightarrow ((p \vee q) \rightarrow r)$ | |
| *49 $((p \rightarrow q) \& p) \rightarrow q$ | |
| *50 $((p \rightarrow q) \& \sim q) \rightarrow \sim p$ | |
| *51 $(p \rightarrow (q \rightarrow r)) \rightarrow ((p \& q) \rightarrow r)$ | |
| *52 $(p \rightarrow q) \rightarrow \sim(p \& \sim q)$ | |
| *53 $(p \rightarrow \sim p) \rightarrow \sim p$ | |
| *54 $((p \rightarrow q) \& (p \rightarrow \sim q)) \rightarrow \sim p$ | |
| *55 $p \rightarrow Mp$ | *56 $Lp \rightarrow p$ |

$$*57 \quad \sim Lp \rightarrow M\sim p$$

$$*58 \quad Mp \rightarrow \sim L\sim p$$

$$*59 \quad ((p \rightarrow q) \& \sim Mq) \rightarrow \sim Mp$$

$$*60 \quad ((p \rightarrow q) \& Lp) \rightarrow Lq$$

$$*61 \quad (p \rightarrow q) \rightarrow \sim M(p \& \sim q)$$

$$*62 \quad Mp \rightarrow \sim (p \rightarrow \sim p)$$

5. The minimal system SIM₁₀

The system SIM₁₀ is based on the following axioms³¹:

$$\text{Ax1)} \quad (p \& q) \rightarrow (q \& p)$$

$$\text{Ax2)} \quad (p \& (q \& r)) \rightarrow ((p \& q) \& r)$$

$$\text{Ax3)} \quad p \rightarrow p$$

$$\text{Ax4)} \quad (p \rightarrow q) \& (q \rightarrow r) \rightarrow (p \rightarrow r)$$

$$\text{Ax5)} \quad (p \rightarrow q) \& (p \rightarrow r) \rightarrow (p \rightarrow (q \& r))$$

$$\text{Ax6)} \quad (p \rightarrow (q \rightarrow r)) \rightarrow ((p \& q) \rightarrow r)$$

$$\text{Ax7)} \quad (p \rightarrow \sim q) \rightarrow (q \rightarrow \sim p)$$

$$\text{Ax8)} \quad (\sim p \rightarrow q) \rightarrow (\sim q \rightarrow p)$$

$$\text{Ax9)} \quad p \rightarrow Mp$$

$$\text{Ax10)} \quad (p \rightarrow q) \rightarrow \sim M(p \& \sim q)$$

$$\text{Ax11)} \quad (p \& p) \rightarrow p$$

$$\text{Ax12)} \quad M(p \& q) \rightarrow Mp \& Mq$$

$$\text{Ax13)} \quad Mp \& Mq \rightarrow Mp^{32}$$

31 It can be shown, that the axioms Ax1, Ax2, Ax4, Ax5, Ax6, Ax7, Ax8, Ax9, Ax10, Ax13, Ax14 and Ax16 (taken separately) are independent of all the other axioms and rules of SIM₁₀.

32 Though (Ax13) is not only a substitution instance of, but also has much resemblance with the general formula of simplification (and thus could be called "general formula of possibility simplification") it should be noted that there are no "from contradictories" formulas of simplification that could be obtained from it by mere substitution. For the rule of substitution allows only formulas like $Mp \& M\sim p \rightarrow Mp$ and $Mp \& M\sim p \rightarrow M\sim p$. On the other hand (Ax13) together with (Ax16) (and some other axioms) indeed leads to a paradoxical formula of simplification (cf. footnote 27).

- Ax14) $(p \rightarrow q) \rightarrow \sim M \sim (\sim q \rightarrow \sim p)$
 Ax15) $((p \rightarrow q) \& \sim Mq) \rightarrow \sim Mp$
 Ax16) $MMp \rightarrow Mp$
 Ax17) $\sim M(\sim(p \& \sim q) \& ((p \& r) \& \sim(q \& r)))$

together with the following definitions for " $\leftrightarrow$ ", " $\vee$ ", " L " and " $\supset$ " (material implication) and calculus rules:

- DEF. " $\leftrightarrow$ " : $p \leftrightarrow q := (p \rightarrow q) \& (q \rightarrow p)$
 DEF. " $\vee$ " : $p \vee q := \sim(\sim p \& \sim q)$
 DEF. " L " : $Lp := \sim M \sim p$
 DEF. " $\supset$ " : $p \supset q := \sim(p \& \sim q)$

Rule of substitution 1: If x is a thesis and y differs from x only in so far as all occurrences of a propositional variable of x are replaced in y by a certain term z , then y is a thesis.

Rule of substitution 2 ("equivalence rule"): If x is a thesis and y differs from x only in so far as some or all occurrences of a certain term z of x are replaced in y by a certain term z' , then y is a thesis, provided that $z \leftrightarrow z'$ is a thesis.

Rule of detachment for implication³³ ("modus ponens rule"): If $x \rightarrow y$ is a thesis and x is a thesis, then y is a thesis.

Rule of detachment for material implication: If $x \supset y$ is a thesis and x is a thesis, then y is a thesis³⁴.

33 The rule of detachment for implication can be introduced by means of the rule of detachment for material implication as soon as it is shown that $(p \rightarrow q) \supset (p \supset q)$ is a thesis, for instead of using the rule of detachment for implication with respect to some theses $x \rightarrow y$ and x we could proceed in the following way:

a) $(p \rightarrow q) \supset (p \supset q)$ thesis of SIM1
 b) $(x \rightarrow y) \supset (x \supset y)$ (a) X $[x/p \ y/q]$
 c) $(x \rightarrow y)$ condition for the application of the rule
 d) $(x \supset y)$ (b) (c) X detachment for material implication
 e) x condition for the application of the rule
 f) y (d) (e) X detachment for material implication

but in order to show that (a) indeed is a thesis of SIM1 the stronger rule of detachment for implication is needed.

34 This rule is (with respect to the SIM1 systems) interchangeable with the *Rule of disjunctive syllogism*: "If x and $\sim x \vee y$ are theses, then y is a thesis", which is known as "Ackermann's γ -rule", since it is used in (Ackermann 1956, 119).

Rule of conjunction: If x is a thesis and y is a thesis, then $x \& y$ is a thesis.

That $SIM1_0$ fulfils the "positive" requirements (iv), (v) and (vi) can be shown by straight forward derivations – the fulfilment of (ii) by a finite adaptation of the Sugihara matrices³⁵, for (i) and (iii) cf. the last section.

6. The system $SIM1_1$

The system $SIM1_0$ is designed for the connectives " $\&$ ", " $\sim$ ", " $\rightarrow$ " and " M ", and does not contain distribution formulas which could establish a connection between the basic operators of $SIM1_0$ and the operators " $\vee$ " and " L " defined in the usual way. To get a system with distributive properties concerning these connectives, i.e. a system, in which

- *63 $L(p \& q) \leftrightarrow Lp \& Lq$
- *64 $M(p \vee q) \leftrightarrow Mp \vee Mq$
- *65 $((p \vee r) \& (q \vee r)) \leftrightarrow ((p \& q) \vee r)$
- *66 $((p \& r) \vee (q \& r)) \leftrightarrow ((p \vee q) \& r)$

can be derived, some additional axioms must be chosen, for instance,

- Ax18) $(\sim M \sim p \& \sim M \sim q) \rightarrow \sim M \sim (p \& q)$
- Ax19) $\sim M \sim (p \& q) \rightarrow \sim M \sim p$
- Ax20) $(\sim (\sim p \& \sim r) \& \sim (\sim q \& \sim r)) \rightarrow \sim (\sim (p \& q) \& \sim r)$
- Ax21) $\sim (\sim (p \& q) \& \sim r) \rightarrow (\sim (\sim r \& \sim p) \& \sim (\sim r \& \sim q))$

which together with the axioms, definitions and rules of $SIM1_0$ constitute the system $SIM1_1$.

It should be noticed, that though (*65) and (*66) are theses in $SIM1$ the following formula, which has some similarities with distribution and is derivable in the system E) is *not* a thesis of

35 Cf. for instance Linneweber-Lammerskitten 1988, 278.

SIM1₁ (or any other member of the SIM1 family) since it does not satisfy the set of matrices SM1.0:

$$*67 \quad (p \& (q \vee r)) \rightarrow ((p \& q) \vee r)$$

Since the idea of distribution concerning conjunction and disjunction is fully expressed by the formulas (*65) and (*66), there must be something else beyond distribution hidden in (*67), but up to now I have no good idea what it is.

7. The system SIM1₂

It is widely agreed among modal logicians that our pretheoretic notions of necessity, possibility and conjunction are such, that a relation of distribution holds between necessity and conjunction, but not between possibility and conjunction, i.e. that (*68) should be a thesis, while (*69) should not,

$$*68 \quad L(p \& q) \leftrightarrow (Lp \& Lq)$$

$$*69 \quad M(p \& q) \leftrightarrow (Mp \& Mq)$$

The reason for rejecting (*69) is that for two contingent propositions p and q (for instance for a contingent proposition p and its negation $\sim p$) the corresponding possibility statements Mp and Mq are both true, while the statement $M(p \& q)$ might be false (as in the example $M(p \& \sim p)$), thus the formula

$$*70 \quad (Mp \& Mq) \rightarrow M(p \& q)$$

in most of the modal systems³⁶ is not a thesis. The same argument can be used for the rejection of (*71) – in both cases the second part of the conjunction in the antecedent is too weak.

$$*71 \quad (Mp \& q) \rightarrow M(p \& q)$$

If, however, this second part is replaced by Lq , the antecedent should be strong enough to imply $M(p \& q)$, thus by adding

36 A different position was held by Lukasiewicz in his paper *A system of modal logic*, 1952.

$$\text{Ax22) } (Mp \& \sim M \sim q) \rightarrow M(p \& q)$$

to SIM1_1 we get a system (" SIM1_2 ") in which the formula

$$*72 \quad (Mp \& Lq) \rightarrow M(p \& q)$$

is a thesis, which cannot be derived in SIM1_1 . In SIM1_2 the rule of necessitation ("If x is a thesis, Lx is a thesis") is introducible and it can be shown, that it is a normal modal system, i.e. that

- (i) every theorem of the PM-calculus (based on " $\sim$ " and " $\&$ ") is a thesis
- (ii) $L(p \supset q) \supset (Lp \supset Lq)$ is a thesis.
- (iii) it is closed under the rule of material detachment.
- (iv) it is closed under the rule of necessitation.

8. The system SIM1_3

One of the disadvantages the system SIM1_2 still has, is that only restricted versions of the following rules³⁷

- (i) *Becker's rule 1*: If $x \rightarrow y$ is a thesis then also $Lx \rightarrow Ly$ is a thesis.
- (ii) *Becker's rule 2*: If $x \rightarrow y$ is a thesis then also $Mx \rightarrow My$ is a thesis.
- (iii) *Becker's rule 3*: If $x \rightarrow y$ is a thesis then also $LMx \rightarrow LMy$ is a thesis.
- (iv) *Becker's rule 4*: If $x \rightarrow y$ is a thesis then also $MLx \rightarrow MLy$ is a thesis.

are introducible³⁸ and that it contains a lot of hidden modalities³⁹, e.g. conjunctions such as $(Lp \& p)$, which cannot be reduced. By adding the following axioms

³⁷ This rule was originally formulated by O. Becker (cf. Hughes 1996, 200).

³⁸ I.e. the rules can only be applied, if in addition to $x \rightarrow y$ also the corresponding simplification formula $x \& y \rightarrow x$ is a thesis.

³⁹ We can interpret any *first degree modal function of a single variable* i.e. a wff that contains (several occurrences of) a single variable and one or more modal operators, such that none of these modal operators is in the scope of any other, as a hidden modality. It has been shown that even a system like S4 has infinitely many distinct modal functions of a single

$$\text{Ax23)} \quad (p \rightarrow (p \& q)) \rightarrow ((p \& q) \rightarrow p)$$

$$\text{Ax24)} \quad (p \rightarrow q) \rightarrow ((p \& q) \rightarrow q)$$

to SIM1_2 we arrive at the system SIM1_3 , in which Becker's rules above are introducible. SIM1_3 has 14 different modalities, namely p , Mp , Lp , MLp , LMp , MLMp , LMLp and their negations. Most of the hidden modalities built up by conjunction (e.g. $\text{Lp} \& p$ etc.) can be reduced using Becker's rules.

9. The matrix based system MS1.0

By means of the following set of matrices SM1.0 , one can show that the requirements (i) and (iii) are fulfilled⁴⁰ by SIM1_0 :

K	*	1	2	3	4	C	2	4	4	4	N	4	M	1
	*	2	2	3	4		1	2	4	4		3		1
		3	3	3	4		2	4	2	4		2		3
		4	4	4	4		1	2	1	2		1		4

(designated values 1 and 2):

But this set of matrices SM1.0 can also be useful in another respect: it can be regarded as defining a limit case for the construction of a family of systems, where each member is an extension of the SIM1_0 system satisfying all of the WGPS-requirements and allowing some additional plausible inferences. For the set of matrices SM1.0 together with appropriate definitions for the matrices of equivalence, disjunction, necessity etc. determines a matrix based system MS1.0 in the

variable (cf. Hughes and Cresswell 1968, 57). With respect to SIM1 conjunctive formulas are, of course, of main interest and I confine myself to these.

40 One has to show a) that the axioms and rules of SIM1_0 satisfy this set of matrices, i.e. that the axioms take only designated values and that the application of the calculus rules preserves this property, b) that the paradoxical formulas under (i) take at least one undesignated value, c) that SIM1_0 satisfies the 2-valued characteristic set of matrices of the PM-calculus extended by identity matrices for the modal notions and matrices for SIM1_0 implication (equivalence) that are congruent with that for material implication (equivalence). SIM1_0 satisfies the 2-valued characteristic set of matrices of the PM-calculus extended by identity matrices for the modal notions and matrices for SIM1_0 implication (equivalence) that are congruent with that for material implication (equivalence). This can be seen at once by replacing the designated values 1 and 2 in the set of matrices above by the truth-value t and the undesignated values 3 and 4 by the truth-value f .

canonical way: a wff x is a thesis of MS1.0, iff x takes a designated value under SM1.0. The matrix based system MS1.0 is a limit case for a sequence (or in general for a net) of systems satisfying the SM1.0 matrices and containing SIM1₀.

As MS1.0 is stronger than each member of the SIM1 family and has the advantage of being decidable, it may appear to be preferable. But though it satisfies all of the positive (iv)-(vi) and the negative requirements (i) and (iii), it contains some weakest (e.g. $M(p \rightarrow p)$) and some strongest formulas (e.g. $\sim M(p \rightarrow p)$). Like all the systems based on a *finite valued* set of matrices MS1.0 has the additional disadvantage that it allows unwanted equivalences. Thus for instance

$$*73 \quad (p \rightarrow p) \leftrightarrow ((q \& (r \& s)) \rightarrow ((q \& r) \& s))$$

is a thesis, though its constituents $(p \rightarrow p)$ and $((q \& (r \& s)) \rightarrow ((q \& r) \& s))$ have a different number of (different) propositional variables.

Institut für Philosophie der Universität Bern
Unitobler
Länggass Strasse 49
CH 3009 BERN

Bibliography

- ACKERMANN W. (1956). Begründung einer strengen Implikation. *The Journal of Symbolic Logic* 21, 113-128.
- ANDERSON A. R., BELNAP N. D., JR. (1990). *Entailment. The Logic of Relevance and Necessity I*. Second printing, with corrections, Princeton: Princeton University Press [First printing 1975].
- ANDERSON A. R., BELNAP N. D., JR. & DUNN J. M. (1992). *Entailment. The Logic of Relevance and Necessity II*. Princeton: Princeton University Press.
- ANGELL R. B. (1962). A propositional logic with subjunctive conditionals. *The Journal of Symbolic Logic* 27, 327-343.
- BURIDANUS I. B. (1976). De Consequentis. In: H. Hubien (ed.), *Iohannis Buridani Tractatus De Consequentis*. Louvain.
- CHENG J. (1991). Logical tool of knowledge engineering: Using entailment logic rather than mathematical. *Proceedings of the ACM Nineteenth Annual Computer Science Conference*. San Antonio, U.S.A., 228-238.
- HAACK S. (1996). *Deviant Logic, Fuzzy Logic. Beyond the Formalism*. Chicago & London: The University of Chicago Press.
- HUGHES G.E., CRESSWELL M.J. (1990). *An Introduction to Modal Logic*. London: Routledge [Reprint with corrections of the first edition London 1968].
- HUGHES G.E., CRESSWELL M.J. (1996). *A New Introduction to Modal Logic*. London: Routledge.
- LEWIS C.I., LANGFORD C.H. (1959). *Symbolic Logic*. New York: Dover Publications, Inc. [Corrected second edition of *Symbolic Logic*, New York (The Century Co.) 1932].
- LINNEWEBER-LAMMERSKITTEN H. (1988). *Untersuchungen zur Theorie des hypothetischen Urteils*. Münster: Nodus.
- LINNEWEBER-LAMMERSKITTEN H. (1994). A survey of the derivability of important implicative principles in alternative systems of propositional logic. In: G. Meggle, U. Wessels (eds.), *Analyomen 1. Proceedings of the 1st Conference*

- "*Perspectives in Analytical Philosophy*". Berlin & New York: W. de Gruyter, 76-87.
- LINNEWEBER-LAMMERSKITTEN H. (1997). Interfaces of logic. In: P. Weingartner, G. Schurz & G. Dorn (eds.), *The Role of Pragmatics in Contemporary Philosophy: Contributions of the Austrian Ludwig Wittgenstein Society*, Vol. 5, Kirchberg 1997, 549-555.
- LUKASIEWICZ J. (1952). A system of modal logic. *The Journal of Computing Systems* 1, 1952, 111-149 [Reprinted in: Lukasiewicz J., *Selected works*, (ed. L. Borkowski) Amsterdam: North-Holland Pub. Co, 1970, 352-390].
- MCCALL S. (1966). Connexive implication. *The Journal of Symbolic Logic* 31, 415-433.
- NELSON E. J. (1930). Intensional relations. *Mind* 39, 440-453.
- PAUL OF VENICE (PAULUS VENETUS) (1988) *Pauli Veneti Logica Magna. Secunda pars, Tractatus De Obligationibus*, edited with an English Translation and Notes by E.J. Ashworth. Oxford: Oxford University Press.
- READ S. (1988). *Relevant Logic*. Oxford: Basil Blackwell.
- ROUTLEY R., PLUMWOOD V., MEYER R.K. & BRADY R.T. (1982). *Relevant Logics and their Rivals I*. Ohio: Ridgeview.
- SUGIHARA T. (1955). Strict implication free from implicational paradoxes. In: *Memoirs of the Faculty of Liberal Arts*. Fukui University. Series I, Humanities and social sciences, 5, 55-59.
- WHITEHEAD A. N. & RUSSELL B. (1927). *Principia Mathematica*. 2nd ed. Cambridge & New York: Cambridge University Press [First edition 1910].

SPÉCIFICITÉ ET POTENTIALITÉS DE LA LOGIQUE FLOUE

Bernadette BOUCHON-MEUNIER

1. Pourquoi la logique floue?

Le terme «logique floue» a deux acceptions: la première, la plus stricte, la définit comme une extension de la logique classique destinée à raisonner sur des connaissances imparfaites. La seconde, très répandue actuellement, fait référence à tous les développements concernant les théories des sous-ensembles flous et des possibilités. Nous nous intéressons ici à la première acception.

La théorie des sous-ensembles flous a été introduite en 1965 par Lotfi A. Zadeh, professeur à l'Université de Californie à Berkeley, internationalement reconnu pour ses travaux en automatique et théorie des systèmes, qui a éprouvé le besoin de formaliser la représentation et le traitement de connaissances imprécises ou approximatives, afin de pouvoir traiter des systèmes d'une grande complexité dans lesquels sont, par exemple, présents des facteurs humains. Elle a pour but d'assouplir la théorie des ensembles classique, en autorisant un élément d'un univers à appartenir partiellement à une classe donnée de cet univers.

Le concept de sous-ensemble flou a été introduit pour éviter les passages brusques d'une classe à une autre (de la classe «bon» à la classe «mauvais» par exemple) et autoriser des éléments à n'appartenir complètement ni à l'une ni à l'autre (à être «médiocres», par exemple), ou encore à appartenir partiellement à chacune (avec un fort degré à la classe «bon» et un faible degré à la classe «mauvais»). La définition d'un sous-ensemble flou répond au besoin de représenter des connaissances impré-

cises (une note d'«environ 15/20» ou un élève «relativement bon»). Le caractère graduel des sous-ensembles flous correspond à l'idée que, plus on se rapproche de la caractérisation typique d'une classe, plus l'appartenance à cette classe est forte (il est sûr qu'avec 19/20, un élève est bon, avec 1/20, il est mauvais, mais à partir de quelle note passe-t-il d'une classe à l'autre et un centième de point d'écart entre deux étudiants peut-il les faire appartenir à deux classes différentes?). Rappelons qu'un sous-ensemble flou A de l'univers X est défini par une fonction d'appartenance $f_A: X \rightarrow [0, 1]$. Si celle-ci prend la valeur 1 en au moins un point de X , elle est dite normalisée.

Ainsi, la théorie des sous-ensembles flous conduit-elle à une méthode de représentation des connaissances imparfaites sur un système ou une situation, c'est-à-dire des connaissances soumises à des imprécisions, à l'utilisation de critères qualitatifs, approximatifs, reflétant une absence de rigueur des observateurs ou une flexibilité inhérente au système.

Au delà de la représentation des connaissances, il est nécessaire de mettre au point des méthodes d'exploitation de ces connaissances. Certaines méthodes peuvent être directement déduites de la théorie des sous-ensembles flous. C'est le cas, par exemple, de méthodes de classification, de traitement d'images, de gestion de requêtes imprécises adressées à des bases de données, de calcul sur des quantités floues... Cependant, bien des méthodes d'exploitation de connaissances nécessitent un raisonnement, l'établissement d'une conclusion à partir de données décrivant une situation. C'est le cas, par exemple, pour les systèmes d'aide au diagnostic, les systèmes experts. Reasonner sur des connaissances imparfaites a été l'un des objectifs affichés très tôt par L.A. Zadeh, avec la volonté de proposer un mode de raisonnement qui soit proche du raisonnement humain, puisque les humains sont aptes à fournir des conclusions précises et pertinentes à partir de descriptions approximatifs, mal connues, imprécises.

Pour raisonner à partir de connaissances imprécises, il faut aussi être capable de traiter des incertitudes exprimant des doutes non probabilistes, à caractère subjectif, ce que ne permet pas la théorie des sous-ensembles flous. Imaginons la règle pré-

cise «si l'élève a au moins 10, il est admis», et cherchons une conclusion sur l'admission d'un élève qui a «à peu près 10» (connaissance imprécise), nous ne pouvons faire mieux qu'énoncer un diagnostic incertain tel que «l'élève est peut-être admis». L.A. Zadeh a donc introduit en 1978 la théorie de possibilités qui complète la théorie des sous-ensembles flous, en proposant une représentation de connaissances incertaines compatible avec celle-ci afin de parvenir à un raisonnement approximatif. Notons qu'imprécision et incertitude sont étroitement liées, «à peu près 10» (imprécis) entraînant par exemple «peut-être plus de 10, peut-être moins de 10» (incertain), et que la compatibilité entre représentation des imprécisions (sous-ensembles flous) et représentation des incertitudes (possibilités) est donc nécessaire.

La théorie des possibilités introduit une double pondération de tout sous-ensemble d'un univers de définition X d'une variable (par exemple la note d'un élève, définie sur $[0, 20]$) par une mesure de possibilité Π et une mesure de nécessité N indiquant respectivement à quel point il est possible que la valeur (inconnue) de la variable appartienne à ce sous-ensemble et à quel point ceci est certain. On démontre que, si la connaissance sur la valeur de la variable est imprécise, et donc représentée par un sous-ensemble flou A de X , alors $f_A(x)$ indique à quel point il est possible que x soit la vraie valeur de la variable, étant donné la connaissance de A supposé normalisé, c'est-à-dire que $\Pi(\{x\})$ vaut alors $f_A(x)$. C'est ce passage entre flou et possibilités qui établit la compatibilité entre les deux théories.

Théories des sous-ensembles flous et des possibilités sont les deux piliers sur lesquels repose le raisonnement approximatif en logique floue.

2. Comment est introduite la logique floue?

La logique floue a été introduite par Goguen (1969) à partir de la définition d'une grammaire formelle de concepts inexacts, qu'il construit en s'appuyant sur la notion de sous-ensemble flou. Elle a ensuite été présentée par L.A. Zadeh (1983) comme

un moyen de raisonner «à la manière des humains», donc en utilisant des connaissances admettant l'imprécision inhérente à la plupart des éléments d'information que les humains manipulent pour raisonner et c'est cette approche qui prévaut actuellement.

Les besoins sous-jacents à l'introduction de logique floue sont liés à la nécessité de raisonner sur des connaissances imparfaites.

– Puisqu'on s'autorise à traiter de l'imprécision, il est nécessaire de manipuler des connaissances vagues, imprécises, approximatives, exprimées avec des mots du langage naturel; remarquons néanmoins que la logique floue n'a pas la prétention de traiter toute la richesse du langage naturel, mais qu'elle prend en compte des granules de connaissances issus d'expressions en langage naturel.

– Il est également nécessaire d'accepter que la vérité soit elle-même modulable: une proposition peut être tout à fait vraie, fausse, relativement vraie, ...

– Le raisonnement humain est capable d'éviter les ruptures brusques entre les situations strictement compatibles avec une règle de déduction et les autres situations: par exemple, un service de recrutement qui exige «15/20 en économie» à un examen acceptera vraisemblablement un candidat avec un excellent dossier qui aura seulement 14,98. De même, un système de raisonnement en logique floue doit exploiter des connaissances «voisines» de celles attendues en prémisse d'une règle de déduction, pour fournir une conclusion éventuellement soumise à une incertitude ou à une faible spécificité si la relation de «voisinage» n'est pas très forte, mais pour fournir une conclusion tout de même, ce que ne ferait pas un système en logique ordinaire.

– Les quantificateurs universel et existentiel sont également insuffisants pour représenter les quantificateurs que l'humain manipule. Par exemple, des expressions comme «dans presque tous les cas» ou «dans la majorité des cas» sont très naturelles et doivent pouvoir être prises en compte dans un système de raisonnement en logique floue.

Ces raisons, parmi les principales pour montrer l'impossibilité de se servir de la logique classique afin de raisonner sur des connaissances imparfaites, conduisent à la logique floue, extension de la logique classique. Des logiques multivaluées ont été introduites et largement développées, à partir des recherches de Lukasiewicz en 1928, pour assouplir la notion de degré de vérité. Mais ces travaux ne s'intéressent pas à l'expression et à la représentation des connaissances imprécises, ils n'ont pas pour vocation de traiter des termes issus du langage naturel et, s'ils ne sont pas étrangers à la logique floue, ils ne sont pas suffisants pour répondre à tous les besoins exprimés plus haut.

3. Qu'est-ce que le raisonnement en logique floue?

Les propositions sont des propositions floues définies à partir d'un ensemble L de variables linguistiques (V, X, T_V), où V est une variable (la qualité, l'âge...), X son univers de définition, T_V un ensemble de caractérisations linguistiques représentées par des sous-ensembles flous de X , et d'un ensemble M de modificateurs linguistiques (transformations de fonctions d'appartenance représentant des adverbes tels que «très», «relativement»...). Reprenant l'exemple précédent, on pourra avoir pour le résultat à un examen V , l'univers $X = [0, 20]$, et choisir $T_V = \{\text{mauvais, médiocre, bon}\}$, ou un ensemble plus finement décrit en ajoutant «nul» et «excellent». On impose généralement que, pour tout x de X , il existe une caractérisation A de T_V telle que $f_A(x) \neq 0$.

Les propositions floues élémentaires sont de la forme « V est A », où A est dans T_V ou construit à partir d'un élément de T_V à l'aide d'un modificateur (par exemple $A = \text{«très } B\text{»}$, avec B dans T_V), ou bien A est une valeur précise de X .

La valeur de vérité des propositions appartient à tout l'intervalle $[0, 1]$ et elle est fournie par la fonction d'appartenance de la caractérisation floue utilisée dans la proposition floue. Une valeur de vérité égale à 1 (respectivement à 0) correspond à une proposition absolument vraie (respectivement absolument fausse).

Des quantificateurs flous peuvent être utilisés, qui sont des sous-ensembles flous de l'univers $[0, 1]$.

Dans le cas particulier où toutes les propositions floues sont booléennes, c'est-à-dire qu'elles sont soit absolument vraies, soit absolument fausses, la logique floue est identique à la logique classique. Les sous-ensembles classiques étant des cas particuliers de sous-ensembles flous, cette identité est naturelle: des connaissances précises, correspondant au fait que la variable V prend la valeur précise x_0 de l'univers X (par exemple «la note est de 11/20», ou encore «la décision est d'accepter»), sont associées à des sous-ensembles flous particuliers de X , de fonction d'appartenance égale à 1 en x_0 et à 0 ailleurs.

Des opérateurs logiques sont définis à partir de normes triangulaires pour la conjonction et de conormes triangulaires pour la disjonction (Schweizer, Sklar 1963). Par exemple, la valeur de vérité associée à la conjonction de propositions « V est A et W est B » est $T(f_A(x), f_B(y))$, en tout x de l'univers X de V et y de l'univers Y de W , pour une norme triangulaire T telle que le minimum. De même la valeur de vérité associée à la disjonction de propositions « V est A ou W est B » est $\perp(f_A(x), f_B(y))$, en tout x de X et y de Y , pour une conorme triangulaire $\perp$ telle que le maximum. La valeur de vérité de la négation de « V est A » est définie à partir d'un opérateur monotone décroissant de $f_A(x)$, généralement $1 - f_A(x)$. Opérateurs de conjonction et de disjonction sont choisis en liaison l'un avec l'autre pour préserver le maximum des propriétés habituelles en logique classique. C'est le cas par exemple pour min et max, pour $T(a, b) = ab$ et $\perp(a, b) = a + b - ab$, pour $T(a, b) = \max(a + b - 1, 0)$ et $\perp(a, b) = \min(a + b, 1)$, ces paires vérifiant les lois de De Morgan pour la négation définie plus haut.

L'implication floue entre deux propositions floues élémentaires « V est A » et « W est B » est une proposition floue concernant le couple de variables (V, W) , dont la valeur de vérité est donnée par la fonction d'appartenance f_R d'une relation floue R entre X et Y définie, pour tout (x, y) de $X \times Y$ par:

$$f_R(x, y) = \Phi(f_A(x), f_B(y)),$$

pour une fonction Φ de $[0, 1] \times [0, 1] \rightarrow [0, 1]$. De nombreuses formes de Φ ont été introduites, à partir d'extensions de l'implication en logique classique. Citons par exemple

$$f_R(x, y) = \max(1 - f_A(x), f_B(y)),$$

qui représente $\neg V$ est $A \vee W$ est B , ou encore l'implication floue de Lukasiewicz

$$f_R(x, y) = \min(1 - f_A(x) + f_B(y), 1).$$

L'intérêt d'utiliser de telles implications apparaît dans la définition du modus ponens généralisé. Considérons une règle floue de la forme «si V est A alors W est B » et une proposition floue « V est A' », avec A' plus ou moins différent de A . Une extension du modus ponens classique doit préserver celui-ci, c'est-à-dire correspondre à l'obtention d'une conclusion « W est B' » telle que B' soit identique à B dès que A' est identique à A . On définit B' par sa fonction d'appartenance:

$$\forall y \in Y \quad f_{B'}(y) = \sup_{x \in X} F(f_A(x), f_R(x, y)),$$

pour une norme triangulaire F compatible avec l'implication floue R relativement à la préservation du modus ponens. On peut ainsi utiliser $F(a, b) = \max(a + b - 1, 0)$ avec les deux implications floues données en exemple.

Ainsi, avec une règle floue telle que «si le résultat de l'examen est bon, alors la candidature est acceptable», et une proposition floue «le résultat de l'examen est relativement bon», où «relativement» est représenté par un modificateur linguistique, le modus ponens généralisé fournira une conclusion «la candidature est B' », où B' pourra s'interpréter par «à peu près acceptable», forme moins spécifique que «acceptable», ou encore «peut-être acceptable, mais ce n'est pas absolument certain», selon l'implication floue utilisée.

Le schéma du modus ponens généralisé est encore valable lorsque prémisses et conclusion de la règle floue sont la conjonction ou la disjonction de propositions floues élémentaires.

Un tel schéma est le plus répandu dans les systèmes de déduction dans le cadre d'une représentation des connaissances par des ensembles flous.

4. Quelles sont les autres possibilités de raisonner dans un environnement flou?

Dans le cas où l'on ne manipule que des incertitudes non probabilistes, indiquant une absence de fiabilité dans la source d'information, on peut se placer dans un cadre possibiliste sans utiliser de sous-ensemble flou. La logique possibiliste, reposant sur la théorie des possibilités, permet alors de raisonner sur de telles connaissances, en utilisant les degrés avec lesquels il est possible et il est certain que les différentes propositions soient vraies. (Dubois, Prade 1984; 1987). Des liens ont été établis entre logique possibiliste et logique modale et la logique possibiliste a des développements en démonstration de théorèmes.

Le cadre de la théorie des sous-ensembles flous et de la théorie des possibilités permet de généraliser à des connaissances imparfaites d'autres formes de raisonnement, telles que:

- le raisonnement par défaut, pour lequel il est possible d'avoir une approche possibiliste, par exemple (Yager 1987),
- le raisonnement temporel, avec une approche basée sur des intervalles flous du temps (Dubois, Prade 1989),
- le raisonnement syllogistique (Zadeh 1985) où l'on combine des règles floues pondérées par des quantificateurs flous, tels que «pour Q_1 des individus, si V est A alors W est B », «pour Q_2 des individus, si W est B alors S est C », pour quel Q peut-on dire que «pour Q des individus, si V est A alors S est C ?», avec par exemple des quantificateurs flous tels que Q_1 = «une majorité», Q_2 = «la presque totalité»...

D'autres types de raisonnement, que l'on rencontre en intelligence artificielle par exemple, peuvent être étendus à des connaissances imparfaites. C'est le cas des méthodes de généralisation, par exemple, qui sont issues du fonctionnement naturel des êtres humains pour obtenir une conclusion sur une situation donnée, étant donné les connaissances déjà acquises et l'expérience passée. On suppose donné un ensemble d'objets ou de

situations connus décrits à l'aide d'attributs (un ensemble de candidats, décrits par leur âge, leur note à l'examen, le nombre de langues parlées...) et affectés d'une classe ou d'une décision (candidature acceptable, candidature inacceptable, par exemple) dans le passé. Cet ensemble constitue une base d'apprentissage, dont on se sert pour identifier la classe ou la décision à prendre en présence d'un nouveau cas (nouveau candidat, par exemple, dont on connaît les valeurs des attributs.

Les notions de ressemblance, d'analogie, de similitude entre situations, sous-jacentes à la généralisation des résultats connus sur la base d'apprentissage à tout nouveau cas, conduisent à penser qu'une représentation floue des connaissances peut être intéressante. De plus, lorsqu'on a à faire avec des connaissances numériques, il est toujours délicat de faire une différence entre deux situations dont les descriptions sont «voisines» et d'établir des seuils brutaux au-delà desquels une décision à prendre doit changer. La difficulté est encore accrue lorsque toutes les situations ne sont pas décrites avec des valeurs d'attribut du même type, par exemple lorsque certaines notes sont indiquées comme «bonnes» et d'autres avec des valeurs précises. Pour toutes ces raisons, on est amené à utiliser une représentation des connaissances floues dans différents schémas de généralisation.

Le premier mode de raisonnement par généralisation, le raisonnement inductif, ou apprentissage à partir d'exemples, recherche les attributs les plus significatifs relativement à la prise de décision souhaitée et les ordonne par ordre d'importance décroissante dans un arbre de décision. Chaque nœud de l'arbre est ainsi associé à un test sur un attribut, chaque arc issu de ce nœud étant associé à une valeur du test. Par exemple, pour un nœud associé à la note à l'examen, chaque branche pourra être associée respectivement à «approximativement supérieure à 14», «approximativement inférieure à 14», représentés par des sous-ensembles flous de $[0, 20]$. Les sommets terminaux de l'arbre correspondent à une décision identifiée. Pour chaque nouveau cas, on compare la valeur de l'attribut aux nœuds successifs de l'arbre avec celle indiquée dans le test. Le résultat est un degré de compatibilité entre 0 et 1, propagé jusqu'aux sommets terminaux de façon à déterminer avec quel degré de satis-

fiabilité une décision peut être prise (Bouchon-Meunier, Marsala, Ramdani 1996).

Le deuxième de ces modes est le raisonnement à partir de cas ou le raisonnement par analogie, proches l'un de l'autre. On cherche dans la base d'apprentissage la situation qui ressemble le plus à la nouvelle situation à traiter et l'on propose une décision à partir de la décision qui a été prise pour cette situation antérieure. Pour déterminer à quel point deux situations se ressemblent, on utilise des relations floues de ressemblance (Bouchon-Meunier, Valverde 1993) ou des mesures de comparaison de valeurs floues (Bouchon-Meunier, Rifqi, Bothorel 1996).

Un autre mode de généralisation est le raisonnement à partir de prototypes. On cherche dans la base d'apprentissage des familles de situations se ressemblant, on en détermine un prototype et la nouvelle situation sera rapprochée du prototype auquel elle ressemble le plus. Ici encore, la comparaison de valeurs floues d'attributs s'impose et la tâche de recherche de prototype peut être envisagée dans un environnement flou, une candidature retenue étant par exemple caractérisée par «plutôt jeune, avec une bonne note générale, et une excellente note en économie» (Rifqi 1996).

5. Conclusion

Il convient d'utiliser la logique floue lorsque des imperfections entachent la connaissance dont nous disposons pour raisonner. Elle est aussi intéressante toutes les fois que des informations numériques et des informations qualitatives, symboliques, linguistiques, doivent être traitées simultanément. C'est même le seul cadre formel dans lequel un tel traitement est possible. C'est également le seul cadre dans lequel puissent être traitées des imprécisions et des incertitudes, et il autorise également le traitement d'incomplétudes dans les connaissances.

Raisonner en logique floue peut avoir plusieurs facettes. La plus communément répandue est la déduction par *modus ponens* généralisé dans une extension de la logique classique. Mais le

fait de rechercher un raisonnement proche du raisonnement humain conduit à d'autres facettes, dans lesquelles l'imprécis, l'approximatif, sont pris en compte au même titre que la ressemblance, l'analogie. La richesse de la théorie des ensembles flous, d'une part, de la théorie des possibilités, d'autre part, conduit à bien des formes de raisonnement, dont toutes n'ont pas encore été explorées profondément, et dont nous avons tenté de donner un panorama aussi général que possible, bien que non exhaustif.

CNRS - LAFORIA
Université Paris VI - Boîte 169
4, place Jussieu -
F - 75252 PARIS CEDEX 05

Références

- BOUCHON-MEUNIER B. (1993). *La logique floue*. Paris: P.U.F., Que Sais-je? n° 2702, 2e éd. 1994.
- BOUCHON-MEUNIER B. (1995). *La logique floue et ses applications*. Paris: Addison Wesley.
- BOUCHON-MEUNIER B., RIFQI M., BOTHOREL S. (1996) Towards general measures of comparison of objects. *Fuzzy Sets and Systems* 84. 2, 143-153.
- BOUCHON-MEUNIER B., MARSALA C., RAMDANI M. (1993). Inductive learning and fuzziness. *Scientia Iranica* 2.4, 289-298, 1996.
- BOUCHON-MEUNIER B., VALVERDE L. (1993). Analogical reasoning and fuzzy resemblance. In: B. Bouchon, L. Valverde, R.R. Yager (eds.), *Intelligent Systems with Uncertainty*. Amsterdam: Elsevier.
- DUBOIS D., PRADÉ H. (1984). The management of uncertainty in expert systems: the possibilistic approach. *Proc. 10th Triennial IFORS Conference*. Dordrecht: North Holland, 949-964.

- DUBOIS D., PRADE H. (1987). *Théorie des possibilités, applications à la représentation des connaissances en informatique*. Paris: Masson, 2e éd.
- DUBOIS D., PRADE H. (1989). Processing fuzzy temporal knowledge. *IEEE Trans. on Systems, Man and Cybernetics* 19, 729-744.
- GOGUEN J.A. (1969). The logic of inexact concepts, *Synthese* 19, 325-373.
- RIFIQI M. (1996). *Mesures de comparaison, typicalité et classification d'objets flous: théorie et pratique*. Thèse de l'Université Paris VI.
- SCHWEIZER B., SKLAR A. (1963). Associative functions and abstract semigroups. *Publications Mathematicae Debrecen* 10, 69-81.
- YAGER R.R. (1987). Using approximate reasoning to represent default knowledge. *Artificial Intelligence* 31, 99-112.
- YAGER R.R., OVCHINNIKOV S., TONG R.M., NGUYEN H.T. (1987). *Fuzzy Sets and Applications, Selected Papers by L.A. Zadeh*. London: John Wiley & Sons.
- ZADEH L.A. (1965). Fuzzy sets. *Information and Control* 8, 338-353.
- ZADEH L.A. (1978). Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets and Systems* 1, 2-28
- ZADEH L.A. (1983) The role of fuzzy logic in the management of uncertainty in expert systems. *Fuzzy Sets and Systems* 11, 199-227.
- ZADEH L.A. (1985). Syllogistic reasoning in fuzzy logic and its application to usuality and reasoning with dispositions, *IEEE Trans. Systems, Man and Cybernetics* SMC 15, 754-763.

LOGIQUES NON MONOTONES

Gérard CHAZAL

Introduction

Définition de la monotonie

Dans le cadre de la logique classique, le calcul des prédicats du premier ordre par exemple, si l'on se donne un ensemble D de formules cohérentes entre elles, on peut, en appliquant les règles d'inférence habituelles, obtenir un nouvel ensemble de formules que l'on notera $\text{Th}(\Delta)$ avec $\Delta \subseteq \text{Th}(\Delta)$. Si l'on ajoute une formule à Δ cohérente avec l'ensemble des formules de Δ , de manière à obtenir un ensemble Δ' toujours cohérent, $\text{Th}(\Delta')$ comportera au moins autant de formules que $\text{Th}(\Delta)$.

Cette propriété de la logique classique est appelée la monotonie. On pourrait la formuler plus simplement de la manière suivante: *si l'on ajoute à un ensemble de prémisses une nouvelle prémisses de telle sorte que la cohérence soit maintenue, le nombre de conclusions ne peut que croître.*

On peut donner une définition plus formelle de cette propriété:

Si $A_1, A_2, \dots, A_n \vdash B$ alors $A_1, A_2, \dots, A_n, A_{n+1} \vdash B$

Problèmes conduisant à adopter des logiques non monotones

1 - raisonner en absence d'information complète

Cette propriété de la logique classique qui, dans le cadre de la déduction, la rend très robuste, présente cependant quelques inconvénients lorsqu'il s'agit de modéliser certaines formes de raisonnement, soit que l'on veuille rendre compte de comporte-

ments humains, soit que l'on souhaite les reproduire avec une machine.

Prenons un exemple simple. Pour venir de Dijon à Neuchâtel, deux routes s'offraient à moi et évidemment je souhaitais prendre l'itinéraire le plus rapide. L'un de ces itinéraires emprunte des routes nationales et départementales, traverse beaucoup de villages, l'autre se fait par autoroute. Les connaissances que j'ai des routes, des autoroutes, des limitations de vitesse en vigueur, etc. me permettent de conclure que l'itinéraire par l'autoroute est plus rapide. Cependant, si j'apprends avant de partir qu'il y a d'importants travaux sur l'autoroute ou un grave accident provoquant d'énormes bouchons, je conclurai que l'itinéraire par les routes nationales et départementales est plus rapide. On peut considérer le problème de deux manières. Premièrement je peux, avant de partir me renseigner sur l'état des routes et autoroutes que je suis susceptible d'emprunter. Dans le cas où j'apprendrais l'existence de travaux ou d'un accident, je serai obligé de retirer ma première conclusion «le trajet autoroutier est le plus rapide» pour la remplacer par la conclusion «le trajet par la route est plus rapide». Nous sommes dans un cas typique de non-monotonie puisque l'ajout d'une prémisse (concernant l'état des routes) m'oblige à retirer une conclusion. Deuxièmement, je peux ne pas avoir les moyens de connaître l'état des routes. Dans ce cas j'en resterai à ma première conclusion et, bien que devant raisonner avec un état de connaissance incomplet, je partirai, quitte à réviser ma décision si cet état de connaissance est modifié. Le raisonnement non monotone va intervenir chaque fois que l'on devra raisonner ou faire «raisonner» une machine avec une base de connaissances incomplète.

De telles situations sont nombreuses dans la vie. Nous tenons fréquemment des raisonnements et aboutissons à des conclusions qui nous servent pour agir alors que nos connaissances sont insuffisantes pour obtenir ces conclusions de manière ferme et définitive avec les règles d'inférence habituelles de la logique classique. On rencontre aussi ce problème dans l'élaboration de la succession des tâches dans un processus industriel, ce qu'on appelle parfois les gammes de montage. La gamme la plus logi-

quement déduite peut être remise en cause par de nouvelles informations (état des stocks, événements sociaux, rupture chez un fournisseur, etc.).

2 - utiliser des règles avec exceptions

Les systèmes experts mis au point par les techniques d'intelligence artificielle utilisent le plus souvent le *modus ponens* avec des règles de la forme *Si A alors B, or A, donc B*. L'élaboration d'un système expert suppose donc que l'on traduise les connaissances de l'expert auquel on veut substituer un ordinateur sous forme de règles conditionnelles. Or, lorsqu'il s'agit de traduire ainsi des connaissances empiriques, on s'aperçoit que beaucoup de règles, comme dans la grammaire française par exemple, comportent des exceptions. Si les exceptions sont peu nombreuses, on peut pallier au problème de manière simple en ajoutant des conditions à la règle. Par exemple la règle *si x est un nom singulier et si x se termine par «ou» alors au pluriel on écrira $x+s$* deviendra *si x est un nom singulier et si x se termine par «ou» et si($x<>$ «caillou» et $x<>$ «chou», et....) alors au pluriel on écrira $x+s$* . Cependant une telle technique devient inutilisable si le nombre d'exceptions est très grand (lourdeur et longueur des règles) ou si les exceptions ne sont pas connues à l'avance. Dans ce cas là, il sera préférable d'utiliser le formalisme d'une logique non monotone.

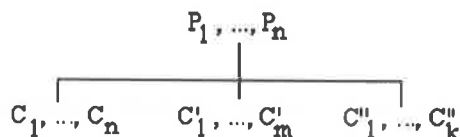
Nous serons donc amenés à utiliser une logique non monotone chaque fois que nous voudrions formaliser un raisonnement conjectural ou devant s'effectuer avec des informations incomplètes, incertaines, évolutives.

1. Caractères généraux des logiques non monotones

La pluriextensionnalité

Si nous considérons un ensemble Δ de connaissances sur les caractères des routes et des autoroutes reliant Dijon à Neuchâtel nous pourrions déduire de cet ensemble deux conclusions incompatibles: *prendre les routes nationales et départementales* et

prendre l'autoroute. Le choix entre l'une et l'autre se fera en fonction de l'ajout à Δ d'une information concernant l'état des routes. D'une manière générale, dans le cadre d'une logique non monotone, à partir d'un ensemble de prémisses données, plusieurs ensembles incompatibles de conclusions possibles peuvent être obtenues.



Cela entraîne un certain nombre de conséquences:

- Pour une formule A , la non-monotonie demande la définition d'une conséquence sémantique permettant d'obtenir des conclusions qui ne sont pas vérifiées pour tous les modèles de A (en prenant «modèle» au sens habituel de la sémantique).
- Des conclusions pourront être consistantes avec les prémisses mais non valides au sens que la logique classique donne à ce terme. D'où une distinction entre la validité au sens classique du terme (une formule est valide si elle se déduit des axiomes et des prémisses) et la simple consistance.
- L'ensemble des conclusions inférées, parmi les ensembles possibles dépendra le plus souvent de l'ordre dans lequel seront effectuées les inférences. (Cf. L'exemple ci-dessous).
- Enfin, il faudra évidemment éviter d'inférer conjointement des conclusions non consistantes entre elles.

Remarque. Pour éviter l'existence de plusieurs ensembles incompatibles de conclusions possibles, on peut utiliser la monotonie restreinte avec un système logique classique. Cette monotonie restreinte correspond à la règle d'inférence suivante:

Si $(A_1, \dots, A_n) \vdash B$ et
 si $(A_1, \dots, A_n) \vdash C$, alors $(A_1, \dots, A_n, B) \vdash C$

Cependant cela s'avère le plus souvent insuffisant pour formaliser des formes de raisonnement du type de ceux que nous venons de voir.

Quelques définitions.

Un certain nombre de notions de la logique classique devront prendre une définition un peu différente de celle que l'on a coutume d'en donner.

Tout d'abord on définira la notion de point fixe de la manière suivante:

Un point fixe est un ensemble stable de formules (de connaissances ou de croyances) pour lequel aucune nouvelle formule ne peut être inférée de manière consistante.

Un théorème sera une formule présente dans toutes les extensions pouvant être inférées.

Une démonstration d'une formule A visera à établir l'existence d'un point fixe pour des prémisses données, qui contienne A

Les règles d'inférence non monotones.

Une logique non monotone suppose donc l'existence, sous une forme ou sous une autre, de règles d'inférence prenant en compte le caractère partiel des connaissances exprimées par les prémisses ou l'existence d'exceptions à une règle générale. Ce seront des règles révisables. Transposées dans un code informatique, un système expert par exemple, le caractère révisable de ces règles supposent qu'elles puissent être bloquées ou inhibées en cours de déduction (par apparition d'une nouvelle connaissance) ou bien qu'elles comportent des conditions d'application dont la vérification puisse évoluer.

Dans l'exemple des trajets routes ou autoroutes, la règle classique serait:

Si x est un trajet autoroutier de A à B alors x est le trajet le plus rapide de A à B.

Avec condition de consistance nous aurions:

Si x est un trajet autoroutier de A à B et s'il est consistant (avec ce qui est déjà connu) que x soit le trajet le plus rapide de A à B , alors x est le trajet le plus rapide de A à B .

Avec condition de non-inférabilité d'une assertion:

Si x est un trajet autoroutier de A à B et si l'on ne peut pas inférer qu'il existe un trajet plus rapide que x de A à B , alors x est le trajet le plus rapide de A à B .

2. Quelques exemples de logiques non monotones

De manière à rendre plus clair ces quelques notions, nous allons présenter rapidement deux exemples de logiques non monotones. Nous verrons après quelles conséquences épistémologiques nous pouvons en tirer.

La logique avec défauts de Reiter.

Définitions.

Reiter (1980, 1987a et b) propose d'ajouter à la logique classique des règles d'inférence d'un type particulier appelées *défauts*. Un défaut se présente ainsi:

$$\frac{\alpha : M\beta}{\gamma}$$

Si α est connu (est une prémisses ou a été déduit des prémisses) et si β est consistant avec tout ce qui est connu, alors on peut inférer γ .

On peut remplacer connu par cru suivant que l'on se place dans une logique de la connaissance ou dans une logique de la croyance.

α , β , γ sont des énoncés prédictatifs implicitement quantifiés universellement.

On aura donc des règles de défaut (ou des défauts) de la forme

$$d: \frac{\alpha(X): M\beta_1(X_1), \dots, M\beta_n(X_n)}{\gamma(Y)}$$

$\alpha(X)$, $\beta_i(X_i)$, $\gamma(Y)$ sont des énoncés bien formés du langage formel utilisé, X , X_i et Y sont des ensembles de variables pris dans l'ensemble des variables du langage.

$\alpha(X)$ est le prérequis du défaut d .

$\beta_i(X_i)$ est la justification du défaut.

$\gamma(Y)$ est le conséquent du défaut.

M appartient au métalangage.

Exemple.

Voici un exemple de défauts emprunté à Reiter:

Une personne habite généralement avec son conjoint

$$d1 \frac{\text{Conjo int}(x, y) \wedge \text{Habite}(y, z) : M \text{Habite}(x, z)}{\text{Habite}(x, z)}$$

Si x est le conjoint de y et si y habite avec z et s'il est consistant avec ce qui est déjà connu de dire que x habite avec z alors x habite avec z

Une personne habite généralement dans la ville de son employeur

$$d2 \frac{\text{Employeur}(x, y) \wedge \text{Ville}(y, z) : M \text{Habite}(x, z)}{\text{Habite}(x, z)}$$

Si x a pour employeur y et si la ville de y est z et s'il est consistant avec ce qui est déjà connu de dire que x habite z alors x habite z

Si l'époux de Joséphine habite Neuchâtel et si son employeur est installé à Paris,

(1) $\text{Habite}(\text{Joséphine}, \text{Neuchâtel})$ et

(2) $\text{Habite}(\text{Joséphine}, \text{Paris})$

sont des formules appartenant à deux extensions incompatibles de la théorie possédant les défauts d_1 et d_2 . Si l'on a, par ailleurs, les prémisses suivantes:

Conjoint(Joséphine,Pierre),
Habite(Pierre,Neuchâtel),
Employeur(Joséphine,IBM) et
Ville(IBM,Paris)

Si l'on applique d'abord le défaut d_1 , on déduira l'énoncé (1) et le défaut d_2 sera bloqué. Si l'on applique d'abord le défaut d_2 , on déduira l'énoncé (2) et le défaut d_1 sera bloqué. On a donc deux extensions:

Conjoint(Joséphine,Pierre), Habite(Pierre,Neuchâtel), Employeur(Joséphine,IBM)
Ville(IBM,Paris) et Habite(Joséphine,Neuchâtel)

Conjoint(Joséphine,Pierre), Habite(Pierre,Neuchâtel), Employeur(Joséphine,IBM)
Ville(IBM,Paris) et Habite(Joséphine,Paris)

Théories normales.

On peut n'utiliser que des défauts dont le conséquent est formé du même énoncé que la justification

$$d: \frac{\alpha(X): M\beta(X_n)}{\beta(X_n)}$$

Dans ce cas on parlera de théories normales. Reiter a montré que de telles théories possédaient les propriétés suivantes:

- elles admettent toujours au moins une extension;
- elles sont semi-monotones, c'est-à-dire que si on ajoute à une théorie normale un défaut, en conservant la normalité, alors la nouvelle théorie ainsi obtenue, toujours normale, admet une extension qui inclut une extension de la théorie de départ.

Quelques problèmes

On se heurte cependant à des problèmes.

1) Soit Δ une théorie normale et f une formule du langage formel utilisé; comment savoir si Δ possède une extension E telle que $f \in E$? Autrement dit, existe-t-il une théorie de la preuve pour les théories normales?

Reiter propose une solution. Soit $\Delta = (D, F)$, une théorie normale, avec D ensemble des défauts et F ensemble des formules de la théorie.

Une séquence finie $D_0, \dots, D_k$ où chaque D_i est un sous-ensemble fini de D , est une preuve de f si et seulement si:

1. $F \cup \{CSQ(D_0)\} \vdash f$
2. $F \cup \{CSQ(D_i)\} \vdash PRQ(D_{i-1})$ pour i de 1 à k
3. $D_k = \emptyset$
4. $F \cup \{CSQ(D_i)\} / 0 \leq i \leq k$ est consistant.

$CSQ(D_i)$: conjonction des conséquents des D_i .

$PRQ(D_i)$: conjonction des prérequis des D_i .

Or pour la construction d'une preuve de f , on partira de $D_k = \emptyset$, puis on construira successivement $D_{k-1}, \dots, D_0$, mais il n'existe pas d'algorithme permettant de passer de D_i à D_{i-1} . La théorie de la preuve dans le cadre de cette logique demeure donc un peu faible.

2) Par ailleurs, on sera parfois amené à utiliser des défauts dont le conséquent n'est qu'une partie de la justification. On aura alors affaire à des théories semi-normales. Cependant les théories semi-normales perdent les propriétés des théories normales (existence d'au moins une extension et semi-monotonie) ce qui rend leur usage plus délicat.

Conclusion

La logique de Reiter fournit donc un formalisme assez proche de ce qui se passe lors d'un raisonnement humain en l'absence d'une connaissance exhaustive ou avec des règles comportant des exceptions. Cependant ce formalisme présente des faiblesses particulièrement s'il doit être utilisé pour des représentations en machine. La sécurité exige des théories normales et celles-ci sont souvent insuffisantes par rapport aux besoins des logiciels envisagés.

Les logiques modales de la connaissance et de la croyance

Généralités

Un autre moyen de formaliser le raisonnement non monotone est d'utiliser le formalisme fourni par les logiques modales en réinterprétant les opérateurs modaux L et M.

LA qui signifiait «il est nécessaire que A» signifiera «A est connu (ou cru)».

MA qui signifiait «il est possible que A» signifiera «non-A n'est pas connu (ou cru)» (donc A est possible et M conserve, d'une certaine manière sa signification initiale).

Les axiomes de la logique modale trouvent aussi une nouvelle interprétation:

Le schéma d'axiome de distribution (axiome K) reste valide:

$$L(A \supset B) \supset (LA \supset LB)$$

Si $A \supset B$ est connu alors la connaissance de A implique la connaissance de B.

L'axiome T devient le schéma d'axiome de la connaissance:

$$LA \supset A$$

Si A est connue alors A est vraie, c'est-à-dire une connaissance est une information vraie.

Le schéma d'axiome 4 devient celui de l'introspection positive:

$$LA \supset LLA$$

Si A est connu alors on sait que A est connu, autrement dit, on sait ce qu'on sait.

Le schéma d'axiome 5 devient celui de l'introspection négative

$$MA \supset LMA$$

Si je ne connais pas non-A alors je sais que je ne connais pas non-A, autrement dit je sais ce que j'ignore.

Ce dernier axiome suppose, si on utilise un système qui le retient, une parfaite connaissance des limites de la connaissance par un agent parfaitement rationnel.

On gardera la règle d'inférence modale de nécessité $A \vdash \text{LA}$ (Si A est valide alors je sais que A).

La sémantique des mondes possibles utilisée avec les logiques modales reste valable dans le cadre de cette interprétation.

Logiques de McDermott et autoépistémique

Des exemples d'utilisation d'une telle interprétation de la logique modale ont été donnés, par exemple, par McDermott et par Moore, par ce dernier sous le nom de logique autoépistémique.

McDermott recherche un système axiomatique universel indépendant des domaines d'application possibles. Le but n'est pas de déterminer telle ou telle extension d'une théorie, mais plutôt de s'intéresser à l'ensemble des formules communes à toutes les extensions possibles.

Comme dans les systèmes modaux habituels on peut bâtir des axiomatiques retenant tout ou partie seulement des axiomes modaux (K, T, 4, 5).

McDermott ajoute une règle d'inférence non monotone sous la forme:

«On ne peut inférer non-A» $\vdash \text{MA}$.

Cette règle permet:

- D'inférer qu'une assertion est possible, c'est-à-dire consistante d'un point de vue logique avec ce qui est déjà connu.
- D'adopter comme vrai un énoncé simplement consistant avec ce qui est connu, en ajoutant l'axiome $\text{MA} \supset A$ et en appliquant le *modus ponens*.
- De conférer au système axiomatique un caractère non monotone puisque, si on a inféré MA, puis B à partir de MA, et si on ajoute une prémisses permettant de déduire non-A, alors on devra rétracter B.

La logique autoépistémique de Moore utilise aussi la logique modale en interprétant LA par «A est cru» et MA par «non-A n'est pas cru».

Comparaison entre la logique de McDermott et la logique auto-épistémique de Moore.

McDermott définit la notion de point fixe de la manière suivante. Soient un système axiomatique modale utilisant tout ou partie des axiomes modaux (S_{TK} ou S_{TK4} ou S_{TK5}) et un ensemble de prémisses Δ .

Γ est un point fixe pour Δ si et seulement si Γ est l'ensemble des conséquences modales par S_{TK} ou S_{TK4} ou S_{TK5} de

$$\Delta \cup \{\neg Lf \mid f \notin \Gamma\}.$$

La logique autoépistémique définit la notion d'expansion stable (comme équivalent de celle de point fixe) de la manière suivante:

Γ est une expansion stable pour Δ si et seulement si Γ est l'ensemble des conséquences modales par S_{TK} ou S_{TK4} ou S_{TK5} de $\Delta \cup \{\neg Lf \mid f \notin \Gamma \cup Lf \mid f \in \Gamma\}$.

Quelques problèmes

Premièrement les difficultés que la logique modale connaît lorsqu'il s'agit de passer du calcul des propositions à celui des prédicats se retrouvent nécessairement dans le cas d'une utilisation de ces logiques pour des logiques non monotones.

Deuxièmement, on s'aperçoit que si l'on utilise un système axiomatique possédant l'axiome 5 $MA \supset LMA$ et si $A \vdash f$ selon la logique de McDermott alors $A \vdash f$ selon la simple logique modale c'est-à-dire sans la règle d'inférence non monotone. Autrement dit il n'y a pas de théorème non monotone si l'on adopte l'axiome 5. McDermott choisit donc un système axiomatique ne comportant que les axiomes T, K et 4.

Troisièmement la notion de point fixe chez McDermott conduit à la remarque suivante de Moore: *Un agent de la logique non monotone de McDermott est omniscient à propos de tout ce qu'il ne croit pas, mais peut tout ignorer de ce qu'il croit.* En fait il manque à la logique de McDermott le $\{Lf \mid f \in \Gamma\}$ qu'ajoute la notion d'expansion stable de Moore.

Quatrièmement, l'axiome T ($LA \supset A$) ne pose pas de problème si LA signifie A est connu. En revanche, si LA signifie A est cru, cet axiome revient à dire que toute proposition crue par un agent est valide donc que tout ce qui est cru est vrai. Or la

non-monotonie exclut la notion de validité ou de vérité au sens de la logique classique puisqu'elle suppose la révisabilité des croyances. Or avec cet axiome rien chez McDermott ne permet de différencier le valide du simplement consistant avec des prémisses. C'est pourquoi Moore propose de rejeter l'axiome T et non l'axiome 5.

3. Enjeux épistémologiques

Nous venons d'avoir un aperçu rapide des formalismes utilisables pour une logique non monotone. J'ai parfaitement conscience du caractère très incomplet de cet aperçu, d'une part quant à l'exposition des logiques de Reiter, de McDermott ou de Moore et, d'autre part du fait que nous avons ignoré beaucoup d'autres travaux. Il aurait probablement fallu parler des techniques dites «de minimisation» telles que l'hypothèse du monde clos, la complétion de prédicats ou la circonscription qui, chacune à leur manière, introduisent la non-monotonie. Cependant, j'espère en avoir suffisamment dit pour pouvoir clairement aborder quelques enjeux épistémologiques de la non-monotonie.

Le vrai et le consistant

Il est tout d'abord intéressant de comparer la conception de la vérité qui ressort de l'usage de logiques non monotones à celle qui est habituellement associée à la sémantique de la logique classique. Nous pouvons partir de la définition que donne Tarski [TAR 44]

X est vrai si et seulement si p

où X est le nom de la proposition «p» dont l'exemple canonique est:

La proposition «la neige est blanche» est vrai si et seulement si la neige est blanche.

Ainsi, selon Tarski, reprenant explicitement une conception d'Aristote, la vérité d'une proposition consiste en son accord (ou sa correspondance) avec la réalité.

Or, toute la logique depuis Aristote, consiste à établir les règles qui permettent de faire passer le caractère vrai d'un énoncé à un autre. Les notions de déduction, d'inférence trouvent implicitement leur définition dans cette conception. En effet dire qu'un énoncé B se déduit d'un énoncé A, c'est dire que A étant vrai, B le sera aussi s'il existe une règle d'inférence reconnue permettant de passer de A à B. Dans le cadre de la logique classique, des énoncés pris comme prémisses étant vrais, soit parce qu'ils consistent en énoncés observationnels intersubjectivement reconnus, soit parce qu'ils ont été déduits selon des règles d'inférence valides d'énoncés observationnels, soit enfin parce qu'ils ont un caractère tautologique, les conclusions qui s'ensuivent selon encore des règles d'inférence valides, sont aussi vraies.

Dans la conception de Tarski, la vérité coule donc de deux sources: l'observation de ce qui est (les énoncés observationnels) et la tautologie, et son cours parcourt le discours correctement dans la mesure où il est guidé par des règles d'inférence admises et éprouvées. Nous sommes donc placés dans une perspective empiriste. Il suffirait d'ajouter que seuls les énoncés ainsi caractérisés sont scientifiques et pourvus de sens pour retrouver la position de Carnap (tout du moins dans ses premiers écrits).

Or il en va assez différemment dans le cadre des logiques non monotones. En effet, un ensemble de prémisses peut-être composé d'énoncés vrais au sens où nous venons de le définir mais le fait qu'il puisse s'en déduire plusieurs extensions mutuellement incompatibles, nous oblige à dissocier la notion de vérité de la notion de consistance ou de non contradiction. Ces deux notions vont de pair dans le cadre de la logique classique puisque les règles de déduction assurent qu'il n'y aura pas de contradiction entre les conclusions et les prémisses d'une part, et entre les conclusions entre elles d'autre part. Tout le problème de la consistance réside donc dans les prémisses ou dans le système d'axiomes adopté dans le cas d'une axiomatisation. Dans un système non monotone la consistance doit demeurer entre les prémisses et les conclusions mais non entre les conclusions.

On peut encore dire les choses d'une autre manière. En logique classique si une déduction est valide, les conclusions pourront être dites aussi valides, et les prémisses étant vraies, la validité des conclusions entraîne nécessairement leur vérité. Il n'en est plus de même en logique non monotone. Accordons encore que les prémisses soient vraies, les conclusions valides ne sont plus nécessairement vraies au sens que Tarski donne au mot.

Limites de l'empirisme et logique de la découverte scientifique

La position empiriste, dans sa forme la plus radicale, celle du premier Carnap par exemple, soulève un certain nombre de problèmes bien connus. Les logiques non monotones, sans apporter une solution définitive et complète peuvent cependant ouvrir quelques voies à la réflexion.

Dans le cadre d'une recherche scientifique, le savant émet des hypothèses, ensemble d'énoncés constituant une théorie au sens logique du terme. Les énoncés de cette théorie ne sont généralement ni des énoncés observationnels, ni des tautologies. Un strict empirisme qualifierait donc de tels énoncés de métaphysiques. Ces hypothèses pourront être confrontées à des procédés expérimentaux donc à des énoncés observationnels, et si ce n'est ces hypothèses elles-mêmes ce seront des énoncés déduits de ces hypothèses selon les règles de la logique. Les énoncés déduits d'hypothèses, comme les hypothèses elles-mêmes, avant toute confrontation à des expériences validantes ou invalidantes (position vérificationniste ou falsificationniste), ne relèvent pas de la catégorie du vrai (seuls en relèvent les énoncés observationnels ou les tautologies), mais d'une simple consistance sur l'ensemble des énoncés considérés (hypothèses, énoncés déduits et éventuellement énoncés constituant le corpus de la science en question). A une hypothèse donnée peut donc correspondre plusieurs extensions au sens des logiques non monotones. Une expérience peut alors permettre d'obtenir un énoncé observationnel qui n'invalidera pas l'hypothèse mais permettra de choisir entre l'une ou l'autre des extensions possibles. Dans le cadre de la logique classique une hypothèse pou-

vait être, suite à une expérience, soit corroborée, confortée (vérifiée ou non falsifiée) soit rejetée. Dans le cadre des logiques non monotones, l'expérimentation peut aussi viser à rejeter une hypothèse, mais tout autant à déterminer quelle extension parmi celles qui sont possibles, correspond à la réalité.

On voit donc, que les logiques non monotones peuvent offrir un cadre formel pour une logique de la découverte scientifique. Bien sûr, le formalisme proposé par la logique des défauts de Reiter ou la logique de McDermott ou la logique auto-épistémique, ne peuvent prétendre épuiser la question de la logique de la découverte scientifique. Cependant, la philosophie des sciences qui s'intéresse à l'élaboration et à la construction des théories scientifiques en lien avec les procès expérimentaux, ne peut pas ignorer de tels formalismes et les services qu'ils peuvent rendre. Nous venons de donner les principes généraux de tels services, il resterait maintenant à les confronter à la pratique scientifique elle-même, sur des cas concrets.

Le sujet de la connaissance

Il y a une dernière conséquence des logiques non monotones qu'il faut maintenant aborder: c'est le rapport qu'entretient la connaissance avec le sujet qui la porte. Pour la logique classique, la vérité ou la fausseté d'un énoncé ne dépend en aucune façon du sujet qui l'émet. Peu importe qui dit que « $2+2 = 4$ » ou que «la neige est blanche». Ces énoncés possèdent une vérité qui peut être *a priori*, innée, le fait d'une réminiscence de type platonicien ou s'établir sur la base d'une observation ou d'une expérience, elle ne dépend cependant absolument pas de l'esprit auquel ils sont innés, ni de celui qui les reconnaît, ni même de l'observateur ou de l'expérimentateur. On peut appeler cela le caractère objectif du vrai et du faux dans le cadre de la logique classique.

En revanche, les logiques non monotones ramènent brutalement le sujet de la connaissance au coeur même de celle-ci. Dans la mesure où elles se prêtent à la formalisation d'une connaissance incomplète ou à celle de la simple croyance, elles se réfèrent nécessairement à l'état cognitif d'un sujet à un moment donné, qu'il s'agisse d'un sujet humain ou de ce qu'on

pourrait appeler d'un «sujet artificiel» dans le cas où elles interviennent dans des systèmes informatiques de gestion de bases de connaissances. On peut montrer cet aspect particulier des logiques non monotones aussi bien avec le formalisme des défauts de Reiter qu'avec les logiques modales de McDermott ou auto-épistémiques.

Dans le cas de la logique des défauts de Reiter, le surgissement du sujet est directement lié à la pluri-extensionnalité. Les règles de déduction appliquées à une théorie comportant des défauts conduisent à l'une ou l'autre des extensions possibles en fonction de l'ordre dans lequel sont appliqués les défauts. Nous venons d'en voir un exemple. Or cet ordre n'est absolument pas déterminé par les différents éléments de la théorie (formules ou défauts). Il renvoie à un élément extérieur à la théorie, à celui qui effectue la déduction et donc en fin de compte à un sujet. Evidemment cela ne signifie pas que l'ordre d'utilisation des défauts soit complètement laissé au hasard. Le sujet déduisant peut obéir à des règles. Par exemple, nous pouvons utiliser l'une ou l'autre des règles suivantes: a) *Utiliser les défauts dans l'ordre de leur écriture*; ou b) *Utiliser prioritairement les défauts qui, statistiquement, ont donné le plus de réussite*; ou encore c) *Elaborer un classement des défauts sur une base expérimentale*; etc. Peu importe, dans tous les cas cela revient à introduire des métarègles. Cependant de telles métarègles conduisent, si elles sont correctement formulées et appliquées, à effectuer, pour une théorie donnée, un choix parmi les extensions, qui sera toujours le même, la déduction étant répétée autant de fois que l'on voudra. L'introduction de telles métarègles reviendrait donc à éliminer la non-monotonie, ou plutôt à la déplacer à une autre niveau, celui du choix de la métarègle utilisée. Il ne s'agit pas de s'engager dans une régression à l'infini en créant des métathéories non monotones qui elles-mêmes appelleraient des méta-métarègles et ainsi de suite. Simplement, on est bien obligé de remarquer que, à quelque niveau que l'on se place, il faut faire intervenir à ce niveau-là une procédure de choix (choix d'un ordre d'application des défauts, choix d'une métarègle, ou d'une méta-métarègle...) et que tout choix renvoie inévitablement à celui d'un sujet qui l'effectue.

Dans le cas des logiques non monotones utilisant le formalisme de la logique modale, la référence au sujet se fait de manière peut-être plus formelle, à la fois plus subtilement et plus explicitement. En effet tout se joue autour de trois axiomes que nous avons rapidement abordés mais qui méritent un examen plus attentif.

L'axiome T de la logique modale qui devient l'axiome de la connaissance $LA \supset A$, peut se lire «Si A est connu alors A est vrai». Bien sûr l'interprétation classique de l'implication donnent trois possibilités de vérité de cette formule, si l'on en fait un axiome. D'abord LA est vraie (A est connu) et A est vrai. Ensuite LA est faux (A n'est pas connu) mais A est vrai. Enfin LA est faux (A n'est pas connu) et A est faux. Le deuxième cas ménage la possibilité qu'un énoncé soit vrai et pourtant non connu comme tel, et par là il semble séparer la question de la vérité d'un énoncé de celui du sujet qui énonce cette vérité. En effet, comment peut-on dire en même temps qu'un énoncé A est vrai et qu'il n'est pas connu? Cela ne peut raisonnablement se faire que si LA signifie en fait «A est connu comme vrai». Alors effectivement il peut se faire que quelqu'un ne connaisse pas A comme vrai (c'est-à-dire le croit faux ou même l'ignore complètement) et qu'en même temps A soit vrai. Mais cette interprétation oblige à introduire un «quelqu'un», c'est-à-dire un sujet de la connaissance. Le premier cas lie explicitement la vérité d'un énoncé à une connaissance correspondante de sorte qu'une connaissance est un énoncé vrai. Enfin, dans le troisième cas, la fausseté de A est impliquée par le fait que A ne soit pas connu. On a donc trivialement un renvoi à un sujet, celui de la connaissance. En résumé, le fait, dans cet axiome d'avoir dans l'antécédent un énoncé soumis à la modalité «être connu» et dans la conclusion le même énoncé nu de toute modalité, établit un lien entre la vérité et le sujet qui a accès à cette vérité.

L'axiome 4 et l'axiome 5, respectivement d'introspection positive et d'introspection négative qui supposent que l'on sait ce que l'on sait et que l'on sait ce que l'on ne sait pas, semblent, à leur façon, dégager la connaissance de son sujet. Cependant à y regarder de plus près ce n'est pas si simple. Tout d'abord cela peut bien être admis pour l'axiome 5 puisque, comme l'a montré

McDermott, un système doté de cet axiome, malgré la règle d'inférence non monotone, ne possède pas de théorème non monotone. Autrement dit, si l'axiome 5, comme nous venons de le dire déconnecte la notion de vérité de celle du sujet d'un énoncé, il fait perdre aussi la non-monotonie. C'est la raison pour laquelle McDermott proposera d'adopter un système axiomatique ne possédant que les axiomes T, K et 4, à l'exclusion du 5. Cependant, comme nous l'avons vu, une des conséquences de la logique de McDermott est, pour reprendre les termes de la critique de Moore, qu'*un agent de la logique non monotone de McDermott est omniscient à propos de tout ce qu'il ne croit pas, mais peut tout ignorer de ce qu'il croit*. Ce qui revient à dire que l'agent de la connaissance est déterminé ou caractérisé par le système d'axiomes choisi. La solution de Moore qui consiste à conserver l'axiome 5 mais à rejeter l'axiome T maintient toutefois la référence à un sujet par le biais de celui de l'introspection positive (axiome 4).

En résumé, la logique non monotone lie la notion de vérité à un moment, un lieu puisqu'elle permet de prendre en compte un ensemble évolutif de connaissances et à un agent. Ce dernier étant tout aussi bien un sujet humain qu'un système artificiel.

Conclusion

Ce très bref exposé est loin de faire le tour des enjeux épistémologiques des logiques non monotones. Une étude beaucoup plus détaillée des techniques de la non-monotonie permettrait d'approfondir les questions qui viennent d'être soulevées et en ferait certainement surgir bien d'autres. Cependant, j'espère avoir montré deux choses. Premièrement les logiques non monotones présentent un intérêt pratique non négligeable pour modéliser certaines formes de raisonnement qui échappent à la logique classique et s'avèrent commodes pour certaines réalisations informatiques relevant de l'intelligence artificielle. Deuxièmement, elles permettent de poser de manière renouvelée des questions épistémologiques relatives à la possibilité d'une logique de la découverte scientifique au sens où l'entendait, par

exemple, Hanson et relatives à la théorie de la connaissance. Si cet exposé est loin d'être exhaustif, j'espère qu'il peut montrer l'existence possible d'un débat à la fois technique et philosophique.

Faculté des lettres
Université de Bourgogne
Bd Gabriel
F 21000 DIJON

Bibliographie

- FROIDEVAUX C. (1986). Axonomic default theory. *Proc. of the 7th European Conf. On A.I. (ECAI-86)*, Brighton, 123-129.
- GREGOIRE E. (1990). *Logiques non monotones et intelligence artificielle*. Paris: Hermès.
- MCDERMOTT D., DOYLE J. (1980). Non-monotonic logic I, *Artificial Intelligence* 13, 41-72.
- MCDERMOTT D. (1982). Non-monotonic logic II: non-monotonic modal theories. *J. of the ACM* 29, 34-57.
- MCDERMOTT D. (1987). A critique of pure reason. *Computational Intelligence* 3, 151-160.
- MOORE R.C. (1984). Possible-world semantics for autoepistemic logic. *Proc. of the AAAI-Workshop on Non-Monotonic Reasoning*. New York: New Paltz, 344-354.
- MOORE R.C. (1988). Autoepistemic logic. In: Smets P. & alii (eds) *Non-Standard Logics for Automated Reasoning*. London: Academic Press, 105-136.
- REITER R. (1980). A logic of default reasoning. *Artificial Intelligence* 13, 81-131.
- REITER R. (1987a). A theory of diagnosis from first principles. *Artificial Intelligence* 32, 57-95.
- REITER R. (1987b). Nonmonotonic reasoning. *Ann. Rev. Comp. Science* 2, 147-186.

TARSKI A. (1974). La conception sémantique de la vérité. In:
Logique, sémantique, métamathématique. Paris: Colin,
265 sq.

LE STATUT LOGIQUE DE LA CONTRADICTION: L'APPROCHE PARACONSISTANTE

Henri VOLKEN

*Nous entrons et nous n'entrons pas dans les
mêmes fleuves; nous sommes et nous ne
sommes pas. Héraclite*

1. Le principe «ex contradictione quodlibet»

Les premières considérations systématiques sur la contradiction remontent, dans notre culture occidentale du moins, aux Eléates. Elles ont permis, selon toute vraisemblance, à Parménide et à ses disciples de développer le style du raisonnement indirect si caractéristique de la pensée mathématique et scientifique. Ce raisonnement est appelé aussi *raisonnement par l'absurde*, par référence au statut négatif qu'il fait jouer à la contradiction. En effet, la première apparition d'une contradiction permet de rebrousser chemin et de nier l'hypothèse qui se trouve au départ. Une contradiction signale donc une impossibilité, une absurdité, un désaccord définitif avec le monde dans lequel et à propos duquel se déroule le raisonnement. Ce monde lui-même est supposé être consistant¹.

Tout à l'opposé on trouve Héraclite pour qui la contradiction est un révélateur de la réalité qui rend celle-ci compréhensible: une chose ne s'explique que par rapport à son contraire. Le devenir est le principe du monde. Le devenir lui-même s'explique par l'opposition et la contradiction. Le monde pour Héraclite est contradictoire.

1 Nous utiliserons le terme de *consistant* pour une théorie ne comportant pas de contradiction et le terme *absolument consistant* pour une théorie non triviale.
Travaux de logique, 11, 1997.

On trouve plus tard chez Aristote une formulation de plusieurs principes réglant le sort de la contradiction, du moins en ce qui concerne son statut logique. Il s'agit notamment du principe de la non-contradiction et du principe du tiers-exclu. Le premier stipule que φ et $\neg\varphi$ ne peuvent pas être vrais simultanément, alors que le second demande que l'un des deux au moins soit vrai. Par la suite, ces principes ont été généralement admis sans réserve par les logiciens et mathématiciens. Héraclite est resté l'*obscur* – comme il avait été surnommé – et son influence se limite au courant dialectique, représenté plus près de nous notamment par Hegel.

Les premiers signes d'une attitude plus critique en ce qui concerne la place de la contradiction dans le raisonnement sont apparus au début du vingtième siècle et plus précisément autour des années 1910. Divers paradoxes sémantiques ont alors ébranlé les «fondements» récents de la logique et des mathématiques, dont le plus célèbre est celui que Russell décrivit dans une lettre à Frege en 1902 et qui porte son nom aujourd'hui. D'autre part à cette époque les discussions autour de la géométrie non euclidienne ne sont pas éteintes. Ces discussions portent en particulier sur le statut des axiomes mais aussi sur l'existence de modèles non standards.

En ce qui concerne la contradiction, trois événements témoignent d'une réaction sérieuse contre les principes aristotéliens. Ces événements ont eu lieu presque en même temps.

Tout d'abord Lukasiewicz (1971) discute le principe de la non-contradiction dans un article où il envisage une logique qui pourrait se passer de ce principe. Il y montre même que la théorie des syllogismes d'Aristote ne nécessite pas ce principe. Ce travail ouvre la voie à de nouvelles formes de logique, en particulier des logiques multivalentes.

Autour de la même date, une discussion animée s'est engagée autour des thèses de Meinong et de sa théorie des objets qui permet la description et l'utilisation d'objets non existants. Les objets peuvent avoir, selon cette théorie, des qualités contradictoires. Cette théorie préfigure déjà la sémantique des mondes possibles de Kripke.

De son côté, le logicien russe Vassiliév introduisit une «logique imaginaire» qui ne supposait pas le principe de non-contradiction, en tout cas pour la partie de sa logique qui touchait le «jugement sur des concepts», par opposition à la partie qui traitait des «jugements sur des faits». Là encore, la démarche imite un peu celle de la géométrie non euclidienne en ce sens que des principes jusque-là admis comme évident, étaient remis en question.

Mais le principe logique fatal à toute tentative de tolérance envers la contradiction est le principe connu sous le nom de «ex contradictione quodlibet». Il peut être formalisé ainsi:

$$\varphi, \neg\varphi \vdash \psi \text{ (ECQ)}$$

et signifie que d'une contradiction on peut déduire n'importe quelle formule. Toute théorie comportant une contradiction explose donc littéralement dans le sens déductif. Une logique dans laquelle le principe (ECQ) est valide est appelée logique *explosive*. Une logique *paraconsistante* est simplement une logique non explosive.

Le principe (ECQ) n'est pas un principe premier de la logique classique ou intuitionniste, mais il est la conséquence de plusieurs principes plus élémentaires. Pour le voir il suffit de regarder les détails de la démonstration classique de Lewis. (Cette démonstration a une valeur emblématique dans tous les textes d'introduction aux logiques paraconsistantes). Voici une version de cette démonstration:

- | | | |
|-----|--|--------------------------|
| (1) | φ | hypothèse |
| (2) | $\neg\varphi$ | hypothèse |
| (3) | $\neg(\varphi \wedge \neg\psi)$ | expansion de (2) |
| (4) | $\varphi \wedge \neg(\varphi \wedge \neg\psi)$ | adjonction de (1) et (3) |
| (5) | ψ | sylogisme disjonctif (4) |

La preuve est basée sur quatre principes logiques: les trois qui figurent explicitement dans la démonstration, expansion, adjonction et syllogisme disjonctif, et la transitivité de la relation de conséquence.

Pour qu'une logique soit paraconsistante, il faut donc, d'après la définition que nous en avons donnée, qu'elle remette en ques-

tion l'un au moins de ces principes. Nous allons brièvement décrire quelques grandes catégories de logiques paraconsistantes en les classant selon les principes évoqués plus haut. Ensuite nous verrons plus en détail deux exemples importants, la logique relevante et la logique des paradoxes.

2. Quelques logiques non explosives

Les logiques qui mettent en cause la règle de l'*expansion* exigent que:

$$\neg\varphi \vdash \neg(\varphi \wedge \neg\psi)$$

Mais par contre dans ces logiques, la contraposée

$$\varphi \wedge \neg\psi \vdash \varphi$$

reste souhaitée.

C'est donc essentiellement l'opération de contraposition qui est mise en question. C'est le chemin choisi par da Costa (1974). Il permet de définir une hiérarchie de logiques C_i à l'aide de conditions sémantiques indépendantes pour φ et pour $\neg\varphi$, ainsi que des conditions supplémentaires. Contraposition et (ECQ) sont invalides dans ces logiques.

Les logiques non adjonctives, celles qui attaquent le principe de l'*adjonction*, sont caractérisées par la condition

$$\varphi, \psi \not\vdash \varphi \wedge \psi$$

L'intuition cachée derrière cette approche est l'idée qu'un discours, ou une base de données, sont l'oeuvre de plusieurs protagonistes, qui pris séparément, sont consistants, mais qui peuvent introduire des contradictions entre eux. La logique discursive de Jaskowski (1969) est de ce type: ni adjonction, ni (ECQ) ne sont valides.

On peut mettre en doute également la *transitivité* de la relation de conséquence. Une tentative dans ce sens a été proposée par Smiley en 1959, dans le sens d'une restriction de la relation de conséquence à partir des inférences classiques:

$\varphi \vdash \psi$ ssi elle est classiquement valide et si $F(\varphi, \psi)$

La condition F, le filtre, peut prendre des formes diverses. Par exemple

$F(\varphi, \psi)$ ssi φ n'est pas une contradiction,
et ψ pas une tautologie.

Dans ce cas on constate que toutes les étapes de (ECQ) sont valides mais que (ECQ) lui-même ne l'est pas. La conséquence n'est donc pas transitive. Les logiques de ce type sont appelées *filtrées*.

Dans la recherche d'une logique non explosive, on peut enfin vouloir éliminer la propriété du *syllogisme disjonctif*. On peut caractériser les logiques dans lesquelles ce principe n'est pas valide, par leur sémantique basée sur la notion de treillis de de Morgan. Il s'agit d'un treillis distributif muni d'un opérateur unaire involutif et anti-monotone, qui fixera le comportement de la négation. La relation de conséquence est définie par:

$\varphi \models \psi$ ssi pour toute interprétation v , $v(\varphi) \leq v(\psi)$

Les logiques de de Morgan sont le plus souvent des logiques pertinentes, bien qu'il existe des exemples intéressants pour lesquels ce ne soit pas le cas.

Nous allons voir un peu plus en détail deux exemples de logiques paraconsistantes qui ont une importance particulière pour la suite.

2.1. La logique pertinente RQ

Cette logique est due à Anderson et Belnap (1975). Elle est née de la préoccupation de définir une opération conditionnelle qui ne présente pas les défauts de la conditionnelle classique. Il est exigé notamment que l'antécédent et le conséquent possèdent en commun au moins une variable propositionnelle, ce qui donne une certaine pertinence («relevance» en anglais) à la conditionnelle.

Nous allons en donner une description syntaxique qui montre le mécanisme de la conditionnelle qui n'est plus vérifonctionnelle, c'est-à-dire qui ne peut plus être définie à l'aide de la négation et de la disjonction. Il est possible de définir la conditionnelle classique, l'implication matérielle, par $\varphi \supset \psi$ ssi $\neg\varphi \vee \psi$. Mais la règle du modus ponens n'est plus valide pour cette dernière opération. Voici les axiomes de RQ:

- R1 $\varphi \rightarrow \varphi$
- R2 $(\varphi \rightarrow \psi) \rightarrow ((\psi \rightarrow \theta) \rightarrow (\varphi \rightarrow \theta))$
- R3 $\varphi \rightarrow ((\varphi \rightarrow \psi) \rightarrow \psi)$
- R4 $(\varphi \wedge \psi) \rightarrow \varphi$
- R5 $(\varphi \wedge \psi) \rightarrow \psi$
- R6 $((\varphi \rightarrow \psi) \wedge (\varphi \rightarrow \theta)) \rightarrow (\varphi \rightarrow (\psi \wedge \theta))$
- R7 $\varphi \rightarrow (\varphi \vee \psi)$
- R8 $\psi \rightarrow (\varphi \vee \psi)$
- R9 $((\varphi \rightarrow \theta) \wedge (\psi \rightarrow \theta)) \rightarrow ((\varphi \vee \psi) \rightarrow \theta)$
- R10 $(\varphi \wedge (\psi \vee \theta)) \rightarrow ((\varphi \wedge \psi) \vee (\varphi \wedge \theta))$
- R11 $\neg\neg\varphi \rightarrow \varphi$
- R12 $(\varphi \rightarrow \neg\varphi) \rightarrow \neg\varphi$

On obtient la logique RM en rajoutant l'axiome R13:

- R13 $\varphi \rightarrow (\varphi \rightarrow \varphi)$

Ceci forme la partie propositionnelle de la logique RQ. Les axiomes concernant les quantificateurs sont:

- R14 $\forall x(\varphi x \rightarrow \varphi t)$ (t est un terme quelconque)
- R15 $\forall x(\varphi \rightarrow \psi) \rightarrow (\forall x\varphi \rightarrow \forall x\psi)$
- R16 $\varphi \rightarrow \forall x\varphi$ (x pas libre dans φ)
- R17 $\forall x(\varphi \vee \psi) \rightarrow (\varphi \vee \forall x\psi)$ (x pas libre dans φ)
- R18 $(\forall x\varphi \wedge \forall x\psi) \rightarrow \forall x(\varphi \wedge \psi)$

Les deux règles de dérivation sont:

- I. Si φ et ψ sont des théorèmes, alors $\varphi \wedge \psi$ l'est aussi.
- II. Si φ et $\varphi \rightarrow \psi$ sont des théorèmes, alors ψ l'est aussi.

La sémantique correspondante est du type RM3. Il s'agit d'une logique à trois valeurs, $\{f, p, v\}$, ces valeurs étant ordonnées alphabétiquement, p et v étant les valeurs désignées: $\{p, v\} = \nabla$. On définit une interprétation I en imposant les conditions

- I1 Pour chaque formule atomique α , $I(\alpha) \in \{f, p, v\}$
- I2 $I(\neg\varphi) = v, p, f$ ssi $I(\neg\varphi) = f, p, v$
- 3 $I(\varphi \wedge \psi) = \min\{I(\varphi), I(\psi)\}$
- I4 $I(\varphi \vee \psi) = \max\{I(\varphi), I(\psi)\}$
- $I(\varphi \rightarrow \psi) = I(\neg\varphi \vee \psi)$ si $I(\varphi) \leq I(\psi)$
- I5 et $I(\neg\varphi \wedge \psi)$ sinon
- I6 $I(\forall x\varphi x) = \min\{a: \text{pour un certain } t, I(\varphi t) = a\}$
- 7 $I(\exists x\varphi x) = \max\{a: \text{pour un certain } t, I(\varphi t) = a\}$

Validité et conséquence sont définies par

$$\Gamma \models \varphi \text{ ssi il n'y a pas de } I \text{ tel que } I(\varphi) = f \\ \text{alors que pour tout } \gamma \in \Gamma, I(\gamma) \in \nabla$$

et

$$\models \varphi \text{ ssi pour tout } I, I(\varphi) \in \nabla$$

2.2. La logique des paradoxes

La logique des paradoxes, LP, est un autre exemple de logique paraconsistante. Nous allons en donner une présentation sémantique. Le 0 et le 1 symboliseront le vrai et le faux. Les trois valeurs logiques seront définies à partir de là. $\{0\}$ sera la valeur faux, $\{1\}$ la valeur vrai et $\{0, 1\}$ la valeur paradoxale vrai-et-faux.

Une interprétation est donnée par le couple $\langle D, I \rangle$, où D est un domaine non vide et I une fonction qui attribue à chaque symbole non logique une dénotation: à chaque constante c , un élément $I(c)$ dans D , à chaque symbole de fonction n -aire f , une

fonction n -aire $I(f)$ sur D , et à chaque symbole de prédicat P n -aire, un couple de sous-ensembles de D^n , appelés extension et anti-extension de P , qui ensemble forment un recouvrement de D . On notera $I(P) = \langle I^+(P), I^-(P) \rangle$. L'extension du prédicat '=' sera toujours la diagonale de D , mais son anti-extension peut recouvrir la diagonale partiellement.

Les conditions de vérité dans cette interprétation sont définies ainsi: (v est la fonction sémantique associée à l'interprétation I)

- $1 \in v(Pt_1...t_n)$ ssi $\langle I(t_1), \dots, I(t_n) \rangle \in I^+(P)$
- $0 \in v(Pt_1...t_n)$ ssi $\langle I(t_1), \dots, I(t_n) \rangle \in I^-(P)$
- $1 \in v(\neg\varphi)$ ssi $0 \in v(\varphi)$
- $0 \in v(\neg\varphi)$ ssi $1 \in v(\varphi)$
- $1 \in v(\varphi \wedge \psi)$ ssi $1 \in v(\varphi)$ et $1 \in v(\psi)$
- $0 \in v(\varphi \wedge \psi)$ ssi $0 \in v(\varphi)$ ou $0 \in v(\psi)$
- $1 \in v(\forall x\varphi)$ ssi $1 \in v(\varphi[x/d])$ pour tous les $d \in D$
- $0 \in v(\forall x\varphi)$ ssi $0 \in v(\varphi[x/d])$ pour un $d \in D$ au moins

La disjonction et la quantification existentielle sont définies de manière standard.

La vérité est définie de manière suivante: φ est vrai dans l'interprétation I ssi $1 \in v(\varphi)$ et φ est faux ssi $0 \in v(\varphi)$. Une formule φ est une vérité logique ssi elle est vraie dans toute interprétation I . On a les résultats importants:

$$\models_{LP} \varphi \text{ ssi } \models_{cl} \varphi$$

et

$$\text{Si } \Gamma \models_{LP} \varphi \text{ alors } \Gamma \models_{cl} \varphi$$

Dans le cas de la conséquence logique, la réciproque n'est pas vraie. (ECQ) est un contre-exemple, comme on peut le vérifier facilement. Il suffit en effet de considérer une interprétation I dont la fonction sémantique v attribue les valeurs 0 à $\varphi, \neg\varphi$ et ψ , ainsi que la valeur 1 à φ et $\neg\varphi$.

3. Arithmétiques inconsistantes

La base logique des arithmétiques que nous allons considérer est donnée par RQ, la logique relevante donnée plus haut et plus particulièrement la sémantique RM3 utilisant trois valeurs de vérité. Les axiomes arithmétiques sont les suivants. Il s'agit essentiellement des axiomes de Peano.

- A1 $\forall x \forall y (x = y \rightarrow x' = y')$
- A2 $\forall x \forall y \forall z (x = y \rightarrow (x = z \rightarrow y = z))$
- A3 $\forall x (\neg x' = 0)$
- A4 $\forall x (x + 0 = x)$
- A5 $\forall x \forall y (x + y' = (x + y)')$
- A6 $\forall x (x \times 0 = 0)$
- A7 $\forall x \forall y (x \times y' = (x \times y) + x)$

Il faut rajouter la règle d'induction

- R1 Si $\varphi 0$ et $\forall x (\varphi x \rightarrow \varphi x')$ sont des théorèmes, alors $\forall x \varphi x$ l'est aussi.

La différence entre l'arithmétique R# ainsi définie et l'arithmétique de Peano réside dans la logique sous-jacente. Ici il s'agit de R, une logique dans laquelle la conditionnelle n'est pas l'implication matérielle.

Nous allons définir une structure pour cette arithmétique qui permettra de montrer la consistance absolue, c'est-à-dire la non-trivialité, de R# par des moyens finis.

Un modèle de type RM3ⁱ est un couple $\langle D^i, I \rangle$ où le domaine est celui des entiers modulo i , et I associe aux symboles d'opération les opérations correspondantes modulo i . Aux formules atomiques $t_1 = t_2$, I attribue la valeur p si $I(t_1) = I(t_2)$ et la valeur f sinon. Pour les formules composées, ce sont les règles de RM3 qui s'appliquent.

Nous appellerons *arithmétique RM3ⁱ* l'ensemble des formules vraies dans cette structure. Une telle arithmétique possède les propriétés suivantes: elle est inconsistante, non triviale, com-

plète, décidable et constitue une extension de $R\#$ pour n'importe quelle valeur de i .

De plus, $RM3^i$ est axiomatisable de la manière suivante: aux axiomes de $R\#$ on rajoute les formules $0 = i$ et pour chaque j tel que $0 \leq j < i$, les formules $0 = j \leftrightarrow 0 = 1$. Les théorèmes de cette nouvelle arithmétique $RM3^{i\#}$ sont exactement les vérités logiques de $RM3^i$.

L'utilisation d'une logique paraconsistante permet de démontrer par des moyens très simples la consistance absolue de $R\#$, contrairement à l'arithmétique de Peano basée sur la logique classique. En effet $I(0=1)=f$, mais tous les théorèmes de $R\#$ obtiennent une valeur désignée.

Dans la logique $RM3^9$ par exemple on peut montrer que la forme suivante du théorème de Fermat n'est pas vraie:

Fermat

$$\forall x \forall y \forall z \forall n (x^n + y^n = z^n \rightarrow (x = 0 \vee x = 1 \vee y = 0 \vee y = 1 \vee z = 0 \vee z = 1 \vee n = 0 \vee n = 1 \vee n = 2))$$

Pour le voir il suffit de prendre les valeurs $x=y=z=n=3$. L'antécédent de la conditionnelle prend la valeur p, alors que le conséquent prend la valeur f. Par conséquent la formule quantifiée prend la valeur f. Dans l'arithmétique $RM3^9$ le théorème de Fermat, dans la forme citée, n'est donc pas valide.

4. Applications et résultats

Le paradoxe de Russell est à l'origine de nombreuses tentatives de formaliser la théorie des ensembles de manière consistante. Ces tentatives impliquent en particulier l'abandon de l'axiome de compréhension non restreint:

$$\exists y \forall x (x \in y \leftrightarrow \varphi)$$

Celui-ci est pourtant souhaité, si l'on veut utiliser cette théorie pour décrire la pratique scientifique courante. Il garantit en particulier que tout prédicat définissable dans la théorie possède une extension.

Les restrictions proposées de l'axiome de compréhension semblent artificielles et souvent ad hoc. Ni la théorie des types, qui restreint le langage de manière peu naturelle, ni l'axiomatisation de Zermelo qui limite l'application de cet axiome à des ensembles déjà existants ne sont très satisfaisants du point de vue de la pertinence de la théorie résultante. L'approche de Gödel-Bernays-von Neumann quant à elle introduit une distinction peu intuitive entre *classe* et *ensemble*. La consistance d'aucune de ces théories restreintes des ensembles n'a pu être établie.

L'approche paraconsistante, qui consiste à trouver une logique sous-jacente, capable de résister à certains paradoxes, constitue une démarche intéressante. Bien sûr, là aussi il y a un prix à payer. La logique qui en résulte, n'a pas toujours des contours familiers, comme nous l'avons vu. Par contre elle permet de retrouver l'axiome de compréhension dans sa forme générale. Dans un travail récent, Brady (1989) a en effet montré la non-trivialité, ou la consistance absolue, de ce qu'il appelle une théorie dialectale des ensembles. Cette théorie est basée sur une logique paraconsistante et possède un axiome de compréhension sans restrictions en plus de l'axiome d'extensionnalité. La logique sur laquelle s'appuie cette théorie est la logique DSQ, une logique relevante de type RM avec des opérateurs modaux. Brady a démontré la non-trivialité de sa théorie en construisant un modèle extensionnel complexe MD. Ce modèle est construit à l'aide de suites transfinies de modèles du type RM3. Dans ce modèle, les axiomes de DSQ ainsi que l'axiome de compréhension sont vérifiés, alors que d'autres formules ne le sont pas.

Une autre application importante de la logique paraconsistante se situe dans le domaine de la sémantique naïve. En effet, une théorie dans laquelle on veut exprimer la notion de vérité par un prédicat appartenant à son langage, est inconsistante. Appelons V le prédicat qui représente la vérité dans la théorie et supposons que ce prédicat satisfasse le schéma de Tarski: $\varphi \leftrightarrow V\hat{\varphi}$ pour toute formule fermée². Il est possible de définir

2 $\hat{\varphi}$ est la représentation du nombre de Gödel de φ .

ensuite une formule ayant la propriété $\theta \leftrightarrow (V\hat{\theta} \rightarrow \psi)$ en appliquant le théorème du point fixe. Une telle théorie est triviale comme le montre le paradoxe de Curry:

- | | | |
|-----|--|--------------|
| (1) | $V\hat{\theta} \leftrightarrow (V\hat{\theta} \rightarrow \psi)$ | Tarski |
| (2) | $V\hat{\theta} \rightarrow (V\hat{\theta} \rightarrow \psi)$ | |
| (3) | $V\hat{\theta} \rightarrow \psi$ | absorption |
| (4) | $V\hat{\theta}$ | modus ponens |
| (5) | ψ | modus ponens |

Les principes logiques en cause ici sont l'absorption et le modus ponens, tous deux principes valides de la logique classique. Une sémantique naïve doit donc être basée sur une logique différente.

Il existe effectivement un modèle non trivial de sémantique naïve, construit sur une logique paraconsistante. Dowden le décrit dans un article de 1984. La démarche de Dowden est intéressante, parce qu'elle combine des idées de la logique LP, présentée plus haut, avec le point de vue de Kripke [Kr] sur le traitement des paradoxes.

On peut définir le paradoxe du menteur par la formule

$$\mu: \exists x(Cx \wedge \neg Vx)$$

dans laquelle C est le prédicat qui s'applique uniquement au nombre de Gödel de μ . L'existence de cette formule est garantie par le théorème du point fixe. Le modèle de Dowden, qu'il appelle J, est construit à partir du modèle K de Kripke en donnant l'extension et l'anti-extension du prédicat V.

$$V^{J+} = \{\hat{\varphi} \mid K \not\models \neg\varphi\}$$

$$V^{J-} = \{\hat{\varphi} \mid K \not\models \varphi\}$$

Dans ce modèle J, on peut montrer les résultats suivants:

- Pour tout φ , ou $J \models \varphi$, ou $J \models \neg\varphi$
- $J \models \mu$, et $J \models \neg\mu$

Si dans le modèle K de Kripke, la vérité est sous-déterminée, – le menteur n'est ni vrai ni faux –, dans le modèle J de Dowden, la vérité est sur-déterminée pour les formules paradoxales comme dans l'exemple du menteur, μ . Dans ce dernier modèle, Dowden a pu prouver le résultat important: dans J, aucune formule classique, c'est-à-dire ne comprenant pas le prédicat V, n'est à la fois vraie et fausse. Ceci montre d'une part sa non-trivialité, mais aussi, d'autre part, qu'il s'agit bien ici d'une extension de la situation classique.

4. Conclusion

Les logiques classique et intuitionniste traitent le problème des contradictions par le moyen radical du principe (ECQ). Celui-ci fait simplement exploser toute théorie inconsistante, avec le désavantage qu'il ne permet plus de faire la différence entre consistance et consistance absolue (non-trivialité). Or nous avons donné des exemples de théories situées précisément dans cette zone entre inconsistance et trivialité. Il existe des modèles de l'arithmétique inconsistants, pour lesquels il existe des preuves finies de consistance absolue! Cela oblige à jeter un regard nouveau sur les théorèmes d'incomplétude de Gödel. Les limitations de la formalisation paraissent moins dramatiques si l'on considère que des théories inconsistantes peuvent être d'un certain intérêt, à condition bien sûr qu'elles ne soient pas triviales. Et certaines de ces théories présentent un intérêt mathématique évident, ce qui n'est pas le moindre des arguments pour en faire l'étude.

Il y a aussi de bonnes raisons pratiques pour ne pas vouloir éliminer toute contradiction de nos considérations logiques. Les bases de données fournissent un bon exemple de situation où le principe «ex contradictione quodlibet» ne devrait pas s'appliquer. Il se peut en effet que des informations soient introduites dont le caractère contradictoire n'est pas immédiatement apparent, mais ne se révèle que plus tard à la lumière de nouvelles données. Cela ne doit avoir pour conséquence que toute affirmation figure désormais dans la base: celle-ci ne doit pas explo-

ser. Cela implique que le mécanisme déductif sous-jacent doit être non explosif, c'est-à-dire paraconsistant.

En ce qui concerne la théorie des ensembles, il est établi que l'on peut en construire une version possédant l'axiome de compréhension sans restrictions. Il existe même des preuves finies de la non-trivialité d'une telle théorie. Il en va de même pour une autre théorie qu'on aimerait conserver pour sa simplicité: la sémantique naïve. Là encore il existe une, et même plusieurs versions absolument consistantes. Le prix à payer en adoptant une telle théorie est celui de l'abandon de certains principes logiques habituels, le *modus ponens*, ou le syllogisme disjonctif par exemple selon les cas. Mais la théorie des types et les versions axiomatisées de la théorie des ensembles de Zermelo, Fraenkel et Gödel, Bernays, von Neumann, qui sont des théories construites pour éliminer le paradoxe de Russell, ne nous obligent-elles pas à abandonner une part de notre intuition pour des limitations souvent *ad hoc*? Une autre possibilité est celle d'opter pour une conditionnelle non extensionnelle, c'est-à-dire qui ne peut pas être définie dans le langage de la négation et de la disjonction. Un tel exemple est proposé par la logique relevante R.

Il est temps de songer à donner un statut logique nouveau à la contradiction. Son rôle de détonateur dans le cadre d'une logique explosive, classique ou intuitionniste, ne semble plus adéquat, pour toutes les raisons indiquées dans ce bref exposé, et bien d'autres qu'on trouvera dans la littérature (cf. Priest 1989). Ce statut nouveau, ce sont les différentes logiques paraconsistantes qui permettent de le lui attribuer. Et si ces logiques dérangent nos habitudes, elles engendrent des théories qui, dans la plupart des cas, vont au-delà des théories existantes dont elles constituent des généralisations. Et en plus de leur intérêt intrinsèque elles nous ouvrent d'autres horizons, aussi bien épistémologiques que mathématiques.

Institut de mathématiques appliquées
BFSH 2
1015 LAUSANNE

Bibliographie

- ANDERSON A.R. & BELNAP N.D. (1975). *The Logic of Relevance and Necessity*. Princeton: University Press, vol. 1.
- BRADY R.T. (1989). The non-triviality of dialectical set theory. In: Priest, Routley, Norman (1989).
- DA COSTA N. (1974). On the theory of inconsistent formal logics. *Notre Dame Journal of Formal Logic* 15, 497-510.
- DOWDEN B. (1984). Accepting inconsistencies from the paradoxes. *Journal of Philosophical Logic* 13.
- JASKOWSKI S. (1969). Propositional calculus for contradictory deductive systems. *Studia Logica* 24.
- LUKASIEWICZ J. (1971). On the principle of contradiction in Aristotle. *Review of Metaphysics* 24.
- MEYER R.K. & MORTENSEN C. (1984). Inconsistent models for relevant arithmetics. *The Journal of Symbolic Logic* 49, 918-929.
- PRIEST G. (1979). Logic of paradox. *Journal of Philosophical Logic* 8.
- PRIEST G. (1987). *In Contradiction: A Study of the Transconsistent*. Dordrecht: Nijhoff.
- PRIEST G., ROUTLEY R., NORMAN J. (eds) (1989). *Paraconsistent Logics*. München: Philosophia.
- SMILEY T. (1959). Entailment and deducibility. *Proc. of the Aristotelian Society*, vol. 59.

QUELQUES ASPECTS ÉPISTÉMOLOGIQUES DE LA THÉORIE CONSTRUCTIVE DES TYPES

Giovanni SOMMARUGA-ROSOLEMOS

Frege, le premier Wittgenstein et après eux toute l'approche métamathématique de la logique ont banni toute odeur épistémologique de la logique. A peu près 100 ans après Frege nous trouvons de nouveau des concepts épistémologiques en logique, à savoir dans la théorie constructive des types. Le but de mon exposé est de rendre attentif au fait qu'aux temps de Frege et du premier Wittgenstein, il y avait déjà une approche philosophique à la logique qui, bien que partageant l'antipsychologisme de ses deux logiciens, n'a pas supprimé le côté épistémologique de la logique; et que cette approche a influencé subrepticement l'émergence des concepts épistémiques en théorie constructive des types. Plus concrètement, cet exposé traitera des points suivants:

- (1) de quelques concepts épistémiques et non épistémiques fondamentaux de la théorie constructive des types (ou de sa métathéorie);
- (2) de la bilatéralité de la logique selon Husserl (dans *Formale und transzendente Logik*);
- (3) de la réfutation du psychologisme par Husserl (dans *Formale und transzendente Logik*);
- (4) d'une approche théorique possible des concepts épistémiques fondamentaux de la théorie constructive des types de Martin-Löf: la psychologie pure ou descriptive de Husserl;
- (5) d'une tentative d'interprétation phénoménologique (ou psychologique pure) des concepts fondamentaux de la théorie constructive des types mentionnés au point (1).

1. Dans une conférence intitulée «Truth and Knowability: On the Principles C and K of Michael Dummett» que Martin-Löf a donné en septembre 1995 en Sicile, il part de l'observation d'une certaine étrangeté du principe K de Dummett et se propose de l'améliorer de telle façon qu'il devienne acceptable. Le clou de cette amélioration est la distinction entre le concept de vérité d'une proposition et le concept de correction (ou vérité) d'un jugement. En ce qui concerne cette distinction, je suivrai l'exposé de Martin-Löf.

Il est presque impossible d'expliquer ces deux concepts de vérité d'une proposition et de correction d'un jugement isolément: ce sont deux concepts qui s'insèrent dans un certain système conceptuel de la théorie des types constructive. Si on veut les clarifier, on est obligé de présenter ce système conceptuel et de voir, comment ces différentes pièces s'adaptent l'une à l'autre et quels rôles ils jouent à l'intérieur de ce système.

Les éléments de base de ce système conceptuel sont donnés dans le tableau suivant:

concepts non épistémiques	concepts épistémiques
proposition	jugement
(objet de) preuve d'une proposition	démonstration (preuve) d'un jugement
vérité d'une proposition	correction (vérité) d'un jugement

La distinction la plus fondamentale qui est à la base de ce système conceptuel est celle entre proposition et jugement. Avant de répondre aux questions: Qu'est-ce qu'une proposition? Qu'est-ce qu'un jugement?, je ferai quelques remarques préliminaires sur ces deux concepts.

Les propositions sont des choses qui sont considérées parfois comme vraies, et parfois comme fausses. Ce sont aussi des choses sur lesquelles les opérateurs logiques opèrent: les

connecteurs opèrent sur les propositions et les quantificateurs sur les fonctions propositionnelles.

Les jugements sont des choses que nous démontrons: à chaque pas d'une chaîne de raisonnement ou de démonstration, nous procédons de jugements démontrés préalablement à un nouveau jugement qui est évident en raison des jugements précédents. Les formes de jugement diffèrent totalement des opérations logiques. Il y a par exemple la forme de jugement affirmative « A est vraie» où A est une proposition, la forme de jugement négative « A est fausse» où A est une proposition; il y a les formes « A est une proposition», et $a:A$, c'est-à-dire « a est une preuve de la proposition A », etc. De même qu'on ne peut pas limiter par avance les opérations logiques, on ne peut pas limiter par avance les formes de jugement non plus: de fait, en théorie constructive des types il y a des formes de jugement qui disent que quelque chose est un type, ou que quelque chose est un objet d'un certain type, ou encore que deux objets d'un certain type sont identiques par définition, etc.

Mais revenons à la question: Qu'est-ce qu'un jugement? Martin-Löf donne la réponse suivante: un jugement est défini en déterminant ce qu'on doit connaître pour avoir le droit de le faire (ou plus brièvement: un jugement est défini par la connaissance qu'on revendique d'avoir quand on fait ce jugement).

Le concept de démonstration est défini en disant qu'une démonstration est ce qui rend un jugement connu ou évident. Dire qu'un jugement est évident, signifie que le jugement a été démontré. Une démonstration d'un jugement est un acte de compréhension, de saisie, d'appréhension d'un jugement. Le concept de correction (vérité) d'un jugement est finalement défini comme suit: Un jugement est correct s'il peut être connu ou évident. C'est-à-dire un jugement est correct, s'il peut être démontré.

Remarquons que évidence et correction se comportent l'un envers l'autre comme l'actualité envers la potentialité. Et donc l'évidence d'un jugement implique sa correction. C'est tout pour ce qui concerne le côté droit du tableau.

Passons maintenant à l'autre question: Qu'est-ce qu'une proposition? Dans une théorie constructive ou vérificationniste de

la signification une proposition est définie par ses conditions de preuve ou par ses conditions de vérification. Ces conditions établissent à quoi ressemble une preuve (un objet de preuve) ou une vérification de la proposition.

En ce qui concerne le concept de preuve (ou objet de preuve) d'une proposition, on doit distinguer d'un côté les preuves ou vérifications de formes telles qu'elles entrent dans les explications de la signification des différents opérateurs logiques, appelées des preuves canoniques, et de l'autre les preuves ou vérifications arbitraires qui peuvent être calculées en des preuves canoniques et qui sont appelées des preuves non canoniques. Les preuves canoniques sont données par les explications bien connues de la signification des opérateurs logiques, les explications dites de Brouwer-Heyting-Kolmogorov. Finalement, A est vraie, où A est une proposition, est défini par: il existe une preuve de A , ou par, une preuve de A peut être donnée, ou encore, A peut être prouvée.

Il ne peut être question d'un jugement évident en lui-même, c'est-à-dire indépendamment de l'activité cognitive de quelqu'un. Parler de jugements évidents en eux-mêmes est aussi absurde que de parler de jugements connus de personne, mais qui sont connus en eux-mêmes. «Être évident» est une expression aussi relationnelle que «être connu», et ces expressions doivent être comprises comme «être évident pour quelqu'un» et «être connu par quelqu'un». Brièvement, il n'y a pas d'évidence hors de l'expérience de cette évidence par quelqu'un. C'est un principe de l'intuitionnisme exprimé par Brouwer qui emploie «vérité» plutôt que «jugement évident» comme suit: Il n'y a pas de vérités non expérimentées en mathématiques.

Or, un jugement n'est rien d'autre qu'une revendication de connaissance, et une revendication de connaissance est toujours une revendication de connaissance faite par quelqu'un. Une démonstration d'un jugement n'est rien d'autre qu'un acte de production d'évidence en vue de cette revendication de connaissance, qui de cette manière rend évidente la connaissance revendiquée, c'est-à-dire le jugement. Il va de soi qu'il ne s'agit pas d'un acte en lui-même, mais de l'acte de quelqu'un. Et la correction d'un jugement signifie que l'évidence peut être fournie pour

une revendication de connaissance, que cette revendication de connaissance peut être démontrée ou justifiée par quelqu'un.

Est-ce que la théorie constructive des types provoque avec cela le fantôme du psychologisme? Notons que l'intuitionnisme de Brouwer possède bien un penchant psychologue ou au moins, il ne s'est jamais défendu clairement contre tout malentendu psychologue. Mais la même chose ne vaut pas pour la théorie constructive des types de Martin-Löf.

2. La bilatéralité de la logique par rapport à la théorie constructive des types a déjà été remarquée par Husserl en 1929 dans *Formale und transzendente Logik*. Il observe que la logique traite des effectuations de la raison dans les deux sens: premièrement au sens des activités et des habitus effectuant, et deuxièmement au sens des résultats effectués et dès lors permanents.

Au deuxième sens, le thème de la logique est les formes variées de formations du jugement et de la connaissance issues des activités pensantes d'une personne occupée à connaître quelque chose. Ce qui sort de ces activités pensantes n'est pas sorti par hasard, mais est intenté par l'acte de penser. La personne pensante s'est dirigée spécialement sur cela et elle l'a objectivement devant soi. Des formes plus complexes de formation du jugement et de la connaissance qui dépassent de loin leurs sphères respectives de présence de conscience, restent néanmoins des formations pratiques auxquelles on peut toujours retourner et avec lesquelles on peut former maintes fois des nouvelles formations, c'est-à-dire de nouveaux concepts, jugements, inférences, démonstrations, théories.

Ces formations objectives n'ont pas seulement l'existence passagère de formations momentanées qui apparaissent dans et disparaissent d'un champ thématique. Elles ont aussi le sens d'une validité objective dans un sens particulier: elles restent identiques dans les répétitions, elles sont toujours à nouveau reconnues à la manière d'existants permanents et elles ont dans la forme documentaire une existence objective. Dans leur durée objective, elles peuvent être trouvées par tout le monde, elles

peuvent être identifiées intersubjectivement et elles existent même si personne ne pense à elles.

Selon le premier sens, le thème de la logique est d'une nature subjective. Il s'agit des formes subjectives par lesquelles la raison théorique produit ces formations. Ce qui est en discussion ici, c'est l'intentionnalité en tant que procès vivant d'où proviennent les formations objectives. Tant que l'intentionnalité dans son procès vivant produit ces formations, elle est «inconsciente», c'est-à-dire elle thématise ses formations, mais elle n'est pas elle-même un thème. Tant qu'elle n'est pas révélée par la réflexion, elle est cachée. La réflexion sur cette intentionnalité en procès la rend thématique, et la constitue comme objet théorique d'une recherche logique de l'orientation subjective. A chaque forme de formation logique objective correspond un système d'intentionnalité effectuant qui pourrait être appelé sa forme subjective. Mais tout cela n'est qu'un aspect d'une recherche logique d'orientation subjective.

Un autre aspect est constitué par l'effectuation subjective particulière qui fait que les formations du jugement et de la connaissance deviennent effectivement conscientes comme quelque chose d'objectif, comme quelque chose de valable d'une manière durable pour la subjectivité, et qu'elles assument le sens d'une objectivité idéale pour une communauté de personnes engagées à connaître.

Husserl conclut que cette bilatéralité de la logique a abouti à des difficultés extraordinaires. Presque tout ce qui concerne la signification de base de la logique, sa problématique et sa méthode ont été contaminés par les obscurités et confusions dues à cette objectivité à partir de l'effectuation subjective qui n'a jamais été comprise ni même mise en question de façon correcte. Nous voilà toujours occupés à comprendre et clarifier cette double face de la logique.

3. Ces deux aspects du premier côté de la logique ont été l'objet de critiques.

Au premier aspect du premier côté de la logique, on objecte une confusion de la logique avec la psychologie. La critique est

la suivante: Ce qui vaut pour toutes les sciences objectives est particulièrement évident dans le cas des sciences naturelles. Jamais personne ne considérerait les vécus de l'expérience de la nature et ceux de la pensée sur la nature comme constituant le domaine des sciences naturelles plutôt que la nature elle-même. La logique étant une science objective, la même chose vaut pour elle. Les vécus de l'expérience de la logique et ceux de la pensée sur la logique appartiennent évidemment à la psychologie.

Husserl répond que ce qui est normalement considéré comme le domaine d'une science objective n'est que sa première sphère thématique. Mais il y a une deuxième sphère thématique qui s'occupe elle de ce côté subjectif psychique de la recherche scientifique de cette science. Je cite Husserl:

La logique traditionnelle dans sa positivité naïve, avec sa manière de concevoir des vérités évidentes dans une immédiateté naïve, se manifeste comme une sorte d'enfantillage philosophique. ... Le caractère non philosophique de cette positivité ne consiste en rien d'autre qu'en ceci: les sciences, par non-compréhension de leurs propres effectuations en tant qu'effectuations d'une intentionnalité effectuant qui reste pour elles non thématique, sont incapables de clarifier le sens d'être authentique de leur domaine et des concepts qui les expriment.

Au deuxième aspect du premier côté de la logique, on a fait l'objection du psychologisme.

Husserl caractérise le psychologisme en général et le psychologisme logique en particulier comme suit. Le psychologisme (logique) consiste dans la psychologisation des formations irréelles de la signification (du genre logique). Que ces formations irréelles sont psychologisées signifie que leur sens d'une espèce d'objets avec une essence spécifique est nié en faveur des vécus subjectifs, c'est-à-dire en faveur de data dans la temporalité psychologique. En d'autres termes, les formations irréelles de la signification (du genre logique) sont identifiées avec des phénomènes de l'expérience interne. Ainsi les concepts, jugements, vérités, démonstrations, théories etc. sont des événements psychiques, c'est-à-dire véritablement des évé-

nements psychiques de jugement, des événements d'évidence ou des démonstrations subjectives-psychiques.

La motivation pour le psychologisme (logique) est la suivante. Le psychologisme (logique) est motivé par l'apparition interne des formations irréelles (du genre logique) des actes de la raison dans la conscience même de ces actes; ces formations se présentent dans la conscience comme quelque chose à l'intérieur d'elles. Par exemple, les formations de composantes de jugements et les formations de jugements entiers se présentent à l'intérieur de l'activité psychique qui se déploie comme des vécus de la conscience. Ces formations irréelles d'actes de la raison ne se présentent évidemment pas dans la conscience comme quelque chose à l'extérieur d'elle, puisqu'elles ne sont pas réelles, elles ne sont pas des objets dans l'espace.

Husserl propose une étude des évidences du réel et des évidences de l'irréel qui clarifierait l'uniformité générale des objectités en tant qu'objectités indépendamment du fait si elles sont réelles ou irréelles. Il commence avec l'observation que l'évidence des objets irréels est, en ce qui concerne leur effectuation, totalement analogue à l'évidence de l'expérience ordinaire dite externe ou interne. Mais c'est seulement à cette évidence de l'expérience ordinaire qu'on attribue, à cause de préjugés, l'effectuation de l'objectivation originale.

Or, l'évidence au sens phénoménologique réfère à l'effectuation intentionnelle de la donation des choses elles-mêmes. Elle est une forme distincte de l'intentionnalité (c'est-à-dire de la conscience de quelque chose) dans laquelle l'objectité qui est consciente, est consciente dans la manière d'être saisie en elle-même, d'être vu en tant que telle. L'évidence dans ce sens peut aussi être appelée la conscience originale.

L'intentionnalité en général et l'évidence, c'est-à-dire l'intentionnalité de la donation des choses elles-mêmes, sont deux concepts étroitement liés l'un à l'autre. Ceci est indiqué aussi par les deux propositions suivantes.

— La première appelée par Husserl légalité fondamentale de l'intentionnalité: toute conscience de quelque chose (intentionnalité) appartient *a priori* à une multiplicité ouverte et illimitée de modes possibles de conscience qui dans la forme

unitaire de la co-validité peuvent toujours être reliés en une conscience unique en tant que conscience de la même chose. A cette multiplicité appartiennent également et essentiellement les modes d'une conscience multiple d'évidence.

– La deuxième proposition concerne une propriété essentielle de toutes les évidences en général: dans la répétition des vécus subjectifs, dans la succession et la synthèse de différentes expériences de la même chose, les évidences rendent manifeste quelque chose qui est numériquement identique et non pas simplement semblable, elles rendent manifeste l'objet qui est saisi maintes fois dans le domaine de la conscience. (Plus spécifiquement: les processus de pensée respectifs à la formation de concepts, de jugements, d'inférences, etc. sont temporellement extérieurs les uns aux autres, ils sont individuellement différents et séparés, tandis que les concepts, les jugements, les inférences, etc. ne sont pas temporellement extérieurs les uns aux autres, mais ils sont atemporels, ils ne sont pas individuellement différents, mais généralement identiques, et non pas séparés, mais numériquement identiques. Et tout cela est révélé par les évidences des formations irréelles du genre logique).

Si on veut comprendre l'effectuation de la conscience et en particulier celle de l'évidence, il faut faire beaucoup de choses: il faut expliquer comment dans l'immanence des multiplicités du vécu se constituent leur se-diriger-vers et ce vers quoi ces multiplicités se dirigent; et il faut expliquer en quoi consiste, dans la sphère intuitive de l'expérience synthétique elle-même, l'objet transcendant. Où par objet transcendant est entendu le pôle d'identité immanent aux vécus particuliers, pôle qui dans son identité surpasse ces vécus particuliers.

Ces deux propositions concernant d'une part une légalité fondamentale de l'intentionnalité et d'autre part une propriété essentielle de toutes les évidences en général rendent suffisamment clair pourquoi le psychologisme (logique) est tout simplement faux:

Il n'y a aucun objet qui peut être identifié avec un certain nombre fini de modes donnés de conscience de cet objet ou avec un certain nombre fini de processus vécus de l'expérience, car, selon la légalité fondamentale de l'intentionnalité, la conscience

d'un tel objet appartient à une multiplicité ouverte et illimitée donc infinie de modes possibles de conscience. (Il est superflu de dire que l'objet est toutefois fondé ou partiellement constitué dans la donation actuelle et évidente des choses elles-mêmes.)

De plus, selon la propriété essentielle des évidences, il est évident que les objets formés dans les mêmes actes ou dans les actes semblables de conscience sont non seulement les mêmes ou semblables, mais ils sont des objets numériquement identiques. Ils sont un, contrairement à leurs apparences variées dans la conscience.

Husserl conclue contre le psychologisme (logique) que non seulement l'objet idéal n'est rien de réel, mais que l'objet réel est intrinsèquement idéal. Il raisonne comme suit. Dans le sens de tout objet saisissable par l'expérience, il y a une certaine idéalité, à savoir l'idéalité générale de toutes les unités intentionnelles à l'opposé des multiplicités d'expériences qui les constituent. Dans cette idéalité consiste la transcendance de toute sorte d'objets par rapport à la conscience de ces objets. Et même la transcendance du réel en tant que réel n'est qu'une formation particulière de l'idéalité, ou plutôt de ce que Husserl appelle l'irréalité psychique, c'est-à-dire de quelque chose qui se présente dans la sphère purement phénoménologique de la conscience.

4. Une approche théorique possible des concepts logiques épistémiques de Martin-Löf est donnée par la psychologie pure ou descriptive de Husserl.

La psychologie pure ou descriptive est l'étude de l'essence du psychique à la différence du physique. Et elle diffère de la psychologie empirique ou expérimentale en ce qu'elle n'est pas une science des faits, mais une science des essences. L'objet principal de la psychologie pure est sans aucun doute l'intentionnalité comme trait fondamental et essentiel de toute vie psychique. La connaissance en psychologie pure est acquise à travers l'évidence de l'expérience interne de la réflexion. Bref, la psychologie pure peut être caractérisée comme analyse purement réflexive des intentionnalités variées de la conscience aussi

bien que de ses implications et ses corrélatifs. Il faut souligner que la psychologie pure étudie le psychique toujours dans l'attitude naturelle et considère le psychique comme étant une couche d'être dépendante des hommes et des animaux.

La suggestion de traiter les concepts logiques épistémiques de la théorie constructive des types en psychologie pure ou descriptive (et pas en psychologie empirique ou expérimentale) a déjà été faite par Martin-Löf, et elle rejoint bien les idées de Husserl à ce propos: selon Husserl, la critique de la connaissance logique au sens subjectif – appelée la critique analytique de la connaissance logique par Husserl – s'occupe du premier aspect du premier côté de la logique et appartient nettement à la psychologie pure ou descriptive.

On peut noter encore que la critique de la connaissance logique au sens subjectif a un autre côté – appelé la critique transcendantale de la connaissance logique par Husserl – qui s'occupe de la subjectivité qui constitue, au sens phénoménologique du terme, le domaine de la logique ainsi que toute effectuation scientifique traitant de ce domaine, c'est-à-dire qui s'occupe du deuxième aspect du premier côté de la logique. Cette deuxième critique de la connaissance logique dépasse la psychologie pure ou descriptive et appartient à la phénoménologie transcendantale.

5. Dans son article de 1931 *Die intuitionistische Grundlegung der Mathematik* Heyting explique et défend le point de vue intuitionniste en proposant de considérer les propositions mathématiques comme expressions d'intentions au sens de la théorie de l'intentionnalité de Husserl. Ensuite il identifie les preuves comme constructions mentales, avec les remplissements d'intentions et il continue à décrire en ces termes la signification des constantes logiques du calcul propositionnel intuitionniste. Dans son cours de 1983 *On the Meanings of the Logical Constants and the Justifications of the Logical Laws* Martin-Löf remarque que Heyting n'a pas simplement emprunté ces termes à Husserl, mais qu'il les a appliqués adéquatement. Puisque Martin-Löf reprend lui-même, par

exemple dans son livre *Intuitionistic Type Theory*, l'interprétation de Heyting, une tentative sera faite ci-dessous de voir jusqu'à quel point un tout petit fragment de la théorie constructive des types pourra être interprété dans la théorie de l'intentionnalité de Husserl ou dans la phénoménologie tout court. Il s'agit de la forme du jugement $a:A$, où A est une proposition et a une preuve (un objet de preuve) de A .

Nous commençons par quelques distinctions et définitions préliminaires de la théorie de l'intentionnalité de Husserl:

- (a) La première distinction est celle entre intention de et intention que, où l'intention de est une intention d'objets et l'intention que est l'intention que quelque chose est le cas¹.
- (b) La seconde distinction est celle entre les intentions simples et les intentions complexes ou catégoriques. Les intentions simples ont une matière à un membre tandis que les intentions complexes ou catégoriques ont une matière à plusieurs membres. C'est-à-dire les intentions catégoriques sont des intentions qui rapportent des intentions multiples à une unité synthétique.
(Notons que selon Husserl toutes les intentions catégoriques sont en dernière analyse fondées dans les intentions simples.)
- (c) La prochaine distinction concerne l'intention vide et l'intention remplissante. L'intention vide n'est qu'une affirmation gratuite ou une présomption par rapport à l'objet intenté, et l'intention remplissante est une présentation pleine et intuitive de l'objet intenté.
- (d) Une intuition est une intention remplissante.
- (e) L'intention catégorique intuitive se distingue de l'intention catégorique signitive en ceci qu'elle est une intention catégorique remplissante, tandis que dans l'intention catégorique signitive, l'objet intenté n'est représenté que par des signes. Parmi les intentions catégoriques signitives, on peut distinguer entre les intentions de signification vides (c'est-

1 Ce n'est pas une distinction de Husserl, mais plutôt une distinction américaine de Charles Parsons et Mark Steiner, reprise par Richard Tieszen. Chez Husserl il n'y a que l'intention d'objets, où les objets peuvent être de différents types. Husserl considère les états de choses aussi comme étant des objets.

à-dire les intentions vides des significations d'énoncés) qui peuvent être remplies intuitivement, et les opérations de calcul mathématique qui ne peuvent pas être remplies intuitivement.

- (f) Une intuition catégorique est une intention catégorique remplissante ou intuitive.
- (g) La dernière distinction concerne l'intuition catégorique et l'intention de signification remplissante catégorique. Une intuition catégorique correspond à une intention catégorique vide, tandis qu'une intention de signification remplissante catégorique correspond à une intention catégorique signitive.

(Notons que selon Husserl les intuitions catégoriques et les intentions de signification remplissantes catégoriques sont parallèles, mais pas identiques; les premiers étant prélinguistiques tandis que les derniers sont évidemment linguistiques.)

Husserl caractérise le rapport entre vérité, énoncé, signification et intuition approximativement comme suit. La vérité d'un énoncé est dérivée de l'effectuation d'une intention de signification remplissante correspondante comme suit:

1) L'intention catégorique signitive correspond dans l'intention de signification remplie catégorique de manière univoque à l'intention de signification remplissante catégorique correspondante; et 2) l'énoncé est une expression linguistique univoque de l'intention catégorique signitive correspondante.

Les intentions vides, remplissantes ou remplies, et catégoriques sont celles qui intéresseront le plus par la suite. Nous allons essayer maintenant d'interpréter phénoménologiquement le jugement $a:A$, ses composantes et quelques concepts liés à ce jugement particulier et aux jugements de la théorie constructive des types en général. Suivant Heyting et Martin-Löf, une proposition est toujours interprétée comme intention, et suivant Martin-Löf un objet de preuve est interprété comme méthode de remplissement d'une intention, mais plus précisément comme suit:

proposition = intention de signification vide catégorique de

- objet de preuve = méthode de remplissement d'une intention de signification vide catégorique de
- vérité = existence d'une méthode de remplissement d'une intention de signification vide catégorique de
- jugement = intention de signification vide catégorique que (c'est une intention de niveau supérieur par rapport à l'intention qui interprète une proposition)
- démonstration = acte d'intention de signification remplie catégorique que
- théorème = produit d'un acte d'intention de signification remplie catégorique que (théorème = jugement démontré)
- évidence = expérience effective d'un acte d'intention de signification remplie catégorique que
- correction = existence (et possibilité) d' (un acte d') une intention de signification remplie catégorique que

J'aimerais terminer cette tentative de redonner l'intuition (c'est-à-dire l'intention de signification remplie) à l'intuitionnisme (comme Tieszen a appelé une de ces tentatives) et en même temps cet exposé avec quelques remarques et observations:

- i) Dans l'intuitionnisme de Heyting et dans la théorie constructive des types de Martin-Löf apparaît une nouvelle espèce d'entités (nouvelle par rapport à la phénoménologie), c'est-à-dire l'objet de preuve. Mais il y a aussi une espèce d'intention en phénoménologie qui disparaît dans l'intuitionnisme, c'est-à-dire l'intention remplissante. L'idée de l'intuitionnisme est que l'exécution de la méthode de remplissement d'une intention de signification vide catégorique donne une intention de signification remplie catégorique. On voit donc que l'intention remplissante est dans quelque manière contenue dans l'interprétation de l'objet de preuve.
- ii) Dans la théorie constructive des types, la correction et l'évidence sont attribuées aux jugements, tandis que, selon cette interprétation-ci, elles sont d'abord attribuées aux démonstrations et seulement secondairement aux jugements. En

plus, la correction semble coïncider essentiellement avec la vérité au sens phénoménologique du terme (ou plus précisément: avec la vérité dans un des sens phénoménologiques du terme).

- iii) L'identification de la démonstration avec l'acte d'intention de signification remplie catégorique telle que je viens de la proposer n'apparaît, au sens strict, pas appropriée: elle ne vaut réellement que pour la démonstration directe, mais pas pour la démonstration indirecte ou la justification. Il y a donc nécessité d'une différenciation ultérieure, dans laquelle des concepts tels que inférence ou validité jouent un rôle important.
- iv) Il est possible que d'autres formes de jugements de la théorie constructive des types liées à la forme de jugement que je viens d'interpréter, ne sont pas faciles à interpréter phénoménologiquement: par exemple les formes $A:\text{prop}$, $A=B:\text{prop}$ ou $a=b:A$.
- v) Ainsi nous revenons au début de cet exposé, c'est-à-dire au tableau des concepts logiques épistémiques et non épistémiques de Martin-Löf. Il semble bien que l'interprétation phénoménologique de ces concepts fondamentaux de la théorie constructive des types annule la distinction de Martin-Löf entre les concepts épistémiques et ceux non épistémiques. Si nous adoptons une telle interprétation phénoménologique, alors tous ces concepts fondamentaux deviennent épistémiques.
- vi) Si cette interprétation phénoménologique tient la route, alors elle ne sert qu'en tant que squelette qui doit encore être complétée par la chaire phénoménologique.

*Philosophisches Seminar I
Albert-Ludwigs-Universität
Freiburg im Breisgau (D)*

Bibliographie

- BERNET R., KERN I., MARBACH E. (1989). *Edmund Husserl. Darstellung seines Denkens*. Hamburg: F. Meiner.
- HEYTING A. (1931). Die intuitionistische Grundlegung der Mathematik. *Erkenntnis* 2, 106-115.
- HUSSERL E. (1992). *Formale und transzendente Logik*. Vol. 7 de *Gesammelte Schriften*. Ed. par E. Ströker. Hamburg: F. Meiner.
- MARTIN-LÖF P. (1984). *Intuitionistic Type Theory*. Napoli: Bibliopolis.
- MARTIN-LÖF P. (1984). On the meanings of the logical constants and the justifications of the logical laws. In: C Bernardi et P. Pagli (edd.), *Atti degli incontri di logica matematica*. II. Siena: Scuola di Specializzazione in Logica Matematica, Dip. di Matematica, Università di Siena, 203-281.
- MARTIN-LÖF P. (1995). Truth and knowability: On the principle C and K of Michael Dummett. Conférence donnée au Deuxième Colloque International de Philosophie de Mussomeli, Sicile, 14-20 septembre 1995. Ms.
- TIESZEN R. (1989). *Mathematical Intuition*. Dordrecht: Kluwer.

Travaux du Centre de Recherches Sémiologiques

Liste des numéros parus

- *1. G. Vignaux: La nouvelle rhétorique. Revue critique et perspectives d'application. 1969-70.
- *2. G. Vignaux: L'argumentation antique: Aristote. Janvier 1970.
- *3. M.-J. Borel: Pour définir l'argumentation. 1969-70.
- *4. F. Bugniet: Remarques sur les notions d'assertion linguistiques et de proposition logique. Septembre 1970.
- *5. M.-J. Borel/G. Vignaux: L'étude de l'argumentation. Séminaire 1969-70.
- *6. G. Vignaux: L'argumentation: bibliographie sélective. Janvier 1971.
- *7. J.-B. Grize: Logique de l'argumentation et discours argumentatif. Mai 1971.
- *8. J.-L. Galay: La rhétorique du discours de philosophie systématique. Essais d'analyse. Mars 1971.
- *9. C. Morier: Charles Sanders Peirce et la sémiotique. Mars 1971.
- *10. G. Vignaux: L'argumentation et le résumé. Mars 1971.
- *11. C. Gillieron/C. Bonnet: Peut-on définir l'argumentation? Avril 1971.
- *12. J.-B. Grize: Notes sur l'ontologie et la méréologie de S. Lesniewski. Mars 1972.
- *13. M. Hirsbrunner/P. Fiala: Les limites d'une théorie saussurienne du discours et leurs effets dans la recherche sur l'argumentation. Avril 1972.
- *14. C. Gillieron/A.-M. Badonnel/J.-P. Iacazzi: Les recherches psychologiques et psycholinguistiques sur la négation et les relations d'opposition. Mai 1972.
- *15. J.-L. Galay: Esquisse pour une théorie figurale du discours. Septembre 1972.
- *16. Y. Oppel: Sémiotique littéraire, à propos de la coordination, répétition et opposition dans un texte littéraire. Mai 1973.
- *17. P. Fiala/C. Ridoux: Essai de pratique sémiotique. Juin 1973.
- *18. M. Hirsbrunner: Pour une critique de la sémiotique de Roland Barthes. Juillet 1973.
- *19. Y. Oppel: Colloque sur l'analyse du discours «Divergences et convergences». Février 1974.
- *20. (Collectif): Logique, argumentation, discours (LAD). Recherche. Septembre 1974.
- *21. (Collectif): Logique, argumentation, discours (LAD). Recherche. Septembre 1974.

- *22. A.-F. Schmid: Philosophie et sciences chez Henri Poincaré: lecture philosophique. Octobre 1974.
- *23. M.-J. Borel: Schématisation discursive et énonciation. Arguments théoriques et approche descriptive (LAD I). Octobre 1975.
- *24. A. Licitra: Les relations interpropositionnelles. Huit types d'après R. Longacre (LAD I). Octobre 1975.
- *25. (Collectif): Discours et structures sociales. Janvier 1977.
- *26. M. Ebel: Langage, histoire, action: les recherches de Jean Pierre Faye. Septembre 1975.
- *27. M. Ebel/P. Fiala: Recherches sur les discours xénophobes I. Juillet 1977.
- *28. M. Ebel/P. Fiala: Recherches sur les discours xénophobes II. Juillet 1977.
- *29. J.-B. Grize: Matériaux pour une logique naturelle (LAD I). Mai 1976.
- *30. D. Miéville/M.-J. Borel/A. Licitra: Discours et analogie (LAD II). Mai 1977.
- *31. J. Moeschler: Contribution linguistique à une sémiotique du cinéma. Mai 1977.
- *32. A. Lecomte: Paraphrase et thématization. Essais d'analyse logique. Décembre 1978.
- *33. (Collectif): Langue et discours I. Colloque Besançon-Neuchâtel, 2-4 octobre 1978. Mars 1978.
- *34. (Collectif): Langue et discours II. Colloque Besançon-Neuchâtel, 2-4 octobre 1978. Mars 1978.
- *35. P. Baldi/J. Moeschler: Comment contrôler le discours: interaction et réfutation dans le débat Giscard-Mitterrand (1974). Juillet 1979.
- *36. (Collectif): Quelques réflexions sur l'explication. Février 1980.
- 37. M. Sanchez-Mazas: Traduction arithmétique des graphes et des relations binaires et applications logiques et informatiques. Juin 1981.
- *38. (Collectif): Le discours explicatif I. Septembre 1981.
- *39. (Collectif): Le discours explicatif II. Septembre 1981.
- 40. C. Wülser: Actes de langage explicatifs. Février 1982.
- *41. (Collectif): Logique naturelle du raisonnement. Avril 1982.
- *42. (Collectif): Linguistique et sémiologie I. Colloque Besançon-Neuchâtel, 5-6 octobre 1981. Juillet 1982.
- 43. (Collectif): Linguistique et sémiologie II. Colloque Besançon-Neuchâtel, 5-6 octobre 1981. Juillet 1982.
- *44. (Collectif): Raisonnements et raisons. Avril 1983.
- 45. F. Albera: Problèmes de l'énonciation au cinéma. Février 1984.
- 46. (Collectif): Construction et transformations des objets du discours I. Colloque Besançon-Neuchâtel, 3-4 octobre 1983. Mars 1984.

47. (Collectif): Construction et transformations des objets du discours II. Colloque Besançon-Neuchâtel, 3-4 octobre 1983. Mars 1984.
- *48. (Collectif): Analyse de texte assistée par ordinateur. Utilisation du logiciel DEREDEC. Janvier 1985.
- *49. (Collectif): Problèmes et méthodes d'une analyse de texte articulant organisation cognitive, argumentation et représentations sociales. Juin 1985
50. (Collectif): Actes du colloque «Dialogisme et Polyphonie», 27/28 septembre 1985. Avril 1986.
- *51. (Collectif): Le discours descriptif. Du texte aux objets de connaissance I. Juillet 1986.
- *52. (Collectif): Le discours descriptif. Du texte aux objets de connaissance II. Juillet 1986.
- *53. (Collectif): La référence. Points de vue linguistique et logique. Mars 1987.
54. D. Apothéloz/J.-B. Grize: Langage, processus cognitif et genèse de la communication. Septembre 1987.
- *55. (Collectif): La schématisation descriptive. Types textuels, formes et fonctions discursives. Janvier 1988.
56. D. Miéville/R. Martin/A. Culioli/G.G. Granger/C. Gillieron/G. Seel/J. Molino/L. Frey/J.-B. Grize: La négation sous divers aspects. Actes du colloque, Neuchâtel 22-23 octobre 1987. Septembre 1988.
- *57. D. Miéville/D. Apothéloz/P.-Y. Brandt/ G. Quiroz/J.-B. Grize: La négation. Contre-argumentation et contradiction. Septembre 1989.
- *58. M. Charolles: De l'art de nager et des différentes manières d'en parler. Septembre 1990.
- *59. D. Miéville/M.-J. Borel/J.-P. Desclés/J. Gasser/P.-Y. Brandt; D. Apothéloz/J. Moeschler/J. Jayez/M.F. Blès: La négation. Le rôle de la négation dans l'argumentation et le raisonnement. Actes du colloque, Neuchâtel 11-12 octobre 1990. Septembre 1991.
60. D. Miéville/D. Apothéloz/P.Y. Brandt: Les organisations raisonnées. Analyse de l'articulation de séquences discursives. Juin 1992.
61. D. Miéville/M. Chavaz/E. Gattico: Relations formelles et non formelles. Septembre 1993.
62. D. Miéville/ C. Tiercelin/ G. Heinzmann/ G. Deledalle/ J. Gasser/ N. Everaert-Desmedt/ J. Réthoré/ M. Balat/ J.-P. Kaminker: Charles Sanders Peirce. Apports récents et perspectives en épistémologie, sémiologie, logique. Actes du colloque, Neuchâtel, 16-17 avril 1993. Avril 1994.

63. D. Miéville, J.-P. Desclés, P. Engel, J.-L. Gardies, J.-C. Gardin, J. Gasser, J.-B. Grize, F. Nef, *Raisonnement et calcul*. Actes du colloque, Neuchâtel, 24-25 juin 1994. Septembre 1995.
64. D. Apothéloz, U. Bähler, M. Schulz (eds), *Analyser le musée*. Actes du colloque international organisé par l'Association Suisse de Sémiotique (ASS/SGS), Lausanne 21-22 avril 1995. Août 1996.
65. D. Miéville (éd.) *Temps, logique et langage*. Actes du Symposaium tenu lors du colloque internation «Penser le temps», Neuchâtel, 8-10 septembre 1996. Avril 1997.

Les titres précédés d'un astérisque sont épuisés.

Les publications disponibles peuvent être obtenues auprès du Centre de Recherches Sémiologiques au prix de **Fr.s 10.-**. Dès le n° 59 **Fr.s. 15.-** (TVA comprise).

Travaux de logique

Liste des numéros parus

1. Denis Miéville: Introduction à la théorie des systèmes formels. Première partie. Septembre 1985 (épuisé).
2. Denis Miéville: Introduction à la théorie des systèmes formels. Deuxième partie. Janvier 1987 (épuisé).
3. James Gasser: La syllogistique d'Aristote à nos jours. Juin 1987.
4. Denis Miéville: Introduction à la théorie des systèmes formels. Première partie. Avril 1991 (réédition du n° 1; épuisé).
5. Denis Miéville: Introduction à la théorie des systèmes formels. Deuxième partie. Avril 1991 (réédition du n° 2; épuisé).
6. Denis Miéville: La négation, une étude logique. Mai 1991 (épuisé).
7. Denis Miéville (éd.): Kurt Gödel. Actes du colloque, Neuchâtel, 13 et 14 juin 1991. Septembre 1992.
8. James Gasser: Introduction à la logique des relations de C.S. Peirce. Novembre 1993.
9. D. Miéville, P. Joray, D. Stauffer, N. Gessler: Etudes logiques. Port-Royal: une logique des idées. L'avènement de la théorie sémantique de la vérité de Tarski. George Boole et l'algèbre de la logique. Décembre 1994.
10. D. Bourquin, P. Joray, N. Gessler, D. Miéville: Analyse catégorielle et logique. Octobre 1996.

Ces publications peuvent être obtenues auprès du Centre de Recherches Sémiologiques au prix de **Fr.s. 10.-** ; dès le n° 7 **Fr.s. 15.-** (TVA comprise).

Couverture: Atelier Seth, Peseux
Impression: Zentralstelle der Studentenschaft der Universität Zürich

