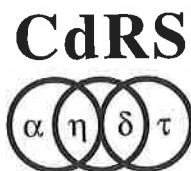


Centre de Recherches Sémiologiques
Travaux de logique
N° 14 — août 2001



MÉRÉOLOGIE ET MODALITÉS.
Aspects critiques et développements.
Actes du colloque - Neuchâtel, 20-21 octobre 2000.
Textes réunis sous la direction de D. Miéville



Université de Neuchâtel

Comité de lecture

Jean-Pierre DESCLÉS, Paris

Gerhard HEINZMAN, Nancy

Frédéric NEF, Rennes

Denis MIÉVILLE, Neuchâtel

Denis VERNANT, Grenoble

Henri VOLKEN, Lausanne

Centre de Recherches Sémiologiques
Université de Neuchâtel
Espace Louis-Agassiz 1
CH-2000 Neuchâtel (Switzerland)

© 2001 by Centre de Recherches Sémiologiques. Tous droits réservés

PRÉAMBULE

La méréologie est la théorie générale des parties, des tous et de leurs relations. La théorie de l'individuation est l'étude de ce qui rend un individu distinct de tout autre individu et identique à soi-même (et l'étude d'une série de notions liées, parmi lesquelles celles de substance, changement, identité à travers le temps, à travers des mondes possibles, etc.).

Au cours de ces dernières années, la méréologie et la théorie de l'individuation ont montré un degré significatif de convergence. Au sein de la méréologie, une alternative fondamentale existe entre les adversaires et les défenseurs du principe d'extensionnalité méréologique. On retrouve cette même alternative, dans la théorie de l'individuation, entre l'idée que les individus matériels sont des processus quadridimensionnels et l'idée qu'ils sont des substances tridimensionnelles étendues dans les trois dimensions de l'espace, mais pas dans celle du temps.

Par rapport à ce double choix, les chercheurs du Centre de Recherches Sémiologiques et de l'Institut de logique de l'Université de Neuchâtel ont voulu comprendre les relations entre ces deux orientations. Ils ont plus particulièrement étudié la relation méréologique à l'œuvre dans certains objets inférentiels, puis les opérations logiques de subordination considérées à l'image d'opérations constitutives de parties. Par ailleurs, ils ont montré de quelle manière développer un système logique permettant d'opérer sur des entités individuelles issues de la catégorie des propriétés et des relations. Ils ont enfin démontré l'universalité du principe d'extensionnalité méréologique.

À l'issue de cette recherche intitulée «Méréologie et individuation» financée par le Fonds National Suisse de la Recherche Scientifique (subside n° 1114-052326,97/1), il était indispensable d'exposer à la critique de logiciens et de philosophes les résultats mis en évidence. En s'exposant ainsi, ils ont suscité de très riches discussions dont les développements sont présentés ici même.

Je tiens par ces lignes à remercier tous les intervenants qui ont contribué à faire de cette rencontre un événement.

Denis MIÉVILLE
Directeur de l'Institut de logique et
du Centre de Recherches Sémiologiques

PROPRIÉTÉS, MONDES POSSIBLES, OBJETS ET PROFILS

PROBLÈMES DE MÉRÉOLOGIE MODALE

Frédéric NEF

Dans la métaphysique analytique récente les propriétés apparaissent comme l'un des meilleurs candidats au titre d'éléments constitutifs de la réalité. S'il y a une structure ontologique à côté de la structure physique ou survenant sur cette structure physique, celle-ci comprend, à titre d'éléments atomiques, des propriétés. Si on poursuit l'analogie, le niveau subatomique est celui des propriétés individuelles, «l'alphabet de l'être» (D.C. Williams).

L'analyse de cette structure métaphysique a fait recours de plus en plus fréquemment, avec des résultats importants, aux formalismes méréologiques. Ceux-ci mettent en avant les relations de tout/partie (*parthood*) et de dépendance (Simons 1987, Casati et Varzi 1999).

Dans ce texte nous ne traiterons pas des rapports de dépendance, ni de fondation¹. Nous essaierons plutôt de répondre à la question suivante: Quels sont les problèmes posés par l'association d'une ontologie formelle des propriétés, et d'une méréologie modale? Nous discuterons quelques-uns des problèmes élémentaires posés par cette association. Nous admettrons éventuellement des propriétés individuelles. Dans la mesure où les propriétés seront considérées relativement à des mondes possibles, les problèmes posés relèveront d'une ontologie intensionnelle.

1 Dans la bibliographie, on se référera aux travaux de K. Mulligan et P. Simons. Travaux de logique 14, 2001.

1. De l'approche ensembliste à l'approche méréologique

Dans une sémantique extensionnelle on ne peut correctement donner l'intension d'un prédicat comme "unicorne" et le distinguer de "créature à un œil".

Dans la sémantique intensionnelle les propriétés sont conçues comme des fonctions des mondes possibles dans les ensembles d'individus:

une expression de type $\langle s \langle e, t \rangle \rangle$ réfère ainsi à une fonction des mondes possibles dans des ensembles d'individus. Maintenant, les ensembles d'individus servent à l'interprétation des prédicats et un prédicat réfère dans différents mondes à différents ensembles. Cette référence multiple d'un prédicat peut être vue comme une fonction des mondes possibles dans des ensembles d'individus et cette fonction peut être pensée comme l'intension d'un prédicat. Toute intension de ce type sera appelée une *propriété*. (Gamut 1991: 121)

La propriété "unicorne", avec comme domaine de départ le monde actuel, donne une extension vide. Dans tel monde possible comme domaine de départ, l'extension n'est pas vide.

On pourrait objecter que le narval est uncorne². Or, on le sait, il n'est pas uncorne; sa défense n'est nullement une corne, c'est sa canine gauche, hypertrophiée.

2 Dans *L'Unique Fondement d'une Démonstration Possible de l'Existence de Dieu* (1763), Kant dans un contexte différent compare aussi la licorne et le narval (je remercie P. Simons et A. de Libera de m'avoir rappelé ce passage). La différence entre ces deux animaux est «qu'on accorde l'existence au narval (licorne de mer) et non à la licorne de terre» (Kant 1980: 326). Kant remarque l'énoncé «un narval est une chose existante» est mal formé, il faut plutôt déclarer: «à un certain animal marin, qui existe, appartiennent les prédicats que je pense en les réunissant dans le concept de narval» (*op.cit.*). La raison de la différence est que «la représentation du narval est une idée tirée de l'expérience» (*ibid*). Pour «prouver le bien-fondé de la proposition que je prononce sur l'existence, je ne chercherai pas cette dernière dans le concept du sujet, car on y trouve que des prédicats de la possibilité, mais j'invoquerai l'origine de la connaissance que j'en ai» (*ibid*). Kant distingue donc la possibilité qui est la réunion des prédicats et l'existence qui est tirée de l'expérience. Le concept est la somme des prédicats, mais dans la représentation il y a plus, il y a «l'idée tirée de l'expérience». Ce passage est une critique de la doctrine wolffienne de l'existence: «*Hinc Existentiam definio per complementum*

Examinons ce contrefactuel:

Si les narvals étaient unicornes, leur canine gauche aurait une taille normale

Il pose le problème classique des relations entre propriétés.

Il est vrai en vertu du principe de modification minimale des mondes impliqués dans le contrefactuel: si le narval est (authentiquement) un corne, il ne peut plus être fallacieusement un corne. Se dresserait sur son front la canine hypertrophiée et la vraie corne – il serait bicorne. Si la première doit s'effacer au profit de la seconde, alors la canine, en vertu du principe cité, retrouve sa taille normale. Le point intéressant est le suivant. Le contrefactuel ne donne l'instruction que de modifier une propriété (et celles qui en dépendent ou qui en résultent, ce qui n'est pas le cas pour la canine). Il ne s'agit pas d'une relation d'implication entre propriétés. La relation en question est plus complexe.

On connaît donc les défauts d'une telle construction. Notamment l'incapacité à distinguer des propriétés qui donnent les mêmes classes d'individus (trilatéral et triangulaire), les difficultés liées à la détermination de propriétés intrinsèques, naturelles etc. Ces défauts qui proviennent tous de ce que la théorie n'individue pas assez finement les propriétés. Elle ne fait que décrire leur fonction de manière extérieure sans aller voir leur structure interne. Elles proviennent, donc, c'est l'hypothèse que l'on fait ici, du fait que cette théorie des propriétés se présente trompeusement comme une métaphysique: il ne s'agit que d'une sémantique, qui n'a pas rompu avec l'approche traditionnelle des propriétés. Approche en termes d'extension – la thèse suivant laquelle les propriétés sont des classes d'individus, même si cette extension n'est pas identique de monde en monde. Bref, trois faits seraient liés: le fait

possibilitatis: quam definitionem nominalem esse paret (§ 191 Log.) & ad recte philosophandum utilem ipso opere experiemur» (Wolff 1736: 143). L'existence du narval est donc plus qu'un «complément de la possibilité» de ses prédicats, c'est quelque chose que je connais par l'expérience et qui s'ajoute au concept. L'existence n'est pas tirée du sujet mais est prédiquée du sujet (c'est le sens de la paraphrase «un certain animal marin, qui existe etc.») Kant distingue donc possibilité logique (non contradiction des prédicats dans le sujet) et possibilité réelle (compatibilité avec les principes d'une expérience possible). Le narval est réellement possible et la licorne est seulement logiquement possible.

que la théorie soit *coarse grained*, la permanence d'un cadre ensembliste et l'absence d'une véritable analyse métaphysique de la structure interne des propriétés. La description des propriétés en termes de mondes possibles, sur une base ensembliste est inadéquate.

Une fois ce constat dressé et rappelé, il existe plusieurs possibilités. L'une d'entre elles est de réinterpréter la théorie des ensembles. Jubien (1989) fait ainsi reposer cette théorie sur la théorie des propriétés, en renversant la relation entre les deux théories, en considérant les ensembles comme des propriétés. Jubien insiste avec raison sur le fait que la métaphysique est première par rapport à la logique, à la sémantique et à la théorie des ensembles. Une telle métaphysique des propriétés entraînerait une nouvelle compréhension de la nature et de la fonction du cadre ensembliste. S'ensuivrait la définition d'objectifs précis en ce qui concerne l'individuation fine des propriétés, les conditions auxquelles on pourrait la décrire. Des théories aussi différentes que la logique intensionnelle de Bealer (1989), la méréologie modale ou les développements de la métaphysique des tropes (Bacon 1995) ainsi que dans une certaine mesure certains développements néomeinongiengs (Zalta 1985), avec la distinction entre propriétés encodées et exemplifiées, vont dans cette direction.

Il ne s'agira pas ici de comparer ces programmes de recherche. Plus modestement, nous traiterons des difficultés que la métaphysique des propriétés doit affronter, si elle souhaite à la fois conserver un composant modal et repenser le cadre ensembliste de départ. On discutera dans cette optique deux questions classiques. Celle de l'essentialisme méréologique – les propriétés peuvent-elles varier de monde à monde? – et celle liée à la fameuse loi de Port-Royal, qui stipule que l'extension et l'intension des propriétés sont en raison inverse. Précisons auparavant le cadre d'analyse, c'est-à-dire la théorie des propriétés en question, dans son aspect grossièrement méréologique.

2. Théorie(s) des propriétés

L'ontologie qui sera admise comme point de départ (ce sera le but d'un autre travail que d'en justifier la pertinence) admet comme concepts fondamentaux les propriétés, les objets, les états de choses et les mondes possibles (cf. Mulligan 2000, Nef 1998). On ne cherche pas ici à déterminer lequel de ces concepts est premier. Il apparaîtra clairement que nous avons des raisons de penser que c'est celui de propriété, mais nous ne discuterons pas ce point; nous nous contenterons de l'admettre. Le cadre qui sous-tend cette amorce de construction à plusieurs niveaux, multi-échelles, est clairement combinatoire. Les objets sont des combinaisons de propriétés; les états de choses sont des combinaisons d'objets et de propriétés (en considérant les relations comme un type de propriétés polyadiques); les mondes possibles sont des combinaisons d'états de choses.

On peut trouver cette construction pléthorique. Toutefois, notre but n'est pas de présenter une théorie minimale et élégante, mais de supposer une théorie assez générale. Ceci pour refléter les difficultés théoriques actuelles, concernant les propriétés, dans un contexte à la fois méréologique et modal.

P. Simons notait en 1987, que concernant ces questions, la situation était opposée à celle de la méréologie non modale: dans ce dernier cas une prolifération de systèmes et un espoir raisonnable de maîtriser les problèmes fondamentaux, de parvenir à une théorie équilibrée. Dans le second cas une abondance de problèmes non résolus et une relative rareté en développements logiques. La situation a évolué depuis cette date en ce qui concerne ce dernier point, mais non en ce qui concerne l'avant-dernier: les problèmes non résolus apparaissent toujours aussi nombreux.

On peut justifier une structure multi-échelle non réductive, en insistant sur le fait que le problème du ciment ontologique, si on autorise une image, se pose de manière différente pour chaque niveau de combinaison. La combinaison des propriétés en objets réclame que ces propriétés tiennent ensemble dans les objets et ne se dispersent pas aux quatre coins de l'espace logique. C'est

après tout un fait ontologique fondamental que les propriétés n'émigrent pas:

...il n'y a rien de plus simple et de plus clair dans toute la philosophie que ceci: qu'un accident n'émigre pas d'un sujet à un autre. (Manzoni, *Les Fiancés*)

La combinaison des objets en états de choses réclame que ces objets s'y rapportent les uns aux autres. La combinaison des états de choses, dans les mondes possibles, réclame l'usage d'une colle ontologique différente. Pour que les mondes aient quelque cohérence ils en réclament une. On admet, dans certains cas, que n'importe quelles propriétés peuvent être combinées, qu'un monde peut être composé de n'importe quels objets et propriétés. D'un certain point de vue c'est absolument correct, si l'on accepte la relativité de l'actualité. Si, par contre, nous sommes actualistes et non possibilistes, au sens où nous faisons varier les propriétés, les objets, relativement à une actualité privilégiée, celle de notre monde, l'*anything goes* se heurte à une intuition massive: les objets et les états de choses ne peuvent pas être recombinaison absolument n'importe comment, comme dans une nuit de la Walpurgis, un *tohu bohu*, un *stacato*.

La structure métaphysique du monde actuel suppose à chaque étape de sa construction un lien ontologique et le *anything goes* n'a pour fonction que de montrer la relativité de cette construction. Revenons au point que nous voulons faire remarquer. Rien *a priori* n'indique que le lien en question sera le même à tous les niveaux. La manière qu'ont les objets de se rapporter les uns aux autres dans un état de choses est différente de la manière que les propriétés ont de se tenir ensemble pour constituer un objet. Par exemple, on peut soutenir que dans le premier cas les objets dépendent des états de choses, tandis que dans le second cas les objets sont fondés sur les propriétés. Si c'est le cas, une construction stratifiée est préférable, car elle conserve quelque chose de nos intuitions métaphysiques premières. Cela dit, si on a des raisons de douter de la pertinence ou de l'importance de ces intuitions, ce raisonnement peut être radicalement contredit.

3. Les rapports tous/parties et les propriétés

La question des rapports tous/parties se pose à un triple point de vue dans une métaphysique des propriétés:

- (i) composition de propriétés par des opérations booléennes: conjonction, disjonction et négation;
- (ii) relation entre propriétés individuelles et propriétés générales;
- (iii) structure interne des propriétés en général.

I. Composition booléenne de propriétés

Il existe dans la méréologie modale des propriétés des structures de composition booléennes³ qui consistent à conjoindre, disjoindre ou nier des propriétés. Les disjonctions de propriété sont soit acceptées (Meixner), soit rejetées (Armstrong).

Si nous admettons des conjonctions de propriétés (comme “rouge et carré”), alors ces propriétés complexes (comme “rouge-et-carré” dont ce serait par défaut du lexique que n’y correspondrait pas un prédicat atomique, un mot comme “*rourré”), alors les conjoints de la propriété complexe (“rouge” et “carré”) seraient des parties authentiques de la propriété complexe. Les concepts de forme et de couleur sont distincts. Ontologiquement forme et couleur sont mutuellement dépendantes. Ne pourrait-il pas y avoir une propriété qui exprimerait la dépendance mutuelle entre la forme et la couleur?

Une question générale qui se pose est celle de la correspondance du langage booléen et du langage méréologique, correspondance qui s’exprimerait par ce type d’équivalence:

$$A \cap B = C \Leftrightarrow (A \wedge B) < C$$

où A, B et C sont des propriétés et < la relation “partie de”.

Soit l’exemple classique $\text{animal} \cap \text{rationnel} = \text{homme}$. Est-ce que la propriété complexe “animal et rationnel” est une partie de la propriété “homme”? Ce qui est certain c’est que

3 Cf. la parenté entre Ontologie et Méréologie lesniewskiennes et algèbre de Boole in Simons 1987: 104.

“homme”_{<ext} “animal” et “homme”_{<ext} “rationnel”. Cependant, pouvons-nous admettre qu’à des conjonctions de propriétés correspondent des propriétés complexes? S’il s’agit de propriétés générales, le problème est connu. Dans le platonisme c’est la difficulté de la participation à des universaux conjoints. S’il s’agit de propriétés individuelles, comment admettre des propriétés individuelles complexes? La sagesse de Socrate est le courage de Socrate peuvent-ils se combiner, se conjoindre en une vertu individuelle complexe?

II. Propriétés générales et propriétés individuelles

Cette dernière distinction est introduite pour être fidèle à une autre intuition métaphysique première, celle de la diversité des choses. Il n’existe pas deux choses strictement identiques. C’est une conséquence du principe d’indiscernabilité des identiques: si nous discernons deux choses qui ne sont pas identiques, alors elles diffèrent par au minimum une de leurs propriétés. La raison de cette diversité pourrait être cherchée dans l’instanciation elle-même des propriétés générales, ou universaux, dans la version “réalisme” (modéré ou pas) “des universaux”, mais cela conduit à des difficultés bien connues, dont on cherche à se sortir en postulant des entités problématiques, comme les essences individuelles non instanciées, les haecécités, etc. (cf. Rosenkrantz 1993). Il est beaucoup plus simple et, apparemment au moins, prometteur d’essayer une autre manière de rendre compte de cette intuition de la diversité, ou, ce qui revient au même, de tirer les conséquences radicales du principe d’indiscernabilité des identiques, c’est de considérer que cette diversité existe au niveau des éléments de construction de la réalité qui sont les premiers dans la construction, les propriétés⁴ et donc d’admettre des propriétés individuelles à la première étape.

Les propriétés générales sont-elles des tous dont les propriétés individuelles seraient des parties, les propriétés générales seraient-elles des ensembles de propriétés individuelles? Cette manière de voir les choses est répandue:

4 Nous nous inspirons ici de la conception itérative de la théorie des ensembles, cf. Boolos 1998.

other philosophers... have envisioned a reduction of properties to sets of tropes. (*Stanford Encyclopedia*, art. «Properties», 20)

Soit par exemple un univers réduit à deux objets x et y et à deux propriétés individuelles F et G , ayant une relation de ressemblance R . La propriété générale H est définie

$$H = F \cap G.$$

Est-ce que $F \wedge G < H$?

III. Structure interne des propriétés

Est-ce que cela a un sens de parler de parties de propriétés? Les objets ont des parties; les objets sont des ensembles de propriétés individuelles. Cela implique-t-il que les propriétés aient des parties? Les concepts ont des parties (les *notae* ou *Merkmale*). Mais peut-on généraliser cette caractéristique des concepts aux propriétés? Les concepts et les propriétés n'ont pas la même fonction, ni la même structure. Frege a montré qu'il ne faut pas confondre les concepts et les propriétés. Les nombres sont des concepts, non des propriétés. Les concepts contiennent seulement une représentation des propriétés essentielles et générales. Les concepts sont généraux et abstraits, les propriétés sont soit générales et abstraites, soit particulières et concrètes. La masse est un concept physique, mais la masse d'une particule est une propriété individuelle concrète. Les espaces conceptuels sont différents des espaces de propriétés. Un concept peut avoir une propriété comme partie, mais pas seulement.

Il existe un cas où dans l'histoire de la logique on a mis en lumière une loi (?) à cheval sur les concepts (ou idées) et les propriétés, c'est la loi de Port-Royal.

4. La loi de Port-Royal

La question des rapports tous/parties pour les propriétés est quelquefois mise en rapport avec la loi dite de Port-Royal⁵. Cette loi stipule ceci (dans la formulation de Meixner):

Pour toutes propriétés F et G, F est une partie intensionnelle de G ssi G est une partie extensionnelle de F.

L'exemple choisi est celui de la définition de "homme" comme "animal $\cap$ rationnel": animal est une partie intensionnelle de humain et humain est une partie extensionnelle de animal.

P. Simons exprime cette loi ainsi:

$$\text{Int}(F) < \text{int}(G) \supset \text{ext}(G) < \text{ext}(F)$$

On peut remarquer que pour Meixner il s'agit d'une équivalence: on a aussi $\text{ext}(G) < \text{ext}(F) \supset \text{int}(F) < \text{int}(G)$.

On peut émettre des doutes sur le fait d'affirmer que cette loi ou ce principe est une loi de la méréologie intensionnelle et même qu'elle est une loi tout court. Les propriétés dans cette loi sont définies de manière extensionnelle. L'intension d'une propriété n'est qu'une partie de concept dans cette loi. Par exemple dire que "rationnel" est une partie intensionnelle de "homme" ne veut rien dire d'autre que le concept "homme" contient "rationnel" à titre de *Merkmal*. Frege (1884, § 53), dans ce contexte, utilise la distinction propriétés/*Merkmale*:

Quand je parle de propriétés qui sont dites d'un concept, je n'entends évidemment pas les caractères (*Merkmale*) qui composent le concept. Ceux-ci sont des propriétés des choses qui tombent sous le concept, non du concept lui-même.

On peut même remarquer que ce n'est pas un hasard si à l'appui de cette soi-disant loi, c'est toujours le même genre d'exemple qui est donné, celui d'une définition réelle. Qu'en irait-il si nous prenions par exemple d'autres propriétés quelconques, qui possèdent une relation soit extensionnelle, soit intensionnelle

5 Cf. *inter alia* U. Meixner 112, P. Simons 1987: 169.

tout/partie? Pourrions-nous impliquer le correspondant intensionnel ou extensionnel suivant les cas? Il existe cependant une raison de penser que cette loi fait partie du composant intensionnel de la méréologie. Cette raison est la suivante: on ne peut interpréter “G est une partie extensionnelle de F” comme “tous ceux qui possèdent la propriété G possèdent la propriété F”, mais comme “tous les possibles qui possèdent une propriété G possèdent la propriété F”.

Ce point est discuté par U. Meixner (*op. cit.*) à propos du contre-exemple de Quine dirigé contre les propriétés, celui qui montre que “créature à un cœur” et “créature à deux reins” étant coextensifs ne sont pas distinguables. L’argument de Quine est le suivant: les propriétés sont, puisqu’on ne peut les identifier, des entités des ténèbres. Elles peuvent être certes définies extensionnellement, en termes de classes (si elles sont définies intensionnellement, l’argument s’arrête car on a montré qu’elles sont des créatures des ténèbres), mais dans le cas où pour deux propriétés distinctes conceptuellement F et G, leurs extensions sont identiques, on ne peut distinguer F et G, donc il faut éliminer les propriétés.

La réponse à cet argument est de poser que les extensions de “créature à un cœur” et “créature à deux reins” ne sont pas identiques dans tous les mondes possibles. Donc que les intensions de ces propriétés sont différentes. Tous les êtres possibles qui possèdent un cœur ne possèdent pas nécessairement deux reins et inversement.

Ce n’est pas comme le fait de posséder un pied droit si on possède un pied gauche: tous les êtres possibles qui possèdent (dans leur plan d’organisation) un pied droit possèdent un pied gauche. La nécessité dans ce cas est d’ordre spatial, il s’agit d’un essentialisme topologique⁶ la nécessité topologique en question diffère de celles touchant les relations de contact, de contiguïté ou d’intériorité (les relations nécessaires entre le *recto* et le *verso*). Dans le cas du pied gauche et du pied droit la nécessité signifie que les concepts spatiaux de gauche et de droit sont dépendants d’une relation de symétrie telle que le pied gauche est l’image en

6 Cf. Casati & Varzi: 155 ss.

miroir du pied droit. Soit I cette opération d'inversion autour d'un axe de symétrie. Soit x le pied droit et y le pied gauche. On a alors

$$I(x) = y \wedge I(y) = x.$$

D'autre part, la relation S , relation de symétrie, porte sur x et y : $S(x,y)$, telle que:

$$I(x) = y \wedge I(y) = x \Leftrightarrow S(x,y).$$

La nécessité en question est que si x est une partie de z , alors nécessairement y est une partie de z , x et y étant définis comme ci-dessus:

$$(\forall x)(\forall y)s)S(x,y) \leftrightarrow \Box \{ S(x,y) \supset [(x < z \supset y < z)] \}.$$

Cela dit, nous ne pouvons pas affirmer que dans tous les mondes l'une des deux propriétés est une partie intensionnelle de l'autre: la relation tout/partie peut varier de monde en monde. U. Meixner affirme cette relation tout/partie nécessaire (*op. cit.*: 251): la propriété F serait une partie de G ssi pour tout x et pour tout monde w si x appartient à G , il appartient à F dans w :

$$(\forall x)(\forall w) F < G \Leftrightarrow x \in_w G \supset x \in_w F$$

ce qui est équivalent à:

$$(\forall x) \Box (F < G \Leftrightarrow x \in G \supset x \in F)$$

ce qui nous reconduit à l'essentialisme méréologique, dont nous pouvons avoir des raisons de douter. Nous exposerons les raisons de soutenir l'essentialisme méréologique avant de le critiquer et nous examinerons alors les issues qui restent ouvertes à une théorie modale des relations tous/parties dans les propriétés.

5. Le principe d'essentialisme méréologique

Le principe de l'essentialisme méréologique (PEM) est:

que pour tout x quelconque, si x a y comme une partie, alors y est une partie de x dans tout monde où x existe (...) chaque tout a les parties qu'il a nécessairement (...) si y est une partie de x , alors la propriété d'avoir y comme l'une de ses parties est essentielle à x . (Chisholm 1989: 66)

Chisholm (1989: 68) oppose à l'essentialisme méréologique l'inessentialisme. Il consiste à affirmer qu'un tout quelconque pourrait être fait de n'importe quelles choses. Par exemple que cette table pourrait être faite du nombre 36 et de la propriété bleue. Ce qui est en cause c'est la rigidité de la constitution matérielle de cette table et en général des artefacts. Cette table aurait pu être constituée de savon de Marseille, mais une fois constituée de ce plateau et de ces pieds, elle est nécessairement constituée ainsi. Ce plateau et ces pieds sont nécessairement constitués de ce qui les constitue: s'ils étaient constitués d'une autre portion de matière ils seraient un autre plateau et d'autres pieds. L'identité de l'objet est la rigidité de sa constitution:

x constitue à t le même objet physique que y constitue à t' =_{def} il y a un z tel que x constitue z à t et y constitue z à t' . (*op. cit.*: 70)

Chisholm repousse l'intuition que nous avons que le PEM est faux; il voit un "conflit d'intuition" dans notre tendance à rejeter le PEM, alors que nous avons de bonnes raisons de le poser, conflit entre une intuition ontologique naïve et une intuition philosophique. Il cite un exemple de raisonnement courant qui semble effectivement mettre en échec ce principe, en montrer le caractère artificiel et la réponse qu'il donne à ce contre-exemple révèle une difficulté profonde de sa théorie, celle de donner des critères plausibles d'identité de l'objet à travers le temps. L'exemple est le suivant:

- (i) ma voiture avait des parties la semaine dernière qu'elle n'a plus et aura des parties la semaine prochaine qu'elle n'a pas encore;
- (ii) le PEM implique que si quelque chose est une partie de ma voiture, cette chose est une partie de ma voiture aussi longtemps que cette dernière existe;
- (iii) donc le PEM est faux.

(i) est une constatation du sens commun, irréfragable.

(ii) est justifié implicitement par une implication que l'on peut formaliser ainsi

$$(x < y \supset \Box x < y) \supset (\forall t) x <_t y$$

et qui repose sur une analogie entre le temps et la modalité. La nécessité implique l'omnitemporalité, l'opérateur temporel G , «il sera toujours le cas que» (ou: $\forall t, t > t_0$) est analogue à l'opérateur modal de nécessité $\Box$ «il est nécessaire que» (ou: $\forall w$).

Pour échapper à (iii) Chisholm est conduit à introduire une distinction problématique entre partie au sens technique de la méréologie et partie au sens courant. Pourquoi? Parce qu'il existerait un usage strict de "partie" dans la mineure, (ii), et un usage populaire de "partie" dans la majeure, (i). En effet dire qu'une partie de ma voiture constituait ma voiture hier (et ne la constitue plus aujourd'hui) veut dire que cette partie constituait ma voiture telle qu'elle était constituée hier (et non que cette partie qui constituait ma voiture constitue ma voiture d'aujourd'hui). Le problème est qu'il y a bien identité référentielle, fondée sur la mémoire, entre ma voiture telle qu'elle était constituée hier et telle qu'elle est constituée aujourd'hui.

Le PEM est certainement une thèse modale (sur la nécessité de l'identité entendue comme l'identité des parties à travers les mondes), mais c'est aussi une thèse temporelle, affirmant l'identité de l'objet à travers le temps. Ce principe revient à dire qu'à chaque instant un objet est constitué de ses parties et qu'il est nécessairement constitué de ses parties (en ce sens c'est un refus de l'objet vague ou incomplet). L'argument qui consiste à proposer comme alternative l'inessentialisme, défini comme le fait de conjindre des propriétés disparates est révélateur. Cet argument montre que

le PEM a également pour fonction de garantir la cohésion des parties à l'intérieur de l'objet, à partir du raisonnement suivant:

A = on ne peut mettre ensemble n'importe quelles propriétés
 $\neg A$ implique le faux
 A implique le PEM
 donc le PEM est vrai

Le PEM a donc trois fonctions distinctes qui consistent à affirmer trois choses distinctes, trois types de nécessité, modale, temporelle et ontologique:

(i) *La nécessité de la constitution matérielle*

Si un objet x est constitué d'une portion de matière y, alors il est nécessairement constitué (C) de cette portion de matière:

$$(\forall x)(\forall y) [C(x,y) \Leftrightarrow \Box C(x,y)]$$

(thèse de la rigidité de la constitution matérielle).

(ii) *La nécessité de la possession des parties à un certain instant et la permanence de cette possession*

$$(\forall x)(\forall y) [\{(x < y) \supset (\Box x < y)\} \supset (\forall t x <_t y)].$$

(iii) *De la cohésion ontologique des propriétés*

Cette cohésion est nommée "co-présence des propriétés" et P. Simons (1994) a donné des arguments pour la compléter par une relation de fondation. En effet, cette co-présence des propriétés dans l'objet est plus qu'une co-tenabilité des prédicats, qui serait réductible en dernière analyse à une non contradiction.

6. Essentialisme méréologique et identité des objets

Hume (1995: 345) déclare à propos de l'identité des objets:

Nous avons une idée distincte d'un objet qui demeure invariable et ininterrompu à travers un changement supposé de temps; cette idée nous l'appelons identité ou mêmeté (*sameness*). Nous avons aussi une

idée distincte de plusieurs objets différents existant successivement et liés les uns aux autres par une relation étroite. (...) Mais bien que ces deux idées, l'identité et la succession d'objets reliés, soient en elles-mêmes parfaitement distinctes et même contraires, il est pourtant certain qu'elles sont généralement confondues dans notre manière ordinaire de penser. (...) Cette ressemblance est la cause de la confusion et de l'erreur et elle fait que nous substituons la notion d'identité à celle d'objets reliés. Même si à un instant donné, nous pouvons considérer la succession corrélatrice comme changeante et discontinue, nous ne manquerons pas de lui attribuer, l'instant après, une identité parfaite et de la regarder comme invariable et ininterrompue.

Hume analyse le mécanisme de constitution du sens populaire de l'identité, qui peut donner lieu aux illusions des théories de la substance. À chaque instant un objet est constitué d'une somme de parties et des changements quelquefois imperceptibles viennent affecter ces parties. En toute rigueur on ne peut parler que d'une succession d'objets reliés. Cependant cette succession est continue et ininterrompue et c'est ce qui conduit à postuler une "identité parfaite" là où il n'y a qu'une identité relative.

La ressemblance mentionnée par Hume peut être réinterprétée dans la métaphysique des propriétés individuelles ou tropes.

Soit un objet o constitué d'un ensemble A de tropes à t :

$$A = \{x_1, \dots, x_n\} \text{ à un instant } t,$$

soit cet objet constitué d'un ensemble B de tropes à t' :

$$B = \{x_1, \dots, x_n\} \text{ à } t'.$$

Par commodité on appellera une occurrence d'objet à un instant t_x un profil temporel ou en bref un profil. Un profil p est donc un ensemble de tropes à un instant t_x :

$$p = (X, t), \text{ où } X \text{ est un ensemble de tropes.}$$

Un objet est une série de profils: $O = (p_1 \dots p_n)$

Trois possibilités se présentent:

$$A \wedge B = \emptyset$$

$$A \wedge B = C$$

$$A = B$$

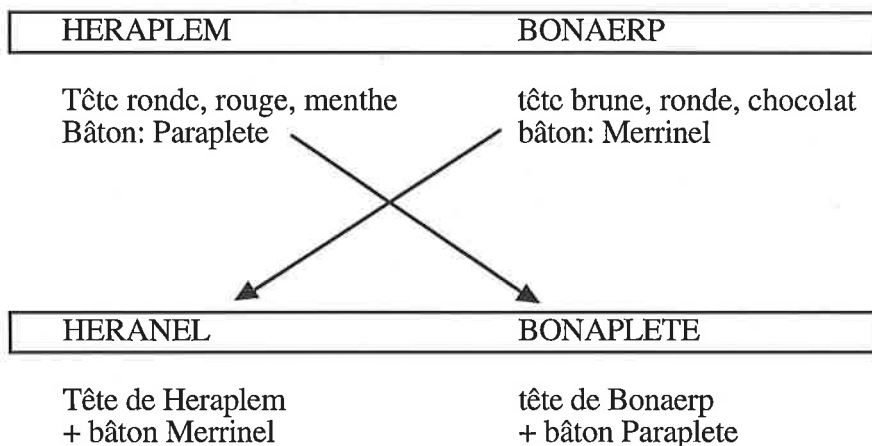
Dans le premier cas il n'y a pas de trope commun, les deux profils correspondants sont déconnectés. Ceci se produit par exemple au terme d'une série de modifications dans les paradoxes de l'identité matérielle, du type du Bateau de Thésée (cf. Engel & Nef 1988), par exemple quand après une série d'étapes marquées chacune par une modification, l'objet ne contient plus rien des parties (ou des propriétés) qu'il avait au départ. On peut penser ici au fameux exemple de la table de Chisholm dont on remplace chaque jour une pièce (un pied ou le tableau): au bout de cinq jours, elle n'a plus rien en commun avec la table de départ.

Dans le dernier cas les profils sont identiques. C'est un problème de savoir si cela peut se produire pour des profils d'objets matériels, car ces derniers sont soumis à une continuelle altération.

Dans le deuxième cas, le plus intéressant, il y a un ensemble de tropes communs aux deux profils, un recouvrement de certains tropes. C'est le cas où l'identité est problématique (dans le premier et le dernier cas, ou pour des raisons négatives ou pour des raisons positives, l'identité ou la non identité ne pose pas de problème). Dans ce cas, il faut distinguer les tropes qui sont communs et les tropes qui se ressemblent. Dans le cas d'une déformation continue d'un trope spatial (par dilatation par exemple), il y a ressemblance du trope de départ (avoir *tel* volume) et du trope d'arrivée (avoir *tel* volume, plus grand que le premier). S'il y a ressemblance des deux tropes, il y a possibilité *a posteriori* de les subsumer sous un universel (ici: "avoir *un* volume, en général"). Dans le cas d'échange *partes extra partem* la situation est différente et les critères sont plus difficiles à établir.

Reprenons l'exemple de D.C. Williams (1953), en le modifiant un peu. Soit deux sucettes HERAPLEM, BONAERP. Elles sont constituées d'une tête et d'un bâton. Heraplem a une tête ronde, rouge et de menthe; Bonaerp a une tête brune, ronde et de chocolat. Le bâton de Heraplem se nomme PARAPLETE et le bâton de Bonaerp MERRINEL; ces bâtons sont similaires (en forme, taille, couleur). Heraplem et Bonaerp sont partiellement

identiques, car leurs bâtons sont identiques et qu'elles possèdent le trope de rondeur en commun. Imaginons maintenant que nous échangeons les bâtons de ces sucettes. Appelons HERANEL et BONAPLETE les produits des substitutions (Heranel = Heraplem avec le bâton Merrinel, antérieurement à Bonaerp, Bonaplete = Bonaerp avec le bâton Paraplete antérieurement à Heraplem). Soit le schéma suivant:



HERAPLEM + MERRINEL = HERANEL
BONAERP + PARAPLETE = BONAPLETE

Est-ce que Heranel et Bonaplete sont partiellement identiques? Oui, parce qu'ils possèdent toujours un trope de couleur en commun. Est-ce que Heranel et Heraplem d'une part et Bonaerp et Bonaplete d'autre part sont partiellement identiques? Oui pour la même raison et parce que de plus ils possèdent des bâtons similaires. Cependant notre intuition est que Heraplem et Bonaerp possèdent une plus grande similarité que Heraplem et Herael (ou Bonaerp et Bonaplete). En effet, même si les bâtons, Paraplete et Merrinel, sont similaires, Heranel a hérité de Merrinel qui était le bâton de Bonaerp. Dans l'histoire de Heranel il y a eu détachement d'une partie, même si c'est pour l'adjonction d'une partie similaire. Bref, est-ce que l'essence individuelle de Heraplem et celle de Heranel sont identiques?

Qu'est-ce que cela montre? Premièrement, que le concept populaire d'identité est effectivement vague ou plus exactement que le schème conceptuel de l'identité à notre disposition dans le raisonnement courant manque de précision.

Deuxièmement, qu'il faut distinguer les conditions externes d'identité, qui peuvent s'accommoder d'un type-identité (les bâtons sont du même type, par exemple, s'il s'agit d'un artefact industriel, de la même marque), tandis que les conditions internes, qui regardent le concept individuel complet d'un objet peuvent sans doute intégrer l'histoire des substitutions de parties, alors même que celles-ci n'ont pas menacé l'intégrité fonctionnelle de l'objet. On retrouve cette distinction, là où elle a une signification existentielle ou vitale, dans le cas où il s'agit de déterminer un dommage suite à un accident. Si dans ce cas une greffe parfaitement réussie permet de retrouver toutes les fonctions de l'organe endommagé ou disparu, on argue quelquefois que la greffe elle-même a représenté un dommage à soi seule. C'est un cas où l'éthique appliquée met en relief une distinction métaphysique entre définition fonctionnelle et définition historique.

7. Conclusion

On peut être perplexe devant les considérations qui précèdent. Quels résultats contiennent-elles? Sont-elles seulement une suite de remarques aporétiques sur les difficultés d'appliquer la relation tout/partie aux propriétés, dans un cadre nominaliste et particulariste modéré? Pourquoi ne pas avoir effectivement renoncé à ces propriétés qui introduisent tant de difficultés?

L'ontologie formelle des propriétés individuelles peut cependant se prévaloir d'un certain nombre de résultats. La relation d'exemplification apparaît moins mystérieuse. La relation des propriétés générales aux propriétés individuelles ouvre la voie d'une analyse de la structure des propriétés. La méréologie modale apparaît vraiment être l'outil adéquat pour une étude de la continuité spatio-temporelle. La nature des paradoxes sur la constitution matérielle des objets est reconnue avec clarté. La dualité

des structures booléennes et des structures méréologiques est démontrée.

Les problèmes métaphysiques qui restent à résoudre font peut-être partie de ce qu'il faut au moins pour l'instant nommer des mystères, en nombre limité: la co-présence des propriétés dans les objets, la relation exacte entre les propriétés particulières et les propriétés générales (si l'on pense que la ressemblance et un mécanisme de constitution comme l'abstraction ne livrent pas le dernier mot sur l'explication d'un phénomène).

Institut de philosophie
Université de Rennes I
 F 35042 RENNES
 e-mail: frederic.nef@univ-rennes1.fr

Bibliographie

- ARMSTRONG D.M. (1978). *A Theory of Universals*. Vol. II: *Universals and Scientific Realism*. Cambridge: Cambridge University Press.
- BACON J. (1995). *Universals and Property Instances: the Alphabet of Being*. Oxford: Blackwell.
- BEALER G. (1989). Fine-Grained Type-Free Intensionality. In: G. Chierchia *et al.*, vol. 1, 177-230.
- BOOLOS G. (1998). *Logic, Logic and Logic*. Cambridge: Harvard University Press.
- CASATI R. & VARZI K. (1999). *Parts and Places: the Structures and Spatial Representation*. Cambridge: MIT Press.
- CHIERCHIA G., PARTEE B.H. & TURNER R. (1989). *Properties, Types and Meaning*. Vol. 1: *Foundational Issues*. Dordrecht: Kluwer.
- CHISHOLM R.M. (1989). *On Metaphysics*. Minneapolis: University of Minnesota Press.
- ENGEL P. & NEF F. (1988). Identité, vague et essences. *Les Études Philosophiques* 4, 475-494.

- FREGE G. (1884). *Les fondements de l'arithmétique: recherche logico-mathématique sur le concept du nombre*. Paris: Seuil. Trad. C. Imbert, 1990.
- GAMUT L.T.F. (1991). *Logic, Language and Meaning*. Vol. 2: *Intensional Logic and Logical Grammar*. Chicago: University of Chicago Press.
- HUME D. (1995). *Traité de la nature humaine. L'entendement*. Paris: GF-Flammarion. Trad. B. Carranger et P. Saltel.
- JUBIEN M. (1989). On properties and property theory. In: Chierchia et al. (eds), 159-177.
- KANT I. (1980). *Œuvres philosophiques*. T. 1: *Des premiers écrits à la critique de la raison pure*. (Sous la dir. de F. Alquié). Paris: Gallimard, Bibliothèque de la Pléiade, n° 286.
- LEWIS D. (1986). *On the Plurality of Worlds*. Oxford: Blackwell.
- LIBERA A. de (1999). *L'art des généralités. Théories de l'abstraction*. Paris: Aubier.
- MEIXNER U. (1997). *Axiomatic Formal Ontology*. Dordrecht: Kluwer.
- MULLIGAN K. (2000). Métaphysique et ontologie. In: P. Engel (éd.), *Précis de philosophie analytique*. Paris: PUF, 5-34.
- NEF F. (1998). *L'objet quelconque*. Paris: Vrin.
- ROSENKRANTZ G. (1993). *Haecceities*. Dordrecht: Kluwer.
- SIMONS P. (1987). *Parts*. Oxford: Oxford University Press.
- SIMONS P. (1994). Particulars in Particular Clothings: Three Trope Theories of Substance. *Phil. and Phenom. Research* 54-3, 553-575.
- SWOYER (1996). Theories of Properties: From Plenitude to Paucity. *Philosophical Perspectives* 10, 243-264.
- WILLIAMS (1953). The Elements of Being. *Review of Metaphysics* vol. 7.
- WOLFF Ch. (1736). *Philosophia Prima sive Ontologia Methodo Scientifica pertractata (...)* Francfort & Leipzig. Rééd. In: éd. J. Ecole, *Gesammelte Werke*, vol. 9, Hildesheim: Olms.
- ZALTA E. (1985). *Abstract Objects*. Dordrecht: Reidel.

IDENTITY AND COUNTERPARTHOOD

Christopher HUGHES

I.

I have never crossed the Himalyas, though I might have done. So there is a non-actual (or, if you prefer, a non-actualized) possible world (or possible state of the world) whose existence is a sufficient condition for the truth of "I might have crossed the Himalyas". At such a possible world, someone crosses some mountains. Is that person me, and are those mountains the Himalyas? Or are they (non-actual) individuals different from me and from the Himalyas?

According to David Lewis (1968), the latter suggestion is the better one:

Where some would say that you are in several worlds, in which you have somewhat different properties and somewhat different things happen to you, I prefer to say that you are in the actual world and no other, but you have counterparts in several other worlds. Your counterparts resemble you closely in content and context in important respects. They resemble you more closely than do the other things in their worlds. But they are not really you. Indeed we might say, speaking casually, that your counterparts are you in other worlds, that they and you are the same; but this sameness is no more a literal identity than the sameness between you today and you tomorrow. It would be better to say that your counterparts are *the men you would have been*, had the world been otherwise.

For Lewis (1973), what an individual might have done (or might have been) is what its counterparts do (or are). (When the

possibility is unrealized, the counterparts will all be in other worlds).

Saul Kripke emphatically rejects counterpart theory:

[For Lewis] if we say, "Humphrey might have won the election" (if only he had done such-and-such), we are not talking about something that might have happened to *Humphrey* but to someone else, a 'counterpart'. Probably, however, Humphrey could not care less whether someone *else*, no matter how much resembling him, would have been victorious in another world. Thus Lewis' view seems to me even more bizarre than the usual notions of transworld identification that it replaces. (Kripke 1972: n. 13)

Even more objectionable is the view of David Lewis. According to Lewis, when we say "Under certain circumstances Nixon would have gotten Carswell through", we really mean "Some man, other than Nixon but closely resembling him, would have gotten some judge, other than Carswell but closely resembling him, through." Maybe that is so, that some man closely resembling Nixon could have gotten some man closely resembling Carswell through. But *that* would not comfort either Nixon or Carswell, nor would it make Nixon kick himself and say, "I should have done such and such to get Carswell through." The question is whether under certain circumstances Nixon *himself* could have gotten *Carswell* through. (Kripke 1971: 148)

Kripke appears to be saying here that, on Lewis' view, if we say "Humphrey might have won the election", we really mean "Someone different from Humphrey, but resembling Humphrey in certain ways, might have won the election": winning the election isn't (on Lewis' view) something that might have happened to Humphrey; it's something that might have happened one of his counterparts. Similarly, if we say "Under certain circumstances, Nixon would have gotten Carswell through", what we really mean is "Under certain circumstances, someone different from Nixon, but resembling him in certain ways, would have gotten someone different from Carswell, but resembling him in certain ways, through." But, Kripke thinks, it is obvious that if we say "Humphrey might have won the election", or "Nixon would have gotten Carswell through", we mean that *Humphrey*

might have won the election, or that *Nixon* would have gotten *Carswell* through.

It is obvious; but why exactly is this a problem for Lewis? Lewis holds that it is Humphrey's counterpart, rather than Humphrey himself, that wins the election (in another world); and it is Nixon's counterpart, rather than Nixon himself, who gets (a counterpart of) Carswell through (in another world). But he does not hold that it is Humphrey's counterpart, and not Humphrey himself, that *might have* won the election; and he does not hold that it is a counterpart of Nixon, and not Nixon, that *might have* gotten Carswell through. As Lewis sees it, it is true that Humphrey himself would have won, under different circumstances, precisely inasmuch as it is true that someone different from Humphrey (his counterpart) wins in an alternative possible world (Lewis 1983: 42). Still, it might be said, in the passages cited, Kripke has gotten hold of a difficulty for Lewis-style counterpart theory. The difficulty is that, intuitively, there is a gap between ordinary modal claims and the counterpart-theoretic claims that Lewis takes those ordinary claims to have the same truth-conditions as. For Lewis, a statement such as

- (a) Nixon might have *Fed*.

is true if and only if someone existing in this world or in another possible world who is like Nixon in such-and-such and ways (someone who resembles Nixon in context and context, and resembles Nixon at least as much as any of his worldmates) *Fs*. Thus for Lewis, (a) is true if and only if

- (b) In some possible world someone like Nixon in such-and-such ways *Fs*.

Now one would have thought that (b) was equivalent to

- (c) It might have been that: someone like Nixon in such-and-such ways *Fs*.

The trouble is that there is no guarantee that (a) and (c) will have the same truth-value.

It depends on what the predicate *Fs* is. Suppose for example, *Fs* = *comes from gametes different from these very gametes* [the

ones Nixon actually came from]. Then, it seems, (a) will be false – since Nixon couldn't have come from gametes different from the ones he actually came from – but (c) will be true – since someone like Nixon in content and context (and no less like Nixon than any of his worldmates) might have come from gametes different from the gametes Nixon actually came from. (For those who doubt the essentiality of origin, similar examples could be provided turning on the essentiality of constitution, or for that matter the necessity of identity)¹.

Here is another way to put essentially the same point. Lewis gives truth conditions for *de re* modal predication in terms of counterparts. As he sees it, the counterpart relation has two features: that $x = y$ is not a necessary condition of y 's being a counterpart of x ; and that x 's having a counterpart that F s is a sufficient condition for the truth of x *might have Fed*. The difficulty is that unless identity is a necessary condition of counterparthood, it is hard to see why x *has a counterpart that F*s should be thought sufficient for x *might have Fed*. Some otherworldly individual y , resembling x in content and context (and resembling x at least as much as its worldmates) F s. Why should that be thought to guarantee that x could have F ed?

Here a counterpart theorist might protest:

Suppose that I have a counterpart that is F . Then there is a way things might have been such that, had things been that way, someone would have been F , and I would have been that someone. (As Lewis puts it in the passage cited earlier, your counterparts are the men you would have been, had things been otherwise). Now if there is a way things might have been such that, had things been that way, someone would have been F , and that someone would have been me, then it follows, as the night does the day, that I might have been F . So there is no mystery about why my having a counterpart who is F entails that I might have been F ; by the same reasoning, there is no mystery about why the converse entailment holds.

1 For an objection of this type, see Salmon (1980: 235, n. 1).

The counterpart theorist may stipulate that “counterpart” is to be used in such a way that “My counterpart is *F*” guarantees the truth of “There is a way things might have been such that, had they been that way, someone would have been *F*, and I would have been that someone.” Given that stipulation, my having an *F*-ish counterpart will guarantee that I might have been *F*. But then it is less than obvious that an individual’s resembling me (sufficiently) in content and context (at least as much as any of its worldmates) guarantees that that individual is my counterpart. Alternatively, the counterpart theorist could stipulate that “counterpart” is to be used in such a way that anything resembling me sufficiently in content and context (at least as much as any of its worldmates) is *eo ipso* my counterpart—whether or not it is numerically identical to me. But then it seems very doubtful that my having an *F*-ish counterpart is necessary and sufficient for the truth of “I might have been *F*”.

Also, inasmuch as Lewis insists both that (a) your (otherworldly) counterparts aren’t you; and (b) your (otherworldly) counterparts are who you would have been, had things been otherwise, this at least suggests that if things had been otherwise, you would have been somebody else. But, as Kripke has emphasized, there is a very strong feeling that under no (possible) circumstances would you have been anybody else.

A last worry about counterpart theory, suggested by Lewis’s remark that you and your (otherworldly) counterparts are no more identical than you today and you tomorrow: suppose it is true that this apple is unripe, but it will ripen. Consider two accounts of what makes it true that this apple is unripe (now), but will ripen (in the future). On the first account, *this apple is unripe, but it will ripen* is true at time t_{now} because this apple exists at t_{now} and is unripe at t_{now} , and at some time t later than t_{now} , this apple exists at t and is ripe at t . On this account, the apple is a “trans-time” individual, existing at both the present time and future times. On the second account, the currently existing (unripe) apple exists only at the present time, though it has what we might call “continuers” at other (future) times. And *this apple is unripe, but it will ripen* is true at time t_{now} because this apple is unripe now, and for some time t later than t_{now} and some continuer c of this

apple, it is true at t that c is ripe. This apple itself is not ripe at any time (it is unripe now, and non-existent elsewhere)--but it is "vicariously ripe" through its ripe continuers².

It seems (to me, at any rate) that there is a presumption in favor of the first account.

I suppose this has something to do with the relative simplicity of the two theories. The first account makes reference to this apple, and to present and future times; the second account makes reference to two (momentary) apples, present and future times, and an additional relation (the continuer-relation holding between (some) pairs of (momentary) apples). In the absence of reasons to think the more complicated account is preferable, I am inclined to think we should prefer the simpler one.

Similarly, it seems to me that there is a presumption in favor of a simpler and more straightforward account of what would make it true that this apple is actually unripe, but it might have been ripe -- an account that involves this apple and a plurality of worlds, rather than a plurality of (worldbound) apples, a plurality of worlds, and an additional relation (the counterpart relation).

Moreover, the second account differs from the first account not just in that it has more working parts, but also in that it provides truth-conditions that have to be defended against "open question" arguments. Suppose that at some future time t , some continuer of this apple is ripe. Given that the continuer relation is non-reflexive, there is a question about why this apple's having a ripe continuer should be considered both necessary and sufficient for the truth of *this apple will be ripe*. By contrast, there is no question about why this very apple's being ripe at a time t later than t_{now} should be deemed necessary and sufficient for the truth of *this apple will be ripe*. And, all things being equal, (putative) truth-conditions that are invulnerable to open question arguments

2 Compare Lewis (1973: 39): «What our Ripov cannot do in person at other worlds, not being present there to do it, he may do vicariously through his counterparts. He himself is not an honest man at any world--he is dishonest here, and nonexistent elsewhere--but he is vicariously honest through his honest counterparts».

are preferable to (putative) truth-conditions that are vulnerable to them³.

II.

Naturally, this is consistent with there being other considerations which overcome the (defeasible) presumption favoring trans-world individuals over counterparts. What might such considerations be?

Suppose that, like David Lewis, you accepted something very like the "foreign country" or "faraway planet" picture of possible worlds that Kripke dismisses. That is, suppose you thought that the actual world was the universe we inhabit, and other possible worlds were other universes, spatio-temporally unconnected to our own. Then, it might seem, you would embrace counterparts, and reject trans-world individuals, on the grounds that an individual can inhabit at most one universe. Kripke himself says (1971: 147) that counterpart theory is the most reasonable view to hold, for one who takes the foreign country view of worlds: "No one far away on another planet can be strictly identical with someone here". And various other philosophers have concurred with Kripke⁴.

-
- 3 In this context, it is interesting to note that when Lewis provides satisfaction conditions for counterfactual conditionals (Lewis 1973: 38), the first set of conditions he considers presuppose that individuals exist in more than one world. He defends his rejection of those conditions by providing a range of arguments purporting to show the preferability of a counterpart-theoretic account of *de re* modal predication to an account involving trans-world individuals.
 - 4 See for example M. Loux (1979: 63-64), and W. Lycan, (1990/1: 221). (Interestingly, although Lewis offers a number of considerations in favor of counterpart theory in *Counterfactuals* (39-41), none involve his conception of worlds as universes. In "Individuation by Stipulation and Acquaintance" (1983: 20-24) he does suggest that someone who thinks of possible worlds as universes has reason to believe in counterparts, lest properties that are not intuitively relations to worlds (or world-time pairs) turn out to be such relations. As Bigelow and Pargetter argue (1990: 168-69), this line of reasoning seems less than conclusive: someone who thinks possible worlds are universes might hold that it is at least as counterintuitive to deny that one and the same thing can exemplify roundness at different times and in different worlds, as it is to suppose that *being round* is a relation between individuals and world-time pairs. Thus Bigelow and Pargetter consider

But why exactly couldn't one thing be in two universes? After all, if (as Lewis supposes) the two universes don't share any times or places, this wouldn't be a case of one and the same thing being in two different places at the same time; it would be a case of one and the same thing being in two different places at two (temporally unrelated, and *a fortiori*) different times⁵.

Sydney Shoemaker (1979) holds that an individual persists through time only if there are causal connections between the earlier stages of that individual's life (or, in the case of an inanimate object, existence) and the later stages; qualitative similarity plus spatio-temporal continuity, in the absence of causal connections, is insufficient for identity through time. In support of this contention, he describes a case in which an absent-minded deity first decrees that a stone tablet will disappear into thin air at just before noon, and then – having forgotten all about his disappearance decree – also decrees that a (qualitatively indiscernible) stone tablet will appear at exactly noon. In this case, Shoemaker says, there are two different stone tablets (the A.M. tablet and the P.M. tablet), even though, to anyone not apprised of the deity's activity, it would look as though there were just one. The reason the A.M. tablet and the P.M. tablet cannot be identified is the absence of any (and, *a fortiori*, of the right sort) of causal relations between the A.M. stages of the existence of the table there before noon, and the P.M. stages of the existence of the table there from noon on.

If the identity of this individual at this time with that individual at that (distinct) time requires that there be causal connections between a this-timely stage of this individual's existence and a that-timely stage of that individual's existence, then individuals existing at different times in different (spatio-temporally disjoint) universes are identical only if there are causal connections between events in spatio-temporally disjoint universes.

it an open question whether a universalist should believe in counterparts or trans-world individuals.

- 5 Loux writes: «It is easy to see why Lewis anxious is so anxious to deny that a single individual can exist in more than one world... If all possible worlds are equally real... then you and I can exist in but a single world; we cannot, so to speak, inhabit two worlds at once». (*The Possible and the Actual*, 47). But why does it follow from the fact that we cannot inhabit two worlds *at once*, that we cannot inhabit two worlds?

But it is very doubtful that someone who thinks that alternative possible worlds are universes spatio-temporally disjoint from the one we inhabit can countenance the idea that there are causal connections between events in disjoint universes. For if there are causal links between, say, what is happening now in this universe, and what is happening in some disjoint universe, then surely the disjoint universe is not an alternative possible world (a *merely* possible world); it is a part of the actual world (spatio-temporally unconnected to, though causally connected to, the part we inhabit)⁶.

So there is a line of thought suggesting that if possible worlds are universes, the individuals inhabiting those worlds are world-bound⁷. But I am uncertain how cogent it is. Shoemaker's thought-experiment seems to show that even if table T is spatio-temporally continuous and contiguous with table T' , it does not follow that $T = T'$. But one can accept this claim, without conceding that the cross-time identity of x at t with y at t' requires causal relations between the t -stage of x (or, for the three dimensionalist, the t -stage of x 's career) and the t' -stage of y (or the t' -stage of y 's career). After all, some philosophers would maintain that it is logically possible for a thing and all its parts to go out of existence at an earlier time, and then, after an interval of non-existence, come back into existence, in spite of the absence of any causal relations between the phase of the thing's career just before its existence, and the phase of the thing's career after it came back into existence. Their intuition is that just as a thing could (uncausedly) pop into existence, or (uncausedly) pop out of existence, a thing that had (uncausedly, say) ceased to exist could (again, uncausedly) pop back into existence after an interval of non-existence. Why not? Perhaps the worry is that that nothing could ground the fact that we have the same entity before and after

6 The idea that there could be causal connections between something happening now in this universe, and something happening in a universe spatio-temporally disjoint from ours is in any case a problematic one, independently of whether or not disjoint universes are alternative possible worlds.

7 Most mediaeval philosophers held that God was the aspatial and atemporal first cause of the universe. Surely no one could (coherently) maintain both that the existence of the universe is the effect of a divine choice, and that God and the relevant choice-event were merely possible individuals and events, found only in alternative possible worlds.

the interval of non-existence. But why couldn't the identity of the pre-demise and post-demise thing be a primitive fact, not grounded in anything else?

In sum: it is not clear that the identity of x existing at t with y existing at t' presupposes the existence of causal relations between the t -stage of x 's existence, and the t' -stage of y 's existence. And if we can have non-causally-mediated cross-time identity between two things in the same universe, it is unclear that we couldn't have it between two things in different universes.

Suppose, though, that we can get from (what I shall call) a "universalist" conception of possible worlds to counterpart theory. Unless (we have reason to think) that possible worlds are universes, that doesn't give us a reason to accept counterpart theory. For unoriginal reasons, I share the (widely held) view that possible worlds are abstracta, and not universes. So I doubt there is much mileage in trying to support counterpart theory by appeal to a universalist conception of possible worlds.

III.

Considerations of a very different sort--involving identity and extrinsicality--might be thought to support counterpart theory. Suppose that in the actual world **G** a certain ship S is built from a set of planks, sinks on its maiden voyage, and gradually falls to bits on the sea floor. Suppose that in an alternative possible world **H**, the same set of planks are put together in the same way at the same time and place by the same shipwright. The resulting ship does not sink on its maiden voyage, but all its original planks come to be dispersed. They are subsequently reassembled by the same shipwright who originally assembled them, in just the way they were originally assembled. Is the ship resulting from the reassembly S ? Well, it depends. Suppose that in **H**, the dispersal of the original planks went hand in hand with the replacement of original planks by new planks (in a structure-preserving way), in the course of ordinary maintenance and repairs. Then, it appears, the ship whose planks were gradually replaced is ship S , and the ship resulting from the reassembly of the original planks is a

different ship (a new ship made from old planks). Suppose on the other hand that in **H**, the plank-dispersal was not accompanied by plank-replacement; the sailors took the ship to bits in order to more easily transport it overland, intending all the while that the shipwright would put the ship back together when they got to the next sea. Then, it seems, the ship resulting from the reassembly of the original planks is the ship *S*. (This assumes that the replacing planks have a dominant claim to constitute the original ship, and the reassembled planks a recessive one. The reader who thinks that the reassembled planks have a dominant claim may modify the story accordingly).

A counterpart theorist might say: whether or not the ship in **H** resulting from the reassembly of the old planks is the ship *S* would have been, had **H** been actual, depends on whether there is another ship in **H** that has a stronger claim to be the ship *S* would have been under those circumstances. This is perfectly compatible with counterpart theory, according to which whether an individual *i* in an alternative world is a counterpart of an actually existing individual *i'* depends, not just on the intrinsic properties of *i* and *i'*, but also on the relational or extrinsic properties of *i*--in Lewis' words, on *i*'s context as well as its content. Suppose on the other hand that individuals are trans-world, so that the ship *S* would have been in such-and-such circumstances is literally identical to *S*. Then it cannot be that whether a certain ship in another world is the ship *S* would have been, depends on whether in that world there is something else that has a stronger claim to be the ship *S* would have been. So--the counterpart theorist might conclude--inasmuch as the would-have-been relation has extrinsic determinants, it is not identity: the ship *S* would have been, had **H** been actual, is not *S*, but a counterpart of it.

The trouble with this argument--as Salmon has pointed out (1982: Appendix I)--is that the champion of trans-world identity (henceforth, "the (trans-world) identitarian") will deny that there is a particular ship in **H** which is the ship *S* would have been if but only if nothing else in **H** has a better claim to be the ship *S* would have been. Instead, the identitarian will say, there is a set of planks in **H** which constitutes the ship *S* would have been (that is, constitutes *S*) at a given time if but only if no other set of

planks in **H** has a better claim to constitute *S* at that time. This has the (at least initially surprising) result that which ship some planks constitute has extrinsic determinants, but the identitarian may find that acceptable (and independently plausible).

Suppose the champion of counterpart theory could tell a story like the story of *S*, in which (i) we are inclined to say that whether or not an otherworldly individual *y* is the individual an actual individual *x* would have been has extrinsic determinants; and (ii) this inclination cannot be explained away by appeal to the idea that constitution has extrinsic determinants. That would be a weighty consideration in favor of counterpart theory. But I do not see how such a story would go. If we could tell a story like the story of *S* that involved simples-things constituted only by themselves--that would do the trick. But I don't know how to do that.

IV.

A different line of reasoning purporting to favor counterparts over trans-world individuals takes as its starting point the idea that two different things cannot have all the same parts at all the same times. Why should this principle favor counterparts over trans-world individuals? Suppose that individuals are trans-world, so that an individual might have been *F* if and only if that individual is *F* in some world. Then *i* and *i'* must be distinct, as long as there is some property *F* such that *i*, but not *i'*, might have been *F*. (If there is a world *w* and a property *F* such that *i* is *F* in *w*, but *i'* is not *F* in *w*, then *i* and *i'* are discernible, and thus distinct). But there seem to be cases in which it is true both that this thing and that thing have all the same parts at all the same times, and that this thing, but not that thing, might have been *F*. An example discussed by Kripke involves a plant that never "grows past" its stem: the plant and the stem have all the same parts at all the same times, but it is true of the plant, and not of the stem, that it might have grown past its stem. Various other examples have been discussed in the literature⁸.

8 See Gibbard (1975), and Lewis, "Counterparts of Persons And Their Bodies" (1968). In Gibbard's example, a (human-shaped) statue and a lump of clay come into and go out of

Actually, one could believe in trans-world individuals, and also think that different things cannot have all the same parts at all the same times, as long as one believes in few enough kinds of individuals. Peter Van Inwagen (1989) holds just this pair of views, because he doesn't think there are such things as stems, or bodies, or lumps of clay, or statues...only simples and living things. But those of us who are ontologically more generous must choose between the view that things with all the same parts at all the same times are identical, and the view that individuals exist in more than one possible world.

Still, given a modicum of ontological generosity, won't we have to deny that (permanent) co-constitution is identity, whether we accept an identitarian or a counterpart-theoretic account of modal predication? If *i*, but not *i'* might have been *F*, then (for the counterpartist) *i* but not *i'* has an *F*-ish counterpart, so *i* and *i'* are distinct.

Since Lewis wants to say that (permanent) co-constitution is identity, he suggests a modification of his original counterpart theory, which allows him to say that (1968). It involves multiple and intersecting counterpart relations, and works like this: one and the same (world-bound) individual can have different sets of counterparts, under different counterpart relations, corresponding to different sortals. For example, it can happen that one and the same world-bound individual is both a person and a body, and has one set of person-counterparts, and a different (though intersecting) set of body-counterparts; or that one and the same world-bound individual, is both a statue and a lump of clay, and has one set of lump-counterparts, a different (though intersecting) set of statue-counterparts. In a sentence such as "This statue might have been *F*", the subject term denotes a world-bound individual, and selects a counterpart-relation--in this case, the statue-counterpart relation. "This statue might have been *F*" will

existence together, and thus have all the same parts at all the same times; the lump, but not the statue, might have survived being squeezed into a ball while the clay was still soft. In Lewis' example, a person and his body come into and go out of existence together, and thus have all the same parts at all the same times; the person, but not the body, might have switched bodies with someone else. (The reader doubtful about the possibility of body-switching may substitute: the body, but not the person, might have outlasted the person).

accordingly be true just in case the individual denoted by “this statue” has a statue-counterpart that is *F*. Similarly, “This lump might have been *F*” will be true just in case the individual denoted by “this lump” has a lump-counterpart that is *F*. If this statue and this lump of clay are (permanently) constituted by all the same parts, then it will be true that this statue = this lump, even though it is true that this lump might have been spherical (because the statue-cum-lump has a lump-counterpart that is spherical), and false that this statue might have been spherical (because the statue-cum-lump has no statue-counterpart that is spherical).

Again, though, these considerations favor counterpart theory only if (permanent) co-constitution is identity. Why suppose that it is? Harold Noonan (1985) has argued that if we don’t, then we shall have to say that purely material things, having the same material constitution at all times are nevertheless distinct, simply in virtue of having different modal, or dispositional, or counterfactual properties. This last view (Noonan holds) «is...surely...astonishing» (1993: 135) it implies that material objects can be *actually* distinct, simply in virtue of the fact that they have different properties in *counterfactual* circumstances. But (Noonan thinks) differences in modal or dispositional properties have to be grounded in differences in “categorical” (non-modal) properties.

Here someone might object that if, say, a statue and a lump of clay are discernible with respect to modal or counterfactual properties, they will also be discernible with respect to some “categorical” properties, such as *being a statue* and *being a lump*⁹. Indeed, it might be maintained, the discernibility of the statue and the lump with respect to modal or counterfactual properties is grounded in their discernibility with respect to *being a statue* and *being a lump*. (It is because the lump is a *lump*, rather than a statue, that it could survive being squeezed into a ball).

9 See Burge (1977: 114). Although Burge does not discuss statues and lumps, or persons and bodies, he considers the case of a positronium atom and a positron-electron pair that have all the same parts at all the same times. After noting that the pair, but not the atom, could have been scattered, he says it would be a mistake to think that the only differences between the positronium atom and the positron-electron pair are modal; the entities in question also differ in kind.

But is it clear that permanently co-composite objects could belong to different kinds? Following Alan Gibbard, let us suppose that our permanently co-composite statue and lump are (respectively) called "Goliath" and "Lumpl". Goliath has everything it takes to be a statue--the right sort of size, shape, origin, and history. But Lumpl has the same size, shape, origin, and history. So why isn't Lumpl a statue too? The identitarian, like the counterpartist, could accept that Lumpl has everything it takes to be a statue--but only at the cost of embracing the (unattractive) view that the sculptor who made Goliath made *two different statues*. The alternative is the (not obviously attractive) view that something with the same composition, size, shape, origin, and history as a statue, may yet not be a statue.

Bracketing the question of whether in the Goliath/Lumpl case, we have two statues, do we really have two things or two (material) objects? It is not obvious that we do. After all, suppose that Goliath is on the table, and weighs one pound. Suppose also (to put it neutrally) that anything on the table that weighs one pound permanently has the same constitution as Goliath. How many one-pound things are there on the table? For the (ontologically generous) identitarian, it would seem, there are (at least) two different one-pound things (two different material objects) on the table--Goliath and Lumpl. But are there? Suppose we take a non-philosopher, show her the table, equip her with a scale, and ask her how many one-pound things there are on the table. It is difficult to imagine her coming up with any answer other than "one". This suggests that the identitarian sees plurality (and difference) where we are naturally disposed to see unity (and sameness).

The identitarian can grant that, when we count how many one-pound things there are on a table, we (ordinarily) count as though permanently co-composite things were identical. Still, he can say, we should not conclude that permanently co-composite things are (in every case) identical. He might say that we are (ordinarily) prone to undercount, although this tendency can be corrected by the right sort of metaphysical education. (In the case at hand, one comes to see that, since Goliath and Lumpl have different persistence conditions, they are two different one-pound things

on the table, so that there are in fact (at least) two one-pound things on the table.) Alternatively, he might borrow an idea originally championed by Peter Geach (1980), and say that counting by identity is just one way of counting. While one may count in such a way that things count as one if and only if they are identical, one may also count in such a way that things count as one if and only if they satisfy some weaker condition--e.g., are permanently co-composite. On this view, someone who says "there is just one one-pound object on the table" is not in error, and what she says is perfectly compatible with Goliath and Lump1's being distinct. Of course, if Goliath and Lump1 are distinct, and each is a one-pound thing on the table, then there are two one-pound things on the table. That is, there are two one-pound things on the table [counting by identity]. Nevertheless, there is just one one-pound thing on the table [counting by some relation weaker than identity, as we are wont to do in non-philosophical contexts]. Finally, the identitarian might say that when we count one-pound things on the table, in that context, "thing" should be understood in a restricted sense, which covers, say, statues, but not bits of clay: given the restriction, there is just one one-pound thing on the table.

Still, the counterpartist may object, what reason is there to think that, in the case under discussion, someone who says there is just one one-pound thing on the table is either undercounting, or counting by some relation weaker than identity, or presupposing a restricted sense of "thing"? Well, the identitarian might say, suppose we modify the case somewhat, so that on the table there is a three-pound statue, made of some bronze that existed before the statue did, and everything on the table that weighs three pounds currently has the same constitution as the bronze statue. Once again, we take a person innocent of metaphysics, show her the table, equip her with a scale, and ask her how many three-pound things are on the table. In this case, just as in the original case, she will probably answer "one". But, the identitarian will continue, it is clear that the statue is not identical to the bronze it is made of, since the statue and the bronze have different origins and histories. So in the modified case, either the person undercounted, or she counted by some

relation weaker than identity, or she was using a restricted conception of "thing". If that is what is going on in the modified case (involving wholly but only temporarily "overlapping" objects)--the identitarian may conclude-- it is incumbent upon the counterpartist to show that it is not what is going on in the original case (involving wholly and permanently "overlapping" objects). Failing that, our practices of counting one-pound objects in a place at a time do not favor the claim that Goliath is Lump (or favor counterparts over trans-world individuals).

In response, the counterpartist might deny that the bronze statue and the bronze it is made of are distinct, on the grounds that two (material) things cannot be in the same place at any (much less all) of the same times. In order to square this with the fact that the bronze existed before the statue did, the counterpartist would have to countenance a less than straightforward account of tensed predication to go along with her less than straightforward account of modal predication. She might, for example, say that in "this statue is bronze", "this statue" denotes a time-bound (as well as world-bound) individual, and selects both a counterpart relation (a way of tracing that (world-bound) individual through logical space) and a "continuer relation" or "temporal counterpart relation" (a way of tracing that (time-bound) individual through time). Then, she could say, "This statue will be [was] *F*" comes out true just in case the individual denoted by "this statue" has an *F*-ish statue-temporal counterpart at a later [earlier] time. And "This bronze will be [was] *F*" comes out true just in case the individual denoted by "this bronze" has an *F*-ish counterpart at a later [earlier] time. That will allow her to say that the statue and the bronze are identical, even though the bronze existed before the statue did.

It is not surprising, though, that Gibbard, Lewis, and most other philosophers who have maintained that permanent co-constitution is identity, have not wanted to say that current co-constitution is identity. For one thing, there is something very odd about the supposition that even though statements such as *this bicycle has existed for twenty years* are true, no individual exists for more than a moment. For another, on the account of (tensed and modal) predication just sketched, there is no such

thing as the history, or the origin of an individual, any more than there is such a thing as the essence of an individual: just as Goliath/Lumpl has one essence *qua* statue, and another *qua* lump, the bronze statue/bronze will have one origin and history *qua* statue, and another origin and history *qua* bronze. And it is very hard to believe that individuals don't have unique origins and histories.

The moral would seem to be that what might be called "counting-based" arguments for the identity of permanently co-composite things are inconclusive.

This leaves the rather different argument to the effect that the identitarian must choose between too many statues, and things that have everything it takes to be a statue, and yet are not statues. The identitarian might try to meet this argument by supposing that in Goliath/Lumpl type cases, there are two statues counting by identity, but only one counting by the relation we normally count by. But he needn't take this (unappealing) line. The identitarian typically distinguishes the 'is' of identity from the 'is' of constitution, and the 'is' of (what I shall call) co-constitution. (For the *locus classicus* of this view, see Wiggins: 1980). In "Hesperus is Phosphorus", the "is" expresses the relation of identity. In "the ring is gold", the "is" expresses a constitution relation (holding between gold and the ring, and more generally between a (kind of) stuff and a thing. In "the copse is a bunch of trees" or "the statue is a piece of bronze" the "is" expresses a co-constitution relation (*being (currently) made of the same matter as*). (Because the 'is' of co-constitution expresses a tensed relation, one and the same copse can be different bunches of trees at different times).

Since he recognizes the co-constitutive 'is', as well as the 'is' of identity, the identitarian can say that anything with the same composition, size, shape, origin, and history as a statue is itself a statue – where 'is' expresses co-constitution, not identity.

Someone who thinks permanent co-constitution is identity runs into the following difficulties in arguing for that view. She can start from relatively uncontroversial premisses (say, about how we count material objects, or material objects of a certain kind; or about whether things just like statues are (in some sense

of 'are') statues). Then--as we have soon--the identitarian can grant the premisses, but insist that the premisses (properly construed) do not entail that permanent co-constitution is identity. Alternatively, she can start from a strong metaphysical premiss, such as the claim that a thing's modal properties supervene on its origin *cum* history *cum* constitution *cum* qualities (so that necessarily, any two things that are indiscernible with respect to their origin, history, constitution, and qualities, are also indiscernible with respect to their modal properties). Then there won't be any daylight between the premisses and the conclusion; but the identitarian will simply reject the premisses. This is not to say that the modal supervenience principle at issue doesn't have a certain intuitive appeal (I think it does); it is just difficult to know what to say to those that find it intuitively unappealing.

On the other hand, there are principles I find initially plausible, which are inconsistent with the modal supervenience principle, and the sufficiency of permanent co-constitution for identity. Suppose we call a fact *soft* just in case its being a fact depends (logically) on how the future goes, and call a fact *hard* otherwise. That I will never walk on the moon is a soft fact; that I have lived in Castello is a hard fact. I find it intuitively very plausible that facts about how many material objects there are in a given place at a given time are hard facts: how can what happens in the future make any difference to how many things are here *now*? Suppose, though, that Goliath is Lumpl, and Goliath is on the table right now. That Goliath is the same material object as Lumpl will be a soft fact (since its factuality depends on whether or not some future event puts an end to Goliath, without at the same time putting an end to Lumpl). So whether we have n or $n + 1$ material objects on the table now will be a soft fact. This unwelcome consequence is avoided if Goliath and Lumpl are conceived of as as different trans-world individuals with difference essences: it will then be a hard fact that Goliath and Lumpl have different modal properties, and are two material objects, rather than one.

Those who would identify Goliath and Lumpl can respond that there is a fact about how many material objects there are on the table now – on the supposition that we are counting by current

total overlap, rather than identity¹⁰. But isn't there a hard fact about how many material objects there are on the table right now, even if we are counting by identity? Aren't Goliath and Lump one *and the same*, or not, irrespective of what happens in the future?

One reason the identitarian is likely to think the answer is obviously 'yes', is that he is likely to believe that material objects are "wholly present", not just at different worlds, but also at different times. Given this (Anselmian)¹¹ picture, it is hard to see how one could avoid thinking that either there are two (currently overlapping) material objects on the table, or just one, and nothing could change that. By contrast, a counterpartist may well suppose that material objects are present at different times by having different temporal parts that are present at those times, in much the way that events are present at different times by having different temporal parts that are present at different times. Given this picture, it is natural to suppose that there can be soft facts about how many material objects there are here now. Suppose a battle is being fought now. It could be that, depending on how the future goes, that battle will turn out to have been part of one war, or two. If so, there won't be a hard fact about how many wars (and how many events) are getting underway now. Similarly, if what is wholly present on the table (or, rather, the current stage of the table) is a temporal part or stage of a (four-dimensional)

10 Notice, though, that if we are counting by current total overlap, then how many material objects there are *hic et nunc* can no more (logically) depend on the past, than it can (logically) depend on the future. But, I have argued elsewhere, how many material objects there are here and now can depend on the past (see my (1997)).

11 The view that material objects are temporal-part-less, and wholly present at all the times they exist, is often associated with Aristotle; and that view certainly does go very naturally with the Aristotelian account of change. As far as I know, though, Anselm is the first philosopher to argue explicitly that material objects and other substances lack temporal parts, and are wholly present at every time they exist, unlike processes or events, that are present at different times by having different parts present at those times. See *Monologion* 21:

If this Nature were to exist as a whole distinctly and successively at different times (as a man exists as a whole yesterday, today, and tomorrow), then this Nature would properly be said to have existed, to exist, and to be going to exist... its lifetime... would not exist as a whole at once but would be extended by parts throughout the parts of time. (1974: 34)

material object, then there won't be a hard fact about whether that stage is a part of two different four-dimensional material objects (a statue, and a lump), or just one (a statue-cum-lump). So there won't be a hard fact about how many material objects are there on the table right now. This suggests that a serious objection to the thesis that permanent co-constitution is identity can be met if, but only if, material objects are not Anselmian.

Peter Van Inwagen has suggested that, in order to defend a four-dimensionalist account of physical objects against a certain objection, the four-dimensionalist must reject an identitarian account of modal predication. I am suggesting that, in order to defend a counterpart theoretic account of modal predication against a certain objection, the counterpart theorist must reject a three-dimensionalist account of persistence. The view that objects are wholly present in different worlds, and the view that objects are wholly present at different times are as it were made for each other, and someone who rejects one is well advised to reject the other.

V.

We have not yet found any arguments decisively favoring counterparts over trans-world individuals. But we have not yet considered what Lewis (and, it seems, Kripke) consider the strongest argument for their view, which could be put as follows:

The identitarian thinks that the person you would have been, had things been different, is just you, and that the would-have-been relation is just identity. (We may speak of the would-have-been relation in order not to prejudge whether the relation is identity or some (qualitatively based or non-qualitatively-based) counterpart relation). But the would-have-been relation does not have the right logical properties to be identity.

Why not? Suppose that a bicycle B_0 is made by a bicycle-maker from (small) parts $P_1, \dots, P_n$ ¹². It seems that the bicycle maker could have made that same bicycle from all the same (small) parts except one (as long as she had put the parts together in the same way, in the same place and time, etc.). On the other hand, it seems that no bicycle maker could have made that very bicycle from a completely different set of (small) parts: a bicycle made from completely different (small) parts would have been a different bicycle¹³. There is a problem, though, about reconciling these two claims. In the actual world G our bicycle-maker makes B_0 from (small) parts $P_1, \dots, P_n$. Since she might have made B_0 from all the same (small) parts but one, there is a possible world H_1 in which she makes B_0 from $P_1, \dots, P_{n-1}, P'_n$ (where P'_n is distinct from but just like P_n ; in what follows, a P'_n will always be distinct from a P_n). Now if she had made B_0 out of $P_1, \dots, P_{n-2}, P_{n-1}, P'_n$, it would then have been true that she could have made B_0 out of $P_1, \dots, P_{n-2}, P'_{n-1}, P'_n$. So, it seems, there is a possible world H_2 in which she makes B_0 from $P_1, \dots, P_{n-2}, P'_{n-1}, P'_n$. Continuing in this way, we reach a possible world H_n in which our bicycle-maker makes B_0 from a set of parts $P'_1, \dots, P'_n$ completely different from B_0 's actual parts $P_1, \dots, P_n$. If there is such a world, then the bicycle maker could after all have made the bicycle that she actually made from one set of parts, from a completely different set of parts. Given that we accept the (intuitively plausible) claim that a bicycle could have had a slightly different set of parts, we seem to be stuck with the (intuitively implausible) claim that a bicycle could have been made from a completely different set of parts¹⁴.

The counterpart theorist can, however, take the better without the bitter. As she sees it, it could have been that B_0 was made from $P_1, \dots, P_{n-1}, P'_n$, since B_0 has a counterpart that was made

12 Assume that all of the parts $P_1, \dots, P_n$ are small, and that the bicycle has no parts disjoint from each of $P_1, \dots, P_n$. Thus even the bicycle's frame can be broken down into many small parts.

13 Here it is important that each of the bicycle's parts overlap with some of (the small parts) $P_1, \dots, P_n$. If, say, the frame and the handlebars did not overlap with any one of $P_1, \dots, P_n$, it could be argued that the bicycle could have (originally) lacked all of $P_1, \dots, P_n$.

14 The example is taken from Chandler (1976), though, as we shall see, the moral Chandler draws from it is quite different.

from $P_1, \dots, P_{n-1}, P'_n$. And it could have been that it could have been that B_0 was made from $P_1, \dots, P_{n-2}, P'_{n-1}, P'_n$, since B_0 has a counterpart that has a counterpart that was made from $P_1, \dots, P_{n-2}, P'_{n-1}, P'_n$. By the same reasoning, it could have been that it could have been that...it could have been that B_0 was made from $P'_1, \dots, P'_n$, since B_0 has a counterpart that...that has a counterpart that was made from $P'_1, \dots, P'_n$. Still, the counterpartist can say, B_0 could not have been made from $P'_1, \dots, P'_n$, since B_0 has no counterpart made from $P'_1, \dots, P'_n$ (even though it has a counterpart that has a counterpart that...has a counterpart made from $P'_1, \dots, P'_n$). The counterpart relation need not be transitive; in fact, if it is similarity-based, one would expect it to not to be (since a lot of little differences can add up to a big difference).

So, the counterpart theorist may say, the bicycle case shows that the would-have-been relation is non-transitive. Corresponding to our sequence of possible worlds $G, H_1, \dots, H_n$, we have a sequence of bicycles $B_0, B_1, \dots, B_n$. Each (non-initial) B_i is the bicycle that B_{i-1} would have been (had H_i been actual); but B_n is not the bicycle B_0 would have been, had H_n been actual. Since the would-have-been relation is non-transitive, it isn't identity.

All this presupposes that a bicycle could not have been made from a completely different set of parts. That assumption is contestable. But I doubt there is much mileage in trying to block the argument by contesting it. The argument depends on the idea that an individual could have been slightly different along a certain dimension (in this case, the parts it was made from and originally made of), but could not have been too different along that dimension. Even if a bicycle could have had been made from a completely different set of parts, there will surely be some individual and some dimension such that that individual could have been slightly different, but not too different, along that dimension. For instance, it seems that a clay statue might have had had an (original) shape slightly different from its actual (original) shape. But it could not have had an (original) shape hugely different from its actual (original) shape. A cup made by a potter could have been originally made of slightly more or slightly

less matter than the amount of matter it was initially made of. But it could not have been originally made of an amount of matter sixteen times greater, or sixteen times less, than the amount of matter it was actually originally made of. Richmond Park might have a slightly different original location--it might have (originally) extended slightly into what is actually Wimbledon Common--but it surely could not have originally been entirely within the city limits of Edinburgh. And so on. Denying all these claims would commit us to an appreciable (and unappealing) degree of hypoesentialism.

Similarly, there is not much mileage in blocking the argument by denying the premiss that a bicycle could have been made from all the same parts except one. If we blocked similar arguments in the same way, we would end up with a version of hyperessentialism--at least with respect to the properties an individual has at its first moment of existence.

Of course, all of the kinds of entities appealed to in arguing for the non-transitivity of the would-have-been relation are at least somewhat ontologically controversial: van Inwagen, for one, would deny the existence of every one of them. But anyone who (like van Inwagen) countenances living beings will face problems concerning the transitivity of the would-have-been relation. After all, it seems possible that someday we shall be able to grow cells individually, and subsequently "assemble" them in such a way as to get an organism. Moreover, it seems that the same organism might have been "assembled" from almost-all-but-not-quite-all the same cells, but couldn't have been "assembled" from a completely different batch of cells. That is enough to call into question the transitivity of the would-have-been relation. (The point may be worth emphasizing, inasmuch as van Inwagen has averred that opposition to, or even skepticism about trans-world identity, is simply the result of one or another sort of confusion). Perhaps we can avoid problems for the transitivity of the would-have-been relation if our ontology includes only simples--say, space-time points *vel similia*. I take it, though, that an ontology this bare is not significantly more attractive than hypo or hyper essentialism.

In any case, most philosophers who have addressed the problem under consideration have conceded the existence of the

entities that give rise to the problem, and rejected both hypoessentialism and hyperessentialism. For example, Kripke (tentatively) accepts that an artefact could have been made from and originally of slightly different parts or matter, and insists that an artefact could not have been made from or originally of completely different parts or matter. So how would Kripke respond to the argument just sketched? Judging from note 18 of *Naming and Necessity*, surprisingly concessively:

...If a chip, or molecule, of a given table had been replaced by another one, we would be content to say that we have the same table. But if too many chips were different, we would seem to have a different one. The same problem can, of course, arise for identity over time.

Where the identity relation is vague, it may seem intransitive; a chain of apparent identities may yield an apparent nonidentity. Some sort of 'counterpart' notion (though not with Lewis' underpinnings of resemblance, foreign country worlds, etc.), may have more utility here. One could say that strict identity applies only to the particulars (the molecules), and the counterpart relation to the particulars 'composed' of them, the tables. The counterpart relation can then be declared to be vague and intransitive. It seems, however, utopian to suppose that we will ever reach a level of ultimate, basic particulars for which identity relations are never vague and the danger of intransitivity is eliminated.

Although Kripke expresses himself rather tentatively ("we would be content to say... we would seem... may seem intransitive... may have more utility"), he seems to concede that there are cases in which what I have been calling the would-have-been relation at least appears to be intransitive, and that a counterpart-theoretic account of modal predication--unlike an identitarian one--can obviously accommodate this (apparent) fact. If it turned out that an identitarian account of modal predication cannot be squared with the idea that a bicycle could have been made from some different parts, but not from completely different parts, or the idea that a table could have been made from an only slightly different portion of wood, but not from a completely different bit of wood, that certainly would strongly favor a

counterpart-theoretic account of modal predication over an identitarian one. And inasmuch as the former account clearly and unproblematically accommodates the (apparent) fact at issue, and the latter account does not, that favors the former account over the latter one¹⁵.

Hugh Chandler has suggested, though (1976), that the identitarian can accept that a bicycle could have (originally) lacked some but not all of its original parts (without contesting the transitivity of identity). His idea is that somewhere in the sequence of worlds that begins with G and ends in H_n , we pass from possible worlds to impossible worlds¹⁶. H_1 is possible relative to G (that is, everything true at H_1 is possibly true at G). And each (non-initial) H_j is possible relative to H_{j-1} . But H_n is not possible relative to the actual world G (although it is possibly possibly....possible with respect to G). It follows that the relation of relative possibility is not transitive, so that any modal logic at least as strong S4 (according to which relative possibility is transitive) is too strong. It will not in general be true that $\Box A \supset \Box \Box A$, or that $\Diamond \Diamond A \supset A$. Moreover, some possible states of affairs are only contingently possible. Suppose, for the sake of simplicity, that a bicycle with n small parts could have been made from a set of small parts fewer than half of which were different from the small parts that bicycle was actually made from, but could not have been made from a set of small parts half of more of which were different from the small parts the bicycle

15 In note 18, Kripke at least suggests that it is the (apparent) vagueness of the would-have-been relation that gives rise to its (apparent) intransitivity. But as Anil Gupta has noted (1980: 102-103) even if we supposed there was a perfectly precise answer to the question, "how many of its original parts could this bicycle have lacked, and still existed?", the would-have-been relation would still be intransitive – unless the answer was "none" or "all". Although the fact that a bicycle could have (originally) lacked some but not all of its original parts provides at least a *prima facie* reason to think that the would-have-been relation is intransitive, it is unclear that it provides a (*prima facie*) reason to think that it is vague (see Salmon 1980: 240-250).

16 Chandler suggests that we do not pass from a (definitely) possible H_j to a (definitely) impossible H_{j+1} ; instead we move from possible worlds to impossible worlds "via a region of indeterminacy" where the worlds are presumably neither definitely possible nor definitely impossible. Thus the relation of relative possibility will have the vagueness, as well as the intransitivity, that some counterpart theorists have wanted to ascribe to the would-have-been relation.

was actually made from. Given that B_0 was actually made of $P_1, \dots, P_n$, it could have been made from $P'_1, \dots, P'_n$. Moreover, B_0 could have been made from $P'_1, \dots, P'_{n/2-1}, P'_{n/2}, \dots, P'_n$. But if it had been, then it wouldn't have been possible for it to have been made from $P'_1, \dots, P'_{n/2-1}, P'_{n/2}, \dots, P'_n$ (since that would entail B_0 's being made from a set of small parts half of which were different from the small parts it was actually made from).

My first reaction to all this was disbelief: it seemed paradoxical that what possibilities there are could depend on which of the possibilities that there are is actualized. Moreover, consider the state of affairs, B_0 's *being made of* $P'_1, \dots, P'_{n/2}, P'_{n/2+1}, \dots, P'_n$. Given our assumptions about how many of its actual parts a bicycle would have to have been made from, this is a (just barely) impossible state of affairs. But it is not, as it were, intrinsically impossible--impossible simply in virtue of being the state of affairs it is. If B_0 had been made from $P_1, \dots, P_{n/2-1}, P'_{n/2}, P'_{n/2+1}, \dots, P'_n$ – as it perfectly well might have been – then B_0 's *being made of* $P'_1, \dots, P'_{n/2}, P'_{n/2+1}, \dots, P'_n$ would have been a (just barely) possible state of affairs. So B_0 's *being made of* $P'_1, \dots, P'_{n/2}, P'_{n/2+1}, \dots, P'_n$ is impossible only because it happens to be just a bit too far from actuality. And how could the (metaphysical) impossibility of a state of affairs be a contingent and extrinsic feature of that state of affairs?

Although I continue to resonate to these intuitions, I am no longer sure how much weight they should be accorded. For the defender of Chandler could say that, however plausible it might be *ex ante* that what is possible does not depend on which possibilities are actualized, or that metaphysical impossibility is an intrinsic feature of a state of affairs, cases involving artefacts show – by appeal to robust and widely shared intuitions – that relative possibility is intransitive, and that both of the (initially plausible) principles just mentioned have counterexamples.

Indeed, two commonly made criticisms of Chandler can be played off against each other. It is often held to be intuitively obvious – or at least very plausible – that (where metaphysical possibility is at issue), whatever is possible is necessarily

possible¹⁷. And it is often said that Chandler's solution to the bicycle problem is *ad hoc*, inasmuch as Chandler provides no independent reason for thinking that relative possibility is intransitive¹⁸. Chandler could reply that (i) the bicycle case and its congeners provide (tolerably) clear counterexamples to the transitivity of relative possibility (sufficiently motivating the denial of transitivity), and (ii) counterexamples to the transitivity of relative possibility are relatively hard to come by – which explains the initial appeal of principles such as “whatever is possible is necessarily possible” (and the related principles discussed in the last paragraph). An analogy: many non-philosophers, and even some philosophers, are initially disposed to regard as intuitively obvious the claim that if this material object and that material object are in the same place at the same time, then this material and that material object are identical¹⁹. But most (contemporary analytic) philosophers deny the claim, and explain its initial appeal in terms of our failure to initially think of the kind of cases that are counterexamples to it. It is open to Chandler to make the same kind of move with respect to the intuitive appeal of the principles incompatible with his treatment of the bicycle case.

Here the champion of counterparts might object:

As Kripke notes, the same sorts of difficulties concerning vagueness and intransitivity arise for identity through time that arise for identity across worlds: a chain of apparent identities can lead to an apparent non-identity. So, for example, a minuscule deformation of the bit of clay a statue will result in the (minuscule) deformation of the statue, and not in the statue's demise. But a long enough series of minuscule deformations will result in the statue's demise, and its replacement by something else²⁰. How is this possible? We cannot answer this question along Chandler's lines (since the temporal analogue of the

17 See, e.g., Plantinga (1974: 51-55).

18 See Gupta (1980: 101).

19 I put the principle this way, rather than in terms of the impossibility of having two material objects at the same place at the same time, to avoid complications concerning whether we are counting by identity or some other relation.

20 I have switched from bicycles to statues, because it is by no means clear that a bicycle could not survive the total replacement of all its original small parts, even if could not have been made from a completely different set of small parts.

relative possibility relation (the earlier-than-or-simultaneous-with-or-later-than relation) is unquestionably transitive. Thus Chandler's approach is unsatisfactory, because it does not provide a general solution to a problem that comes in both a modal and a temporal version.

It is true that Chandler's approach can not be brought to bear on the (superficially similar) statue problem. But what would an account providing a uniform solution to the modal and temporal problems look like? One such account would hold that identity through time is genuine identity, and avoid the paradox by maintaining that some of the relevant alleged (cross-time) identities are false. On this approach, in the statue case described above, we deny that for all times t and $t + \epsilon$, if the clay statue exists at t , and the clay the statue is made of still exists at $t + \epsilon$, and is only minimally deformed with respect to its shape at t , then the statue existing at t = the statue existing at $t + \epsilon$. At some point, when the statue gets too far away from its original shape, a slight deformation of the clay it is made of becomes fatal to it. On an alternative account, we would say that an identity statement may be completely true, or completely false, or something in between. And we would say that some of the relevant identities are less than completely true: at some point, a slight deformation results in its no longer being completely true that the statue survives.

If we treat the temporal case along these lines, a parallel treatment of the modal case would look like this: identity across worlds is genuine identity, and some of the alleged (cross-world) identities are false. We deny that for any worlds H and H' , if the bicycle existing at H differs from the bicycle existing at H' only with respect to one small part, then the bicycle at H = the bicycle at H' . Alternatively, we say that identities may be completely true, completely false, or something in between, and we deny that for any worlds H and H' , if the bicycle existing at H differs from the bicycle existing at H' only with respect to one small part, then it is completely true that the bicycle at H = the bicycle at H' .

I don't find either of these uniform treatments of the modal and temporal problems especially appealing: I am inclined to think the two problems, in spite of their apparent similarity, have different

solutions. But for present purposes, what matters is that the defender of counterpart theory cannot champion either of them, inasmuch as they presuppose that identity across worlds is genuine identity. If the counterpartist wants a uniform treatment of the temporal and modal cases, then she is going to need to suppose that there is a cross-time counterpart relation, as well as a cross-world counterpart relation, and that the cross-time counterpart relation, like the cross-world counterpart relation, is non-transitive. On this approach, in the statue case, we will have a series of times $t_0, \dots, t_n$, and a series of individuals $i_0, \dots, i_n$, where i_0 is a statue, each i_{j+1} is the individual that i_j will be (when t_{j+1} is present), but i_n is *not* the individual that i_0 will be (when t_n is present). The trouble is that, on this approach, it won't in general be true that if it will be that it will be that...this is F , then it will be that this is F . (For it won't follow from the fact that this thing has a (future) "temporal counterpart" that has a (future) temporal counterpart...that is F , that this thing has a (future) temporal counterpart that is F). Surely, though, if a statue is going to be going to be spherical (or going to be going to be going to be spherical, or going to be going to be going to be....spherical), then it is going to be spherical.

To put essentially the same point in a slightly different way, treating the bicycle case and the statue case in a uniform counterpart-theoretic way is unappealing, because whether or not the would-have-been relation is transitive, the will-be relation certainly is. (It is often objected to "memory theories" of personal identity that, on such theories, it can happen that P_3 at t_3 is the person that P_2 at t_2 will be, and P_2 at t_2 is the person P_1 at t_1 will be, but P_3 at t_3 is not the person that P_1 at t_1 will be. Few philosophers would be tempted to defend the memory theory on the grounds that the will-be relation is intransitive).

The moral is that the counterpart theorist is in no position to impugn the adequacy of Chandler's solution to the bicycle problem on the grounds that what is wanted is a solution that can also solve the statue problem: the counterpartist theorist, just as much as Chandler, is committed to the problems' having different solutions.

Having considered four lines of argument for counterpart theory, we can take stock. Kripke and others have argued (to my mind, persuasively) that, in virtue of its simplicity and straightforwardness, an identitarian account of modal predication is the “default” option, in the absence of considerations favoring a counterpart-theoretic one. As Kripke concedes, such considerations might be forthcoming. And some will find some of those considerations (e.g., what I have called the modal supervenience principle) persuasive. But none of the considerations touched upon here are likely to persuade someone initially impressed by the simplicity and straightforwardness of an identitarian account that a counterpart-theoretic account is, all things considered, preferable to it.

King's College London
Strand
London WC2R 2LS
e-mail: chughes3@nd.edu

REFERENCES

- ANSELM ST. (1974). Monologion. In: J. Hopkins & H. Richardson (eds.), *Anselm of Canterbury*. Toronto: Edwin Mellen Press.
- BIGELOW J. & PARGETTER R (1990). *Science and Necessity*. Cambridge: Cambridge University Press.
- BURGE T. (1977). A Theory of Aggregates. *Noûs* 11, 97-117.
- CHANDLER H. (1976). Plantinga and the Contingently Possible. *Analysis* 37, 152-159.
- GEACH P. (1980). *Reference and Generality*. Ithaca: Cornell University Press (3rd edition).
- GIBBARD A. (1975). Contingent Identity. *Journal of Philosophical Logic* 4, 187-221.
- GUPTA A. (1980). *The Logic of Common Nouns*. New Haven: Yale University Press.

- HUGHES C. (1997). Same-Kind Coincidence and the Ship of Theseus. *Mind* 106, 53-67.
- KRIPKE S. (1971). Identity and Necessity. In: M. Munitz (ed.), *Identity and Individuation*. New York: New York University Press.
- KRIPKE S. (1972). Naming and Necessity. In: D. Davidson & G. Harman (eds.), *Semantics of Natural Language*. Dordrecht: D. Reidel: 253-355.
- LEWIS D. (1968). Counterpart Theory and Quantified Modal Logic. *Journal of Philosophy*, 113-126.
- LEWIS D. (1971). Counterparts of Persons and Their Bodies. *Journal of Philosophy*, 203-211.
- LEWIS D. (1973). *Counterfactuals*. Cambridge: Harvard University Press.
- LEWIS D. (1983). Individuation by Acquaintance and Stipulation. *The Philosophical Review*, 3-33.
- LOUX M. (1979). *The Possible and the Actual*. Ithaca: Cornell University Press.
- LYCAN W (1990/1). Two--No Three--Conceptions of Possible Worlds. *Proceedings of the Aristotelian Society*.
- NOONAN H. (1985). The Closest Continuer Theory of Identity. *Inquiry* 28.
- NOONAN H. Constitution is Identity. *Mind*, 133-145.
- PLANTINGA A. (1974). *The Nature of Necessity*. Oxford: Clarendon Press.
- SALMON N. (1982). *Reference and Essence*. Oxford: Basil Blackwell.
- SHOEMAKER S. (1979). Identity, Properties, and Causality. In: P.A. French, T.E. Uehling (Jr.), H.K. Wettstein, (eds.), *Midwest Studies in Philosophy*. Morris: University of Minnesota Press, 321-342.
- VAN INWAGEN P. (1989). *Material Beings*. Ithaca: Cornell University Press.
- WIGGINS D. (1980). *Sameness and Substance*. Cambridge: Harvard University Press.

INDIVIDUATION AND MEREOLOGICAL UNIVERSALISM

Pierdaniele GIARETTA

The claim that composite material objects exist can be supported by providing many examples of ordinary material objects, like tables, cars, dogs, pencils. However, it is well known that many philosophers would deny that some of the ordinary material objects accepted by common sense really exist while others would claim that objects exist which are definitely non-ordinary. Existence of non-ordinary objects follows from a view called "Mereological Universalism" by Van Inwagen, i.e. the thesis that any disjoint objects necessarily compose an object. Such a view is rejected by Van Inwagen in *Material Beings* (1990). His argument has been criticised and taken to be inconclusive. Indeed it fails to prove that the thesis of Mereological Universalism is false, and I will try to account for its failure by arguing that one of its premises lacks sufficient motivation. At the same time I will emphasise that Van Inwagen might provide some reasons for it which depend on thinking that general assertions of existence are justified only under certain conditions which do not seem to be satisfied by the thesis of Mereological Universalism, when this thesis is not read in the way indicated by the disputable premise.

1.

Both the thesis and the rejecting argument are formulated in terms of a notion of composition which Van Inwagen himself introduces. This notion is defined by using plural variables, i.e. variables which are meant to be bound by plural quantifiers, so that composition can be qualified as a multigrade relation, i.e.

– by and large – as a relation between a plurality of entities and an entity. The definition runs as follows:

the X compose y =_{df} the X are all parts of y and no two of the X overlap and every part of y overlaps at least one of the X

where X is used as a plural variable instead of Van Inwagen's "the x s"¹. Thus the composite object turns out to be a fusion of the composing entities, which, according to the given definition, are required to satisfy the constraint of being disjoint. The constraint of being disjoint, which the composing entities are required to satisfy, avoids what Lewis and Varzi would call double counting, i.e. counting a part more than once when overlapping among the composing entities is allowed. It should be stressed that in such definition nothing implies that the composite object, if any, is unique.

The thesis of Mereological Universalism is so formulated:

It is impossible for one to bring it about that something is such that the X compose it, because, necessarily (if the X are disjoint), something is such that the X compose it (Van Inwagen 1990: 74).

Van Inwagen explains that according to this thesis «one can't bring it about that the X compose something because they already do; they do so "automatically"» (*ibidem*). The explanation goes on by pointing out the following analogy:

Just as, according to the theory of sets, there has to be associated with the X a certain abstract object, a set that contains just the X , so, according to the theory we are considering, there has to be associated with the X a certain concrete object, a sum of the X . (*ibidem*).

In other words, just as a set exists if its elements exist, so a composite object exists if the composing entities exist². In the

1 Van Inwagen 1990: 29. Also in the following "the x s" is replaced with " X ".

2 After Boolos' interpretation of second order quantifiers in terms of plural quantification it has become clear that plural variables may and, indeed, should not be understood as set-theoretic variables. It is not even true that, in general, for any X there is a corresponding

former case an abstract entity, in the latter a concrete one, is affirmed as existing, but in both cases its existence is qualified as automatic. Such an analogy is supported by a particular conception of sets, which could be called the combinatorial one and was first specified by Black, in contrast with Cantor's conception:

...Cantor's formula [by a 'set' we understand any assembly into a whole *M* of definite and well-distinguished objects *m* of our perception or thought], stripped to essentials, runs quite simply: "A set is an assembly into a whole of (well-defined) objects". Here, the phrase "assembly into a whole" certainly suggests that something is *to be done* to the elements, in order for the "whole" or "the unified thing", which *is* the set to result. But *what* is to be done...? ...The truth is that once the elements of a set have been identified, *nothing* need or can be done to produce the corresponding set. (Black 1971: 621)

The idea that analogously with the case of sets nothing need or can be done to get the composite object from the composing entities plays a major role in the justification of the most crucial assumption of Van Inwagen's argument. Let us look at the premises of the argument, as listed by Van Inwagen (1990: 75).

The first two are:

- A. I exist now and I existed ten years ago.
- B. I am an organism (in biological sense), and I have always been an organism.

We can say that (A) and (B) introduce an organism *p* which existed at a certain time *t* and has continued to exist until time *t'*, ten years after. Let us neglect the objections which could be raised against (A) and (B)³.

set containing just the *X*. Like Van Inwagen, I will neglect this important logical problem in the following, since it is quite plausible that any logical constraint on the existence of sets would not substantially affect the point of the drawn analogy.

3 Van Inwagen himself does partially neglect such objections.

The third premise states:

- C. Every organism is composed of (some) atoms (or other) at every moment of its existence.

In a note Van Inwagen specifies that he means «chemical atoms – atoms in the modern, scientific sense – not mereological atoms», but he immediately adds that «the argument could as easily be stated in terms of mereological atoms or “simples” as in terms of chemical atoms» (1990: 288). I will assume that mereological atoms are referred to.

The fourth and the fifth premises state, respectively:

- D. Consider any organism that existed ten years ago; all of the atoms that composed it ten years ago still exist.
E. Consider any organism that exists now and existed ten years ago; none of the atoms that now compose that organism is among those that composed it ten years ago.

(C), (D), and (E) are taken to express empirical facts, even if, when referring to chemical atoms, unstable atoms could falsify (D) and a few of the atoms could still compose the same organism after ten years. Van Inwagen neglects such subtleties. So do I.

From (C), (D), and (E) it follows that the object p , introduced by (A) and (B) as an organism existing at t and at t' , is composed by some atoms X at t , such atoms X exist also at t' but at t' p is composed by some different atoms Y .

Let us now introduce the sixth crucial premise

- F. If Universalism is true, then the X cannot ever compose two objects. That is, the X cannot compose two objects either simultaneously or successively. More formally, if Universalism is true, then it is not possible that $\exists y \exists z \exists w \exists v$ (the X compose y at the moment w , and the X compose z at the moment v , and y is not identical with z).

(F) implies that, if Mereological Universalism is true, then the atoms X composing p and only p at t compose p and only p also at t' . But it follows from the other premises that it not true that p

is composed by the atoms X at t' . So the antecedent of the sixth premise is false, i.e. Mereological Universalism is false.

2.

The most controversial assumption is the sixth one as Van Inwagen himself acknowledges. He seems to appeal to the consideration that if, given any objects, however taken or situated, nothing is relevant to the existence of an object composed by them but the existence of the composing entities, then nothing should be relevant to the identity of the composed object but the identity of the composing entities. Such a consideration relies on the analogy with sets, as combinatorially conceived. Let us give a closer look at this analogy in order to see whether and in which sense it holds and then to evaluate whether it really supports Van Inwagen's sixth premise.

Following Pollard's approach (1996) let us introduce the predicate "comp" as standing for a many-one relation of composition in the multigrade sense⁴. "Are all the elements of" could be one of its interpretations, "are all the parts of" another one. The first is the set-theoretical interpretation and the latter is to be understood as a mereological interpretation. In both interpretations it holds that:

$$I. \quad \exists X (X \text{ comp } a \wedge X \text{ comp } b) \rightarrow a=b$$

"Are all the parts of" is not the only possible mereological interpretation of "comp". The predicate "comp" can be interpreted in Van Inwagen's defined sense of sum (y is a sum of the X =_{df} the X are all parts of y and every part of y overlaps at least one of the X). Then (I) still holds and amounts to *Uniqueness* of sum, i.e. extensionality holds for mereological fusions too.

Does Van Inwagen accept (I)? Van Inwagen is committed to (I) when "comp" is interpreted as "are all the parts of", since it follows from two principles he accepts, i.e. reflexivity and antisymmetry of parthood (1990: 55). If the predicate "comp" is

4 "comp" is not Pollard's predicate, but it is used in a similar way as Pollard's predicate.

interpreted in Van Inwagen's defined sense of sum or in Van Inwagen's defined sense of composition, i.e. if (I) is such that it amounts to *Uniqueness* of sum or of composite object, Van Inwagen does not assume it. However *Uniqueness* of sum is a theorem of classical mereology and, if the predicate "comp" is interpreted in Van Inwagen's defined sense of composition, he might see (I) as something implicit in the analogy with sets which, in his opinion, is suggested by the thesis of Mereological Universalism. Indeed, if reference to time is left aside, the sixth premise presents (I) as a consequence of the thesis of Mereological Universalism⁵.

(I) only excludes that some things compose different objects, not that some things compose an object and also some other things compose it. So (I) does not allow us to take the individuality of the composed object as reduced to the individuality of the composing entities. To get such a reduction the following principle is needed:

$$\text{II. } \forall X \forall Y ((X \text{ comp } c \wedge Y \text{ comp } c) \rightarrow X=Y)$$

where " $X=Y$ " is to be taken as an abbreviation of " $\forall z (Xz \leftrightarrow Yz)$ ". (II) implies that different pluralities of things do not compose the same object. Only assuming it together with (I) can it be said that identity questions among composite objects reduce to identity questions among the composing entities.

(II) holds in the set-theoretical interpretation, but not, in general, in the mereological interpretation. However if X and Y plurally stand for atoms, and everything is a sum of atoms, then (II) holds also in the mereological interpretation. So we can conclude that in order to specify the identity of a mereological composite object, nothing else has to be done than to specify its

5 On p. 55 Van Inwagen (1990) seems to accept principles such that when the thesis of Mereological Universalism is added to them, a system is got which is as strong as the system presented by P. Simons (1987: 37). This system is presented by Simons as though including *Uniqueness* as a theorem and so is intended by him as a version of the classical extensional mereology. It follows that Van Inwagen would be formally justified in assuming the sixth premise, at least if time is not taken into account or only existence and uniqueness at a time are taken into account. However it is not so! Simons' axioms are satisfied by the model he describes at p. 28, but this model does not satisfy *Uniqueness*.

composing atoms. This really seems to suggest that nothing else is relevant to the identity of the composed thing than the identity of the composing atoms, in the sense that identity questions among composite objects reduce to identity questions among the composing atoms.

Is this enough to support Van Inwagen's sixth premise? No. In his argument Van Inwagen uses a temporally relativised notion of composition of the form " X compose z at t ". This notion can be taken to be defined by temporally relativising the definition of composition introduced by Van Inwagen. Regarding (II), temporal relativisation leads to:

$$\text{II}_t \quad \forall X \forall Y \forall t ((X \text{ comp } c \text{ at } t \wedge Y \text{ comp } c \text{ at } t) \rightarrow X=Y)$$

(II)_t could be taken for granted, but to justify the sixth premise Van Inwagen needs to show, in particular, that the thesis of Mereological Universalism supports something like:

$$(*) \quad \forall X \forall Y \forall t \forall t' ((X \text{ comp } c \text{ at } t \wedge Y \text{ comp } c \text{ at } t') \rightarrow X=Y)$$

(*) does not logically follow from (II)_t, not even from (II)_t in conjunction with the temporally relativised analogue of the thesis of Mereological Universalism:

$$\text{III. } \forall X \forall t ((\text{disjoint}(X) \wedge \forall z (Xz \rightarrow z \text{ exists at } t)) \rightarrow \exists z (X \text{ comp } z \text{ at } t))$$

where the predicate "disjoint" is to be taken as appropriately defined. It is fairly obvious that (II)_t and (III) do not logically imply (*), either when "comp" is interpreted as Van Inwagen's defined notion of composition⁶. (II)_t and (III) do not conflict with

6 That can be formally shown by interpreting the predicate "comp" as a suitable relation of occupation in such a way that, when X and Y are restricted to atoms, (II)_t expresses the true proposition that if at any time the atoms X and the atoms Y occupy the same region of space, then they are the same and (III) expresses the true proposition that at any time disjoint atoms existing at t occupy a region of space, but, of course, what (*) expresses in this interpretation, i.e. that a region occupied by atoms at different times is occupied by the same atoms, is not generally true.

Objections could be made that the occupation relation does not satisfy all principles adopted by Van Inwagen on the composition relation. Not, for example:

the fact that the object p , as described on the basis of the first five premises, was composed by the atoms X at t and by the different atoms Y at t' . Let us remember that also the atoms X are supposed to exist at t' and, by the thesis of Mereological Universalism, they compose an object at t' , say q . As far as (II₁) and (III) are concerned, it is not excluded that q existed at t and it too was composed by the atoms X . If contingent identity is rejected, hence it is taken that

$$a=b \text{ at } t \rightarrow a=b$$

such a case is excluded by the following temporally relativised version of (I):

$$\exists X (X \text{ comp } a \text{ at } t \wedge X \text{ comp } b \text{ at } t) \rightarrow a=b \text{ at } t$$

Now, let us assume that contingent identity is rejected, so that the temporally relativised version of (I) reduces to:

$$(I_1) \quad \exists X (X \text{ comp } a \text{ at } t \wedge X \text{ comp } b \text{ at } t) \rightarrow a=b$$

(I₁) is implied by the consequent of (F) and so should be supported by (III), i.e. the thesis of Mereological Universalism. How is it possible? Van Inwagen might have an intuitive partial justification. Let us fix a certain time t . If at t nothing but the composing entities is relevant to the identity of an object

If each of the X has a surface and the X compose y , then y has a surface and the surface area of y is less than or equal to the sum of surface areas of the X .

If each of the X has a mass and the X compose y , then y has a mass and the mass of y is the sum of the masses of the X .

If each if the X occupies a region of space and the X compose y , then y occupies the sum of the regions occupied by the X . (Van Inwagen 1990: 44).

However also sets and membership do not satisfy all the principles concerning composition. Thus the objection against the interpretation in terms of occupation could be legitimately raised by one intending to argue for (*) on the basis of (II₁) and (III) only if he is willing to withdraw the analogy with sets as a reason supporting the sixth premise. Another objection, which could be raised from a common sense point of view, is the following: the occupation relation of the occupying atoms to the occupied region is, so to say, more extrinsic than the relation of the composition relation of the composing atoms to the composed object. Things look as though they are related to their composing atoms more strongly than they are to the region they occupy.

composed by them at t , there cannot be at t two objects composed by them. Let us grant the truth of this conditional and of its antecedent⁷. Can this help Van Inwagen in justifying (*), i.e. that a composite object does not change composition?

(I₁) excludes that objects having the same composition at a given time are different, but does not exclude that a composite object changes composition through time, that is it does not exclude that (atomic) compositions, at different times, are different, while the composite objects are identical. For, even if it emerges each time that identity of composite objects is reduced to the identity of the composing entities, since at every time composite objects are the same if and only if their atomic components are the same, nothing is implied with regard to the identity among objects picked up at different times. Analogy with sets could not be appealed to in order to exclude that a composite object can change composition: sets are not temporal entities, indeed the very concept of their existence in time and through time is, at least, something which cannot be immediately grasped or intuited. Thus (I₁) does not provide any support to (*) and assuming the analogy with sets does not appear to be appropriate.

3.

If (I₁) does not provide a full specification of the identity of a composite object and analogy with sets cannot be appealed to in order to complete such specification, how, otherwise, could the identity of a composite object like p be specified? A quick mathematical answer could consist in taking the identity of the composite object as given by a set representing the composition of the object along its whole existence. Such a set should be constituted by couples of the form $\langle n, X \rangle$, where n is a time and

7 The conditional «if at t nothing but the composing entities is relevant to the identity of an object composed by them at t , there cannot be at t two objects composed by them» is not to be confused with the wrong thesis that the temporally relativised version of (I) is a formal consequence of the temporally relativised version of the unrestricted principle of Composition, called summation by Van Inwagen, when parthood is only taken as satisfying assumptions on p. 55 of Van Inwagen (1990). On this point recall note 5.

X plurally stands for atoms, and, if (I₁) is to be satisfied, it should not contain any two different couples having the same first member⁸.

Of course, the quick mathematical answer is not an interesting answer to the individuation question. What we are looking for is a principled understanding of the way in which composite objects are individuated. To this purpose it is natural to look at something which is given together with the composing atoms and which can change through time. Van Inwagen gives an example in terms of blocks, not of atoms, but his example is meant to illustrate a point concerning atoms. Van Inwagen tacitly assumes that for any blocks existing at any given time, their arrangement is unique. In his own words Van Inwagen's example is the following:

Consider an object that is composed of the blocks at t , when they are widely scattered and moving rapidly in relation to one another. How long does it last? Only two answers seem possible. (1) It doesn't last at all; it exists only at t . (2) It lasts as long as its constituent blocks do. Any compromise between these two answers would be intolerably arbitrary: If the blocks "automatically" compose an object, then either any rearrangement of the blocks must destroy that object, or else no rearrangement could destroy it. (1990: 77)

He goes on observing that the former answer has implausible consequences and so he concludes that «Universalism... cannot countenance the supposition that at two different times – or at one time – the X compose two different objects» (*ibidem*).

It follows that the following claim which Rea takes as quite plausible should be rejected:

(α) For some X , what the X compose depends upon how the X are arranged. (Rea 1999: 202)⁹

8 Objections could be made that to specify the identity of a composite object its alternative compositions in different possible worlds also have to be taken into account. It is fairly obvious that an analogous, only more complicated, mathematical answer could be given and so I set this objection aside.

9 Here and below symbols are used which are slightly different from the originals.

McGrath (1998) puts the rejection of (α) down to the acceptance of the following principle:

(P1) If the X automatically compose something at t , there is exactly one object o they automatically compose at t , their sum (McGrath 1998: 118).

where “the X automatically compose something at t ” means “the mere existence of the X at t is necessarily sufficient for there being, at t , a y such that the X compose y ” and “the X automatically compose object o at t ” means “the mere existence of the X at t is necessarily sufficient for the X composing o at t ” (*ibidem*). Rea seems to agree with McGrath that (P1) is accepted by Van Inwagen and takes its acceptance as a mistake. (P1) looks as an instance of the invalid principle that necessarily something being F implies something necessarily being F , and Rea observes that one can hold that «for some X , the mere existence of the X is necessarily sufficient for their composing something but it is not necessarily sufficient for their composing what they now in fact compose» (Rea 1999: 203).

Is Rea right in diagnosing what is going on in Van Inwagen’s argument? Let us recall that already in the comments accompanying the presentation of the thesis of Mereological Universalism Van Inwagen assumed that if nothing is relevant to the existence of a composite object but the existence of the composing entities, then nothing else but the identity of the composing entities is relevant to the identity of the composite object. Afterwards he seems to take into account the hypothesis that something which is always given together with the composing entities contributes to determining the identity of the composed object. Arrangement of the composing atoms is what he mentions in this connection. He refers to a rather simple example, the blocks example, and with reference to it, he makes an assumption which, given the nature of the example, can be generalised in the following way:

G. For any X the arrangement of the X is always relevant to the identity of the object composed by the X or for any X the arrangement of the X is never relevant to the identity of the object composed by the X .

Van Inwagen does explicitly state a consequence of (G) which can be generalised as follows:

- (G') For any composite object, either any rearrangement of the composing atoms destroys the object, or else no rearrangement of the composing atoms destroys it.

The rejection of (α) is not drawn by Van Inwagen from (P1), as McGrath and Rea claim, but from (G') by showing that the alternative that any rearrangement of the composing atoms destroys the composite object is absolutely implausible. Granting that (G) implies (G'), the soundness of Van Inwagen's argument appears to depend on the following two points:

- 1) whether (G'), if not (G), is sufficiently motivated;
- 2) whether the arrangement of the composing atoms is the only other thing which might determine the identity of a composite object.

In reference to the motivation of (G'), Van Inwagen affirms that any compromise between the two alternatives is intolerably arbitrary. It is not clear what Van Inwagen has in mind. I will try to provide some possible reasons for (G'), without claiming that they are what Van Inwagen thinks of when he says that any compromise between the two alternatives is intolerably arbitrary.

Let us suppose that, as a matter of fact, for some *X* any rearrangement of the *X* destroys the object composed by the *X* and that for other *X* no rearrangement of the *X* destroys the object composed by the *X*. Is it really justified to hold *a priori* that in any case a composite object exists, independently of the effect of the rearrangement of the composing entities on the existence of the composite object? Some analytic philosophers think that

- 1) existence cannot be justifiably affirmed without possessing a conception of the identity of the entities affirmed to exist¹⁰ and that

10 A conception of identity does not need and should not be provided by identity criteria. Insofar as identity criteria are framed in a reductivistic perspective, they cannot provide a general notion of identity.

- 2) the way in which a material being comes into existence, goes on existing and ceases to exist characterises its identity.

If such theses are accepted, then, in particular, existence of material beings should be affirmed only in the context of some specification of their persistence through time. Van Inwagen might think that when such a general affirmation of existence is made as the thesis of Mereological Universalism, nothing general and non-arbitrary can be said about the question of identity. For example, if it is said that the composite object exists when, and only when, the composing atoms form a shape satisfying certain requirements, the specification of these requirements appears to be difficult to achieve. Some changes of shape imply a change of object, but which ones? Is it possible to give a unique general answer? On the other hand a criterion based only, for example, on the number of rearrangements, would surely give unacceptable results. In both cases it does not seem possible to identify, *a priori* and in general, those rearrangements, if any, which imply the destruction of the composite object.

Regarding the question whether the arrangement of the composing atoms is the only other thing, besides the composing atoms, which might determine the identity of a composite object, a very old philosophical idea should be taken into account. It is the idea that a composite object is the result of composing certain entities in a certain way. So an object is individuated not only by the composing entities, but also by the manner of composition¹¹. However, this notion can be appealed to in order to admit different entities constituted at the same time by the same composing entities only if mereological extensionality, as formulated by (I_c), is not assumed as valid for the general not qualified relation of composition¹². But Van Inwagen does not take into account the hypothesis that mereological extensionality as formulated by (I_c) might be false, but only the hypothesis that the thesis that a composite object can change composition might be false. So the starting point of Van Inwagen's discussion of (*)

11 Arrangement can surely be taken into account as a manner of composition, but the notion of manner of composition is more general.

12 This relation can be identified with the union of the different relations of composition arising from the different manners of composition.

seems to be that manner of composition can affect only how long the composite object persists through time. Is Van Inwagen justified in taking into account only this possibility? If Van Inwagen grants that arrangement could be relevant to the question of persistence, he should also grant that it is relevant to the question of which objects exist at a given time. It cannot be excluded, without any argument, that there are objects which at a given time are distinguished only by the different ways of persisting.

It may be useful to mention that Fine (1994) presents two manners of composition as generating different kinds of composite entities, i.e. the aggregates and the compounds. Fine's manners of composition are connected with different kinds of persistence of the composite entity – the aggregate exists when and only when at least one of the aggregated entities exists and the compound exists when and only when all of the components exist – and they can affect the same composing entities at the same time, that is at a given time the same entities can compose both a compound and an aggregate. From a three-dimensional point of view, an ordinary material object is neither an aggregate nor a compound¹³. So other coinstantiable manners of composition are also to be admitted. Is the admission of a plurality of manners of composition, possibly constituting the only difference among some material entities existing at a given time, logically consistent with the thesis that the mere existence of any material entities implies the existence of a material entity composed by them? The answer depends on which mereological principles are presupposed. It is not consistent, if a system of principles for classical mereology is assumed. Van Inwagen does not presuppose it, but only weaker principles which are formally consistent with the thesis that there are or there can be different material entities constituted at a given time, or perhaps also at all times, by the same atoms¹⁴.

13 If composite objects are taken to be aggregates, then an ordinary material being should exist before and after the temporal interval of its existence as normally conceived. On the other hand, if composite objects are taken to be compounds, then an ordinary material being should exist for only a temporal subinterval of its existence as normally conceived.

14 See note 5.

However, if different entities exist or can exist which are composed by the same entities, it does not seem justified to affirm, as in Van Inwagen's opinion Mereological Universalism claims, that the existence of the composing entities is the only ground for the existence of a composite entity. Or, at least, the assertion that a composite object exists should go together with a specification of the kind of composition spoken of. The thesis of Mereological Universalism, as formulated by Van Inwagen, excludes that such specification is needed and this can suggest, but does not imply, that there is only one kind of composition. Van Inwagen takes the kind of composition as analogous to the way in which a set, combinatorially conceived, is determined by its elements, and so thinks that, if composition is to be conceived in this way, it is impossible that there be, at any time, different objects composed by the same atoms. Of course, the analogy with sets may be questioned. However, it is not easy to see on which other basis it can be affirmed that for any given entities there is an object composed by them.

To conclude, if the analogy with sets is presupposed by the thesis of Mereological Universalism, Van Inwagen should not grant the possibility that something like the manner of composition affects persistence and so the identity of the composite object: the question has already got a negative answer on the basis of the analogy with sets. However, analogy with sets is not appropriate when concrete persisting entities are spoken of.

On the other hand, if the analogy with sets is not presupposed by the thesis of Mereological Universalism, then Van Inwagen has not proven that if Mereological Universalism is true, then there is at most one way in which entities are composed at a given time, but it should be granted that it is not very clear how the identity of a composite object is in general to be conceived. Assuming that for any composing entities a composite object always exists requires that the identity question be answered, but a general answer does not seem to be possible. When composite objects are not taken as analogous to sets, the difficulty of providing such a general answer may be traced back to the very unrestricted principle of existence of a composite object, relying on the theses that existence cannot be affirmed without a

conception of identity and that conditions of persistence characterise identity.

4.

Wiggins seems to have an answer to the question: what kind of entities does Mereological Universalism affirm as existing? We are told that the composite objects which are affirmed as existing by the unrestricted principle of Composition are special entities, i.e. aggregates, but aggregates are not all the material entities. There are ordinary composite objects which are not aggregates: they only correspond to aggregates, since they are composed of entities which the corresponding aggregates are composed of. So it seems that he would accept the thesis of Mereological Universalism, but he would not read it as conceptually implying that mereological extensionality is valid on the whole domain of material entities on which composition is defined. This seems to block the possibility of providing a general uniform notion of composite object. On the other hand Wiggins is known as the main proponent of the theses that existence cannot be affirmed without a conception of identity and that conditions of persistence characterise identity. In the end does it emerge that Wiggins lacks a unified conception of composite object? Let us first sum up Wiggins' approach to the subject of mereological composition.

Wiggins, who writes before Boolos' introduction of plural quantification, adopts Tarski's definition of sum, and takes aggregates to be sums. Such definition makes reference to the elements of a non-empty class and not to entities plurally quantified over. Wiggins seems to accept the existence of aggregates without any restriction, but does not take *Mereological Extensionality*, i.e. the principle that fusions of the same things are identical, as universally valid. *Mereological Extensionality*, and so the axioms of mereology, are taken as valid on a restricted domain of entities which Wiggins calls collections in mereological sense or aggregates. Since Tarski's definition of fusion also applies to entities which according to Wiggins are outside this

domain, it seems to follow, a bit paradoxically, that mereology does not concern all fusions but only some of them¹⁵.

Wiggins' notion of aggregate cannot be simply provided by the Tarskian definition of sum. For he makes the point that, for instance, the jug and a suitable corresponding collection of bits of china-clay are different entities and only the latter is taken to be an aggregate, even if such bits are part of the jug and no part of the jug is disjoint from each of them and, similarly, such bits are part of the collection and no part of the collection is disjoint from each of them. What is the relation between the two entities? Wiggins (1980: 31) claims that the jug is made of china-clay or is constituted by the collection of china-bits.

Against Shoemaker Wiggins also holds that there is nothing wrong in saying that the statue is composed of a piece of bronze and that the piece of bronze is composed of a statue. This affirmation is made and commented in the following way:

In my usage of “compose” and “constitute”, a usage suggested by mereology (but which I believe to agree with the underlying logic, stripped of pragmatic accretions, to which the expressions conform in English) there is nothing actually wrong in saying what Shoemaker thinks it is not right to say [[the statue] is composed of a *piece* of bronze and ... the piece of bronze is composed of a statue]. And identity is even a special case of constitution. For any x , x compose x . But neither x composes y nor y composes x nor the conjunction of these sentences actually entails “ $x=y$ ”. (Wiggins 1980: 197)

Wiggins can consistently reject *Antisymmetry of Composition*, i.e. the implication:

If x composes y and y composes x , then $x=y$

where “compose” seems to be synonymous with “is a part of”, just because he does not take *Mereological Extensionality* as universally valid. For such an axiom allows us to deduce “ $x=y$ ”

15 Van Cleve (1986: 150-151) emphasises that Wiggins should specify on which grounds it is possible to isolate the fusions which mereology is about, but he also thinks that Wiggins does not provide such a specification. I disagree. Wiggins has at least something positive to say in this regard, as I will point out in the following.

from “ x is a part of y ” and “ y is a part of x ”, and so can justify the definition of “ $x=y$ ” as “ x is a part of y and y is a part of x ”¹⁶.

Rejection of *Mereological Extensionality* entails that it is not possible to speak of *the* entity which is composed by given entities: composition is not enough to identify only one entity. Then what is the relation between composition and identity according to Wiggins? In the above quoted passage Wiggins says that identity is a special case of constitution and this could suggest that “constitutes” has to be taken to mean “is a proper part or is identical to”. However adopting this definition allows to prove *Antisymmetry of Composition* very easily and then to prove also *Mereological Extensionality* from the other principles of classical mereology. So it seems that it is not possible for Wiggins to accept the suggested definition of constitution. For he holds that the collection of bits of china-clay constitutes the jug without being a proper part of the jug, but the above definition would imply that the collection is identical to the jug.

The relation between the material entities which are aggregates and the material entities which are ordinary objects is not so clear. The same entities can compose an aggregate and an ordinary object, the aggregate and the ordinary object being different objects, but it is at least puzzling that the primitive mereological relation of part-of is assumed to apply both to aggregates and to ordinary objects. Aggregates and ordinary objects of different kinds are distinguished by their different ways of persisting, and Wiggins identifies what grounds the way of persisting of aggregates with the existence of the composing entities. For he says:

...the whole rationale of their [i.e. Lesniewski or Goodman] theories of aggregates was the unique determination of aggregates by their constituents and their exhaustion by these. No other mode of determination is provided for, or conceivable. (If another mode were devised it would be a principle of individuation for substances, or wholes not necessarily subject to mereology). (1986: 307)

16 See Wiggins 1980: 93.

However Wiggins does not claim that nothing but the composing entities is relevant to the identity of any object composed by them, as the thesis of Mereological Universalism seems to involve in Van Inwagen's reading. That concerns only the entities called aggregates. So Wiggins can reject that any given entities compose at most one thing: for every ordinary thing there is an aggregate, a different entity, which is composed of the same atoms.

Leaving aside the non-principled way in which the restriction is introduced, a problem arises from thinking of aggregates as concrete entities which we can be causally related to. To think of the atoms composing an ordinary object and also as composing an aggregate one needs to abstract from the relations they have with each other, but it is doubtful that a concrete entity can result from such an abstraction, over and above the composing atoms. Could it be claimed that no other entity is affirmed to exist over and above the composing atoms, more or less as suggested by Lewis' thesis of Composition as identity? This would imply that the unrestricted principle of Composition is literally false or, at least, that speaking of sums or fusions can be dispensed with. Perhaps aggregates are only reifications of pluralities which enhance brevity and expressive power in ontological discourse. Wiggins himself does not really seem to need them in order to propose his conception of a continuant as an entity provided with a way of persisting which it can share with other entities.

Dipartimento di Filosofia
P.za Capitaniato 3
35139 Padova (Italia)
 e-mail: pierdaniele.giaretta@unipd.it

References

- BLACK M. (1971). The Elusiveness of sets. *Review of Metaphysics* 24, 614-636.
 BOOLOS G. (1984). To Be is to Be a Value of a Variable (Or to Be some Values of some Variables). *Journal of Philosophy* .

- FINE K. (1994). Compounds and aggregates. *Noûs* 28, 137-158.
- LEWIS D.K. (1991). *Parts of Classes*. Oxford: Oxford University Press.
- MCGRATH M. (1998). Van Inwagens's Critique of Universalism. *Analysis* 58, 116-121.
- POLLARD S. (1996). Sets, Wholes, and Limited Pluralities. *Philosophia Mathematica* 3, 42-58.
- REA M. (1998). In Defence of Mereological Universalism. *Philosophy and Phenomenological Research* 58, 347-360.
- REA M. (1999). McGrath on Universalism. *Analysis* 59, 200-204.
- SIMONS P. (1987). *Parts: a Study in Ontology*. Oxford: Clarendon Press.
- VAN CLEVE J. (1986). Mereological Essentialism, Mereological Conjunctivism, and Identity Through Time. *Midwest Studies in Philosophy* 11, 141-156.
- VAN INWAGEN P. (1990). *Material Beings*. Ithaca: Cornell University Press.
- VARZI A. (2000). Mereological Commitments. *Dialectica* 54, 283-305.
- WIGGINS D. (1979). Mereological Essentialism: Asymmetrical Essential Dependence and the Nature of Continuants. *Grazer Philosophische Studien* 7/8, 297-315.
- WIGGINS D. (1980). *Sameness and Substance*. Oxford: Blackwell.

L'UNIVERSALITÉ ET L'INCOMPLÉTUDE DE LA MÉRÉOLOGIE EXTENSIONNELLE CLASSIQUE*

Andrea BOTTANI

I. Une chose appelée «méréologie extensionnelle classique»

Depuis Simons 1987, il est devenu habituel de parler d'une chose appelée «Méréologie extensionnelle classique» (dorénavant plus brièvement M.E.C.). Il s'agit en grande partie d'une idéalisation de quelques traits généraux et communs de certains systèmes formels de calcul des parties et des tous essentiellement équivalents et intertraduisibles – la *Méréologie* de Lesniewski et le *Calcul des Individus* de Leonard et Goodman – présentés entre la fin des années 20 et le début des années 40¹. Il y a plusieurs manières de séparer ce qui est essentiel, constitutif et déterminant d'une méréologie extensionnelle classique de ce qui ne l'est pas. Selon Lewis 1990, qui parle tout simplement de «méréologie», les assomptions essentielles sont les trois suivantes:

1. *Transitivité*

Si x est une partie de quelque partie de y , alors x est une partie de y .

* Ce travail fait partie d'une recherche intitulée «Méréologie et Individuation» financée par le *Fonds National Suisse de la Recherche Scientifique*, subside n° 1114-052326.97/1. Je désire remercier Denis Miéville, Pierre Joray et Nadine Gessler pour les nombreuses discussions qui ont influencé ce travail.

1 Voir Lesniewski 1927-30, Leonard & Goodman 1940.

2. *Composition non restreinte*

Chaque fois qu'il y a des choses, il y a aussi une fusion de ces choses.

3. *Unicité de la composition*

Il ne se passe jamais que le mêmes choses aient deux fusions différentes².

2. est aussi connue comme «principe de somme» ou «principe d'existence de la composition», 3. peut être considérée comme strictement liée au bien connu *principe d'extensionnalité méréologique*:

4. N'importe quels x et y sont identiques si et seulement s'ils ont exactement les mêmes parties³.

2 Lewis 1991: 74. Les assumptions 1-3 peuvent être formulées en symbolisme comme suit:

1. $(\forall x)(\forall y)(\forall z) (x < z \wedge z < y) \rightarrow (x < y)$;
2. $(\forall x)(\forall y)(\exists z) (z = y + x)$
3. $(\forall x)(\forall y)(\forall z) z = x + y \rightarrow (\forall j) (j = x + y \rightarrow j = z)$,

où «<» signifie «est une partie de», au sens où mon nez est une partie – non exhaustive – de ma tête et ma tête est une partie – exhaustive – de moi-même, et «+» signifie «est une fusion de», au sens où un arbre est la fusion de ses racines, de son fût et de ses branches. Un problème concerne le fait que les renoncés formels 2. et 3., différemment de leurs analogues informels mentionnés hors de note, emploient une relation rigidement binaire de somme. Ce problème, volontairement mis de côté ici, peut être facilement résolu en employant des quantificateurs pluriels.

2. et 3. semblent concerner la notion de fusion plutôt que celle de partie. En fait, la première notion est définissable en les termes de la deuxième. Une entité étant la fusion d'un nombre d'entités si et seulement si elle a comme parties toutes ces choses, et toutes ses parties partagent des parties avec ces choses. Cela entendu, 2. et 3. peuvent être formulées de manière équivalente sans employer la notion de fusion mais seulement celle de partie, comme suit:

2'. Chaque fois qu'il y a des choses x et y , il y a *au moins une* chose z telle que x et y font partie de z et rien ne peut partager une partie avec z sans partager une partie avec x ou bien avec y .

3'. Chaque fois qu'il y a des choses x et y , il y a *au maximum une* chose z telle que x et y font partie de z et rien ne peut partager une partie avec z sans partager une partie avec x ou bien avec y .

3 En symbolisme: 4'. $(\forall x)(\forall y)((x=y) \leftrightarrow (\forall z) (z < x \leftrightarrow z < y))$. D'une façon informelle, le lien entre 3. et 4. peut être éclairci ainsi. Puisque des entités x et y ont exactement les mêmes parties si et seulement si elles sont la fusion des mêmes entités, alors 3. dit que des choses différentes ne peuvent jamais avoir exactement les mêmes parties; autrement dit,

Lewis traite 1.-3. comme des axiomes, mais il est aussi possible de dériver au moins quelques-uns d'entre eux d'une base axiomatique différente⁴. Ce qui importe selon Lewis, c'est qu'il s'agit d'assomptions qui, sous forme d'axiomes ou de théorèmes, doivent être incluses dans tout système logique pour qu'il puisse être traité comme une version de M.E.C. De ce point de vue, donc, si on appelle C la conjonction de 1.-3., on peut dire que C constitue le *noyau essentiel* de M.E.C.

Deux questions philosophiques principales peuvent se poser à propos de C. Premièrement, on peut s'interroger sur la validité universelle de C. C est-elle une caractérisation de la relation de partie comme elle s'applique à toute entité ou seulement à quelque champ partiel d'entités? Est-elle valable pour tout ce qui existe? Serait-elle vraie si les variables x , y et z étaient interprétées comme ayant un domaine absolument non restreint (c'est-à-dire le domaine universel)? Deuxièmement, on peut s'interroger sur le caractère exhaustif de C. C est-elle capable d'impliquer tout ce qu'une théorie méréologique devrait impliquer des relations entre les parties et les tous? Est-elle capable de donner une image exhaustive de la logique des relations méréologiques⁵?

si des choses ont exactement les mêmes parties, alors elles sont identiques. Ce qu'il reste à montrer pour obtenir 4. est l'autre direction de la conditionnelle: si des choses sont identiques, alors elles ont exactement les mêmes parties. Dans un sens général, cette conditionnelle n'est rien d'autre qu'un exemple de la loi de Leibniz, donc valable indépendamment de la logique spécifique de la relation de partie. Dans un sens méréologique plus spécifique, toute chose doit être partie d'elle-même en vertu de la même notion de partie. Donc, chaque fois qu'un x est identique à un y , x doit être partie de y et y doit être partie de x . Mais cela ne peut pas se passer si x et y n'ont pas les mêmes parties. D'un point de vue formel, la démonstration de 4'. dépend du choix d'un système formel particulier, avec certaines notions méréologiques primitives et certains axiomes.

4 À ce propos, voir Simons 1987, 25-41.

5 Il est tout à fait évident que ces deux questions ne sont pas absolument semblables: on peut dire quelque chose d'universel à propos d'une relation sans dire n'importe quoi d'exhaustif à son propos. Par exemple, dire «pour n'importe quelles entités x , y et z , dire si x est frère de y et y est frère de z , alors x est frère de z », est tout à fait universel mais absolument pas exhaustif à propos de la relation de fraternité. Il est universel parce qu'il vaut pour n'importe quels objets x , y et z pouvant entrer dans la relation de fraternité. Il n'est pas exhaustif parce qu'on peut très bien savoir que la relation de fraternité est transitive sans savoir absolument de quelle relation il s'agit.

Si C n'est pas universelle, alors C est fautive des entités d'au moins quelque domaine particulier. Dans ce cas, puisque C est une conjonction, on peut essayer de localiser l'erreur dans un conjoint ou un autre et de restaurer la vérité universelle de C en l'abrégant, c'est-à-dire en coupant le conjoint (ou les conjoints) faux. De ce point de vue, C dit *trop* et le problème consiste à la réduire d'une façon ou d'une autre. D'un autre côté, si C n'est pas exhaustive, alors C ne dit pas *tout* à propos des relations partie-tout, elle ne suffit pas à caractériser complètement ces relations, donc elle dit *trop peu*. Comme il est clair qu'on peut dire en même temps *trop* et *trop peu* de n'importe quoi (on peut en même temps dire quelque fausseté et ne pas dire toute la vérité à propos de n'importe quoi), C pourrait même être non universelle et non exhaustive. Les trois autres possibilités demeurent aussi tout à fait possibles: abstraitement, C pourrait être non exhaustive et non universelle, ou exhaustive et non universelle, ou universelle et non exhaustive, ou universelle et exhaustive.

À ce propos, il est bien connu qu'il existe un dissentiment philosophique assez radical. Les positions les plus connues sont les suivantes:

- (i) C n'est *ni exhaustive ni universelle*. C vaut comme exhaustive et universelle à l'intérieur d'un domaine d'entités «méréologiquement constantes» qui ne peuvent pas changer leurs parties (il s'agit des entités quadri-dimensionnelles douées de parties temporelles, comme par exemple les promenades, des entités tridimensionnelles non persistantes, s'il y en a quelques-unes, et des entités tridimensionnelles persistantes qui ne peuvent pas changer de constitution, comme par exemples, peut être, les électrons)⁶. Mais C n'est ni exhaustive ni universelle dans tout domaine comprenant nos

6 Une entité peut être traitée comme *quadri-dimensionnelle* si elle est douée de parties temporelles, comme *tridimensionnelle* si elle est dénuée de telles parties. Une entité quadri-dimensionnelle occupe l'espace de la même façon que le temps: elle s'étend dans le temps en ayant des parties temporelles différentes à des instants différents, juste comme elle s'étend dans l'espace en ayant des parties spatiales différentes dans des lieux différents. Au contraire, une entité tridimensionnelle s'étend dans l'espace mais pas du tout dans le temps: elle existe à chaque instant particulier dans sa totalité, en «traversant» le temps sans s'étendre en lui. Les entités tridimensionnelles persistantes ont été souvent appelées «continues».

«continus» ordinaires (tels que les chats, les arbres et les personnes). D'un côté, C ne dit pas comment la relation méréologique entre un chat et ses parties se comporte, d'un autre côté ce qu'elle dit à propos des relations de partie est tout simplement faux de la relation entre un chat et ses parties. Donc, C ne dit pas assez et en même temps elle dit trop, elle est le noyau d'une théorie méréologique à la fois trop faible et trop forte. Cette position a été typiquement défendue, parmi d'autres, par Wiggins et Simons⁷.

- (ii) C est *quasi-universelle et quasi-exhaustive*. Elle est valable pour toutes les entités (mais non pas pour les quasi-entités) et dit tout ce qui doit être dit des relations de partie entre les entités (mais non pas des relations de partie entre les quasi-entités). Cette thèse a été typiquement soutenue par Chisholm. Chisholm appelle les entités *entia per se* et les quasi-entités *entia successiva*. Un *ens successivum* est tout à fait dénué d'existence propre. Il s'agit tout simplement d'une certaine série de entia per se doués d'existence propre et absolument conformes à 1.-3. Nos objets ordinaires (par exemple les chats) sont typiquement des *entia successiva*.
- (iii) C est tout à fait *universelle et exhaustive*. D'un côté, C dit tout ce qui doit être dit de la relation de partie, de l'autre elle est ontologiquement neutre, étant valable pour tout ce qui existe ou pourrait exister: les entités quadri-dimensionnelles, les entités tridimensionnelles non persistantes, et les entités tridimensionnelles persistantes. Cette thèse a été soutenue par Quine et Lewis⁸.

Dans ce qui suit, je défends encore une autre position. Ma thèse sera que C n'est pas exhaustive (elle ne suffit pas à donner une logique exhaustive des relations de partie), mais à l'encontre des apparences, elle est tout à fait universelle, étant valable pour tout ce qui existe, y compris les entités ordinaires du sens commun. Les assumptions 1.-3. restent tout à fait vraies, quel que soit le domaine associé aux variables x , y et z . Par conséquent, C doit être complétée, mais absolument pas réduite.

7 Voir Wiggins 1980; Simons 1987.

8 Voir Quine 1976; Lewis 1986, 1991.

II. Un dilemme méréologique familial

Les systèmes classiques de la méréologie extensionnelle avaient été originellement conçus pour traiter une relation atemporelle et amodale de partie. Mais lorsqu'on sort d'un univers temporellement et modalement aplati et qu'on essaie de tenir compte de notre univers ordinaire fait de chats, d'arbres et de personnes, ces systèmes semblent demander des corrections étendues et profondes. Il est bien connu que le problème principal (mais peut être pas le seul) concerne l'assomption 3. de C, et plus généralement le principe d'extensionnalité méréologique, c'est-à-dire la thèse que n'importe quels objets x et y sont les mêmes si et seulement s'ils ont exactement les mêmes parties (voir le principe 4. susmentionné). 4. fonctionne comme un théorème dans tout système classique de la méréologie, mais il ne semble pas être valable pour nos objets ordinaires, temporellement et modalement épais. Moi-même, par exemple, ne semble pas obéir au principe d'extensionnalité méréologique, parce que je suis «méréologiquement inconstant»: je perds et gagne continuellement des parties au cours de mon métabolisme sans cesser d'exister et tout ce qui fait partie de moi ne fait pas partie *nécessairement* de moi (pour donner un exemple facile, il y a certainement des mondes possibles où je – juste moi-même – n'ai pas à ce moment une molécule qui en ce moment m'appartient dans le monde réel). D'un autre côté, tout réarrangement de mes parties matérielles n'est pas encore moi-même, comme Monsieur Guillotin le savait très bien. Donc je n'ai pas de bonnes manières méréologiques, tout au contraire mon impolitesse méréologique semble être tout à fait radicale. En effet, le principe d'extensionnalité méréologique a la forme d'une équivalence et je ne m'en tiens ni à une direction de l'équivalence ni à une autre: je peux être identique à quelque chose qui n'a pas les mêmes parties que moi et je peux être différent de quelque chose qui a exactement les mêmes parties que moi.

Le fanatique de l'ordre méréologique peut réagir au désordre méréologique de la même façon que le fanatique de l'ordre social réagit au désordre social: en envoyant au cimetière les responsables du désordre, c'est-à-dire en niant toute existence (ou toute existence *propre*) à presque tout notre univers ordinaire fait de

chats, d'arbres, de personnes et de bâtiments. Le fanatique de la liberté méréologique peut réagir à l'ordre méréologique de la même manière que le fanatique de la liberté sociale réagit à l'ordre social: en niant les règles de cet ordre même. Mais parfois, dans l'accrochage entre règles et subversion, on peut commencer à entrevoir les contours généraux d'un ordre moins rigide. Et cet ordre pourrait être le suivant: a) le principe d'extensionnalité méréologique demeure valable pour les sommes méréologiques (entités quadri-dimensionnelles incluses), mais doit être nié pour les «continus» ordinaires tels que moi-même; b) tous les autres principes de C demeurent universellement valables.

Le problème de cette manœuvre est qu'elle semble tout à la fois trop forte et trop faible. Trop forte parce que l'extensionnalité méréologique a une force intuitive considérable *en elle-même*. En effet, il semble y avoir une tension entre la thèse que les objets ordinaires ont des parties et la thèse qu'ils ne sont pas (identiques à) la somme de ces parties. Il n'y a pas de sens à dire: «A a des parties, c'est vrai, mais elle n'est pas la somme de ses parties» de la même manière qu'il n'y a pas de sens à dire «A a des membres, c'est vrai, mais elle n'est pas la classe de ses membres». D'un autre côté, si on sait ce qu'est une partie et ce qu'est une somme, *on voit* qu'une somme A et une somme B ne peuvent pas avoir les mêmes parties tout en étant des sommes différentes, ou avoir des parties différentes tout en étant la même somme (comme *on voit* qu'un ensemble A et un ensemble B ne peuvent pas avoir les mêmes membres tout en étant des ensembles différents, ou avoir des membres différents tout en étant le même ensemble). Donc, si A et B ont des parties, alors A et B sont identiques si et seulement si A et B ont les mêmes parties.

Donc, nier l'extensionnalité méréologique pour les continus semble trop fort. Mais nier l'extensionnalité méréologique *uniquement* pour les continus (et non pour les sommes) semble trop faible, parce qu'insuffisante à rendre compatible C avec une ontologie de continus. En effet, l'assomption 2. de C – *le principe de somme* – dit que, pour n'importe quels objets *a* et *b*, il y a quelque chose qui est la somme de *a* et *b*. Donc, il y a des sommes de continus. Prenons alors Clincoln. Clincoln est la somme (spatio-temporellement éparpillée) de Lincoln et Clinton. Étant

une somme, Clincoln ne peut perdre aucune de ses deux parties (ni Lincoln ni Clinton) ni en gagner d'autres (différentes de Lincoln ou de Clinton). Mais le fait est que chacune des parties de Clincoln (c'est-à-dire Lincoln ou Clinton), ayant un métabolisme, perd et gagne une grande quantité de parties. Or, pour l'axiome de transitivité, tout ce qui fait partie de quelque partie de Clincoln fait partie de Clincoln. Donc, Clincoln peut perdre et gagner des parties. Mais Clincoln est une somme (de continus) et non pas un continu. Je vois trois solutions possibles. *a*) Abandonner le principe d'extensionnalité méréologique pas seulement pour les continus mais également pour les sommes méréologiques (au moins pour les sommes méréologiques de continus). *b*) Abandonner l'axiome de transitivité, selon lequel tout ce qui fait partie d'une partie de x fait partie de x . *c*) Abandonner le principe de somme, au moins pour les continus, et suivre l'idée qu'il y a des choses telles que rien ne soit leur somme (c'est-à-dire nier «l'universalisme méréologique»). Peut-être que la dernière solution est la plus plausible. En effet, admettre qu'une somme méréologique *a* puisse être identique à une somme méréologique *b* même si *a* et *b* n'ont pas exactement les mêmes parties, semblerait dangereusement conduire à soutenir que les sommes ne sont pas des sommes. Quant au réquisit de transitivité, il semble profondément enraciné dans notre notion intuitive de partie.

Il existe une autre raison de penser que la négation de l'extensionnalité méréologique ne peut pas être arrêtée au domaine des continus. Il est assez commun d'observer que, même si les continus avaient des parties temporelles, ils ne seraient pas *identiques* à des sommes de parties temporelles, parce que les continus ont des propriétés modales qu'aucune somme méréologique ne peut avoir – par exemple ils auraient pu durer moins, c'est-à-dire manquer d'avoir des parties temporelles qu'en effet ils ont, tandis que les sommes méréologiques n'auraient pas pu. Mais si cela est vrai des continus, alors cela doit être vrai aussi des événements et des processus. Un processus a également des propriétés modales différentes de la somme de ses parties temporelles: un match de football, par exemple, aurait pu durer quelques minutes de plus ou de moins (si l'arbitre en avait décidé autrement) ou une maladie aurait pu durer plus ou moins (si la thérapie avait été diffé-

rente). Tout ceci montre qu'il n'est pas facile du tout de maintenir le principe d'extensionnalité méréologique pour les processus, tout en l'abandonnant pour les continus. En d'autres termes, il n'est pas facile de nier la validité universelle du principe d'extensionnalité méréologique sans passer directement à l'affirmation de son *invalidité* universelle.

Une façon des plus brillantes de résoudre ce dilemme méréologique (abandonner notre conception intuitive de partie ou notre univers de tous les jours) serait de montrer qu'il n'y a aucun dilemme: réconcilier l'ontologie ordinaire à l'extensionnalité méréologique en démontrant qu'il n'existe aucune opposition entre les deux. Tant le quadri-dimensionnalisme de Lewis que la théorie des *entia successiva* de Chisholm peuvent être vus comme des tentatives de ce type⁹. Cependant, en laissant de côté le problème de l'évaluation de ces positions, plusieurs philosophes pensent que cela reviendrait à se débarrasser de notre ontologie de tous les jours au lieu de la rendre compatible. Bien que cela semble plus facilement acceptable à propos de la théorie des *entia successiva* qu'à propos du quadri-dimensionnalisme, j'ignorerai ce genre de questions. Au contraire, j'emploierai ce qui me reste de temps pour explorer brièvement une stratégie différente dans le but de réconcilier l'extensionnalité méréologique avec une ontologie d'entités tridimensionnelles. Cette stratégie ne demande aucune modification de notre ontologie de continus et aucune réduction de la base axiomatique de M.E.C. Elle se fonde essentiellement sur l'idée que, lorsqu'on entre dans un univers temporellement et modalement épais tel que celui que nous habitons, la relation absolue ou éternelle de partie devient ni vraie ni fausse de plusieurs couples de choses.

III. Méréologie extensionnelle classique et bivalence

Il y a deux façons de se demander si certains individus spatio-temporellement épais – par exemple les occupants de cette salle – satisfont ou non le principe d'extensionnalité méréologique. D'un côté, on peut se demander si dans cette salle il y a des individus x

9 Voir par exemple Lewis 1986; Chisholm 1973.

et y tels que le couple ordonné $\langle x, y \rangle$ se trouve à un instant dans l'extension du prédicat absolu «être une partie de» et à un autre en dehors de cette extension. De l'autre, on peut se demander s'il y a dans cette salle des individus x et y tels que le couple ordonné $\langle x, y \rangle$ se trouve dans l'extension du prédicat temporel «être une partie de à t' » et en dehors de l'extension du prédicat «être une partie de à t'' ». Dans le premier cas, on demande si un prédicat de M.E.C. est vrai ou faux des occupants de cette salle, mais, dans le deuxième, on demande si certains prédicats de la *méréologie élargie* (M.E.C. plus un appareil d'opérateurs modaux et temporels) sont vrais ou faux des occupants de cette salle.

Je commencerai par répondre à la première question. Elle seule concerne la M.E.C. et la notion absolue ou éternelle de partie. D'ailleurs, le principe d'extensionnalité méréologique tire sa force intuitive justement de la notion absolue ou éternelle de partie et c'est dans sa formulation absolue que le principe semble tout à fait irrévocable (et donc difficilement sacrificable à une ontologie de continus ordinaires). Admettons donc que dans cette salle il y ait des individus x et y tels que le couple $\langle x, y \rangle$ se trouve à l'instant t' dans l'extension du prédicat absolu «être une partie de» et à l'instant t'' en dehors de cette extension. Alors, à l'instant t' il y a dans cette salle une chose qui a une certaine partie et à l'instant t'' une chose qui n'a pas cette partie et ces choses sont la même chose. Donc dans ce cas 4. – le principe d'extensionnalité méréologique – ne serait pas vrai des occupants de cette salle, c'est-à-dire qu'il serait faux si ses variables sont interprétées sur un domaine comprenant les individus de cette salle. (Bien sûr, ce n'est pas le seul cas où 4. est faux des occupants de cette salle. Il serait faux même si dans cette salle il y a un x et un y tels que $x \neq y$ et pour tous les z « z fait partie de x » est vrai à t' si et seulement si « z fait partie de y » est vrai à t'' , mais pour des raisons de simplicité je laisserai cela tout à fait de côté).

Je ne crois pas que le principe d'extensionnalité méréologique soit faux des occupants de cette salle parce que je ne pense pas qu'il ait, dans cette salle, un x et un y tels que le couple ordonné $\langle x, y \rangle$ se trouve à un instant dans l'extension du prédicat absolu «être une partie de» et à un autre en dehors de cette extension. Prenons, par exemple, moi-même et une molécule de sucre qui

fait partie de moi à un certain instant de mon histoire métabolique. Et demandons- nous si le couple <cette molécule, moi-même> est dans l'extension de la relation absolue, temporellement non qualifiée, *être partie de* à cet instant de mon histoire métabolique. Traitons cette question de manière différente: supposons que l'on dise que je suis dans l'extension de la relation absolue, temporellement non qualifiée, d'*être partie de*. Est-ce que cet énoncé serait vrai maintenant? Ce qui semble tout à fait clair est que, si l'on dit cela maintenant, ou hier, ou demain, ou le siècle passé, on ne dit rien de vrai ni rien de faux: «cette molécule être une partie de moi» n'est ni vrai ni faux, et cet énoncé reste ni vrai ni faux quel qu'il soit l'instant auquel on le prononce. En effet, le couple <cette molécule, moi-même> n'est ni à l'extérieur ni à l'intérieur de l'extension de la relation absolue de partie. Ce que le passage à un univers temporellement et modalement épais comporte pour la relation absolue de partie, c'est tout simplement le fait que plusieurs couples ordonnés d'entités d'un tel univers ne tombent ni à l'intérieur ni à l'extérieur de cette relation¹⁰. Donc, à l'encontre

10 J'ai soutenu que le couple <cette molécule, moi-même> ne tombe ni à l'intérieur ni à l'extérieur de l'extension de la relation *être partie de* (et donc ni à l'intérieur ni à l'extérieur de la relation *n'être pas partie de*). Autrement dit, soit l'énoncé «cette molécule est partie de moi-même», soit sa négation ne sont ni vrais ni faux. En fait, je ne crois pas que cette thèse soit indéniable, parce qu'il existe au moins deux autres approches tout à fait incompatibles.

1. On peut soutenir que le couple <cette molécule, moi-même> tombe en dehors aussi bien de l'extension de la relation *être partie de* que de l'extension de la relation *n'être pas partie de*. De ce point de vue, en supposant que «R» abrège «être partie de», l'énoncé «cette molécule R moi-même» ainsi que l'énoncé «cette molécule (non R) moi-même» sont faux, mais les énoncés «non (cette molécule R moi-même)» et «non (cette molécule (non R) moi-même)» sont vrais tous les deux. Cette approche oblige à abandonner le «principe d'obversion» selon lequel chaque fois qu'un individu ne satisfait pas un prédicat, il satisfait son contraire (par exemple, si quelque chose n'est pas pair, alors il est non pair); notez ainsi $\neg(Pa) \equiv (\neg P)a$ (cf. Frey 1987; et Miéville 1991). Le problème avec cette approche, c'est qu'il est toujours possible d'employer la «négation faible», et pas seulement la «négation forte», pour obtenir un prédicat à partir d'un autre. Supposons que «-R» soit un prédicat tel que n'importe quel couple tombe dans son extension si et seulement si il tombe en dehors de l'extension de «R» (ce qui n'équivaut pas à dire: si et seulement si il tombe dans l'extension de «-R»). Puisque l'approche en question suggère que le couple <cette molécule, moi-même> tombe en dehors de l'extension de R, alors il doit tomber dans l'extension du prédicat «-R» («-R» doit être vrai de ce couple, puisque «R» est faux de lui). Mais quelles raisons peut-il y avoir pour

des apparences, le temps n'autorise rien à entrer ou à sortir de l'extension de la relation de partie (il s'agit d'un espace fermé). Il se passe tout simplement que plusieurs couples de choses ne sont ni dans cet espace ni au dehors. Autrement dit, la relation absolue de partie n'est pas bivalente lorsqu'elle est interprétée sur un univers de continus¹¹. Bien sûr, il y a des couples qui tombent de

soutenir cela? Pourquoi dire que le couple <cette molécule, moi-même> tombe dans l'extension de «-R» et en dehors de l'extension de «R»? En effet, tout ce qu'on voit est qu'à certains instants il semble tomber dans l'extension de «-R» et à d'autres dans l'extension de R. Mais cela était déjà évident à propos de «R» et «-R» et c'était justement pour cette raison qu'on avait l'impression que le couple <cette molécule, moi-même> n'était ni dans l'extension de «R» ni dans celle de «-R». Donc, il ne semble pas y avoir plus de raisons de dire que le couple tombe dans l'extension de «-R» que de dire qu'il tombe dans l'extension de «R».

2. On peut soutenir (cf. Johnston 1987) qu'il y a plusieurs manières pour une chose (un couple de choses, un triple de choses etc.) de posséder une propriété: elle peut la posséder-à- t^1 , la posséder-à- t^2 , etc. (où «à- t^1 », «à- t^2 »,... sont des modificateurs adverbiaux). Il y a également une quantité de relations différentes de satisfaction entre séquences d'entités et prédicats (la t^1 -satisfaction, la t^2 -satisfaction,...), et une quantité de types différents de vérité (la t^1 -vérité, la t^2 -vérité,...). Donc, si cette molécule est une partie de moi à t^1 mais non pas à t^2 , le couple <cette molécule, moi-même> possède à t^1 la propriété absolue *être partie de*, mais ce n'est pas le cas qu'il la possède-à- t^2 . Il t^1 -satisfait le prédicat «être partie de» mais il ne le t^2 -satisfait pas. Et l'énoncé «cette molécule est partie de moi» est t^1 -vrai mais t^2 -faux. Une discussion de cette approche est de loin trop complexe pour être abordée ici, je me bornerai donc à un simple aperçu. Je ne crois pas qu'on doive être scandalisé par l'idée générale d'une *vérité-à-un-certain temps*, au contraire cette idée semble absolument naturelle dans le cas des énoncés doués de temps verbal. Par exemple, il est tout à fait évident de dire que l'énoncé «Andrea Bottani est en train de travailler à l'ordinateur» est vrai maintenant mais était faux il y a quelques heures. La raison qui permet d'énoncer ceci est que prononcer à ce moment l'énoncé serait vrai *simpliciter*, tandis que prononcer il y a quelques heures ce même énoncé serait faux *simpliciter*. Il est clair qu'un énoncé dénué de temps verbal comme «cette molécule être partie de moi» ne se comporte pas de la même façon que «Andrea Bottani est en train de travailler à l'ordinateur». Il diffère parce que personne ne peut dire quelque chose de vrai ou faux en prononçant (maintenant, demain ou il y a une année) l'énoncé «cette molécule être partie de moi» (soutenir le contraire équivaut à nier toute différence entre le temps verbal présent et l'absence de temps verbaux). Mais si une énonciation à l'instant t de cet énoncé n'est ni vraie ni fausse, ne peut pas y avoir aucune raison pour soutenir que cet énoncé est t -vrai ou t -faux.

- 11 Cet usage du prédicat «bivalent» est loin d'être orthodoxe ou normal. Généralement, la bivalence est traitée comme une propriété de langages: un langage est bivalent si et seulement si tous ses énoncés (toutes ses formules bien formées) sont vrais ou faux. Ici la bivalence est traitée comme une propriété de propriétés: une propriété est bivalente si et seulement si n'importe quelle chose tombe à l'intérieur ou à l'extérieur de son extension

manière déterminée à l'extérieur de la relation absolue de partie (par exemple le couple <la Mer Méditerranée, moi-même>), et d'autres couples qui tombent de manière déterminée à l'intérieur de cette relation (par exemple, le couple <mon cerveau, moi-même>). Peut-être que les sommes méréologiques sont des entités telles que n'importe quel couple d'entre elles tombe de manière déterminée à l'extérieur ou à l'intérieur de la relation absolue de partie (mais peut-être que non, comme l'exemple de Clincoln semble l'avoir bien montré).

En ce qui concerne la non-bivalence, le cas de la relation absolue de partie ne se détache pas de ceux de plusieurs autres propriétés absolues ou éternelles. Prenons la propriété absolue, temporellement non qualifiée, lewisienne, *être plié*. J'ai la propriété d'être maintenant plié. Je n'ai pas la propriété d'avoir été plié il y a une heure. Mais je n'ai pas, ni ne manque d'avoir, la propriété d'être plié *simpliciter*, c'est-à-dire que je ne suis ni à l'intérieur ni à l'extérieur de la propriété *être plié*. Comme pour la relation absolue de partie, la propriété *être plié* est non bivalente: certaines choses tombent dans l'extension de la propriété, d'autres n'y tombent pas et d'autres encore ne font ni l'un ni l'autre.

C'est cela qui nous permet de sauver la loi de Leibniz, comme cela nous permet de sauver l'extensionnalité méréologique. (En fait, la loi de Leibniz semble être une proche parente du principe d'extensionnalité méréologique: elle montre que x et y sont identiques si et seulement s'ils ont les mêmes propriétés, alors que le principe d'extensionnalité méréologique dit qu'ils sont identiques si et seulement s'ils ont les mêmes parties). En effet, proposons l'argument suivant. Dans cette salle, il y a un x qui se trouve à un instant dans l'extension de «plié» et à un autre en dehors de cette extension. Donc, dans cette salle, à un certain instant, il y a une chose qui est pliée et à un autre une chose qui n'est pas pliée et ces choses sont la même chose. En conclusion, ou les occupants de cette salle ne sont pas ceux que nous pensons, ou la loi de Leibniz n'est pas vraie des occupants de cette salle. La non-bivalence de certaines propriétés peut être employée pour bloquer cette attaque à la loi de Leibniz tout comme elle peut être employée pour

(autrement dit, si et seulement s'il n'y a aucun objet tel qu'elle n'est ni vraie ni fausse de lui).

bloquer l'attaque parallèle à l'extensionnalité méréologique. Les deux attaques, en un certain sens, sont identiques. Elles ont la même structure et doivent être traitées de la même manière.

Un point important concerne le rapport entre les propriétés «simples» ou temporellement non qualifiées d'un objet et ses propriétés «permanentes». Je ne pense pas que toute entité pliée d'une façon permanente soit pliée *simpliciter* – bien que je pense que toute entité pliée *simpliciter* soit pliée d'une façon permanente. Supposons qu'une chaise pliable ait toujours été pliée. Cela signifie tout simplement que, pour n'importe quel instant t auquel la chaise a existé, la chaise a eu la propriété *être plié à t* . Donc, la propriété *être plié d'une façon permanente* n'est pas une propriété «simple» ou temporellement non qualifiée, mais quelque chose de semblable à une liste de propriétés temporellement qualifiées. Si *être plié simpliciter* équivaut à *être plié d'une façon permanente*, il n'y aurait rien comme une propriété temporellement non qualifiée. Mais si je dis «il est une personne (*simpliciter*)», la vérité de ce que je dis ne dépend aucunement du temps (au contraire, c'est la vérité de ce que je dis à limiter le pouvoir du temps de déterminer des changements: si quelque chose est une personne *simpliciter*, elle ne peut jamais être quelque chose de différent d'une personne ou cesser d'être une personne, quoiqu'elle puisse, bien entendu, cesser d'exister). Si quelque chose possède à un instant t la propriété *être une personne simpliciter*, on n'a pas besoin d'attendre et de contrôler ce qui lui arrive dans le futur (attendre et contrôler si elle reste une personne à chaque instant de son histoire) pour être sûr qu'elle possède véritablement cette propriété. En général, seules les propriétés absolues d'une entité sont ses propriétés essentielles. Et seules les relations absolues de partie sont des relations essentielles de partie. C'est-à-dire que n'importe quel couple $\langle x, y \rangle$ est dans l'extension de la relation temporellement non qualifiée de partie si et seulement si x est une partie essentielle de y ¹².

12 Tout comme un objet x peut être P d'une façon temporellement permanente (c'est-à-dire P -a- t pour tout instant t auquel il existe) sans être P *simpliciter*, x peut être P d'une façon modalement permanente (c'est-à-dire P -a- t -dans- w pour tout instant t auquel il existe dans tous les mondes possibles dans lesquels il existe) sans être P *simpliciter*. Prenons par exemple moi. Étant plié à quelques instants et non plié à quelques autres, je ne suis pas

IV. Extensionnalité méréologique et ontologie ordinaire: un pseudo conflit

Si on tient compte du fait que la relation de partie traitée par le principe d'extensionnalité méréologique est la relation absolue ou temporellement non qualifiée de partie (la relation dyadique ...*être partie de*...), le conflit entre le principe et le sens commun s'avère purement illusoire. Même si une personne peut avoir des parties matérielles (disons quelques molécules) à un instant t^1 et ne pas les avoir à un instant t^2 ($t^1 \neq t^2$), ce n'est pas une situation où des choses ont des différentes parties *éternelles* ou *absolues* tout en étant la même chose (et même pas une situation où des choses ont des différentes parties *temporaires* tout en étant la même chose, parce que l'adulte que je suis et l'enfant que j'étais ont exactement les mêmes parties temporaires: si certaines molécules étaient des parties de l'enfant à un instant t^1 , ces mêmes molécules étaient des parties de l'adulte au même instant). Il s'agit plutôt d'une situation où trois choses – une molécule, une personne et un instant – sont dans une relation triadique (*être partie de*... à...) et trois autres choses – la même molécule, la même personne et un différent instant – ne sont pas dans la même relation triadique. Quant au couple de la molécule et de la personne, il n'est ni dans l'extension de la relation dyadique *être partie de* ni au dehors

plié d'une façon temporellement permanente et je ne suis pas non plié d'une façon temporellement permanente. Certes, je suis plié ou non plié d'une façon temporellement et modalement permanente, parce que, pour tout instant t de mon existence et tout monde w auquel je suis présent, je suis plié-ou-non-plié-à- t -dans- w . Néanmoins, je ne suis pas plié-ou-non-plié *simpliciter*. En effet, je ne tombe ni à l'intérieur ni à l'extérieur de l'extension de la propriété «être plié» pas plus qu'à l'intérieur ou à l'extérieur de la propriété «être non plié». Donc, je ne tombe ni à l'intérieur ni à l'extérieur de l'extension de la propriété «être-plié-ou-non-plié», et cette propriété ne peut pas être une propriété «simple» (c'est-à-dire temporellement et modalement non qualifiée) de moi. Si les propriétés «temporellement et modalement non qualifiées» d'une entité sont ses propriétés essentielles, alors il y a une différence entre le concept de propriété essentielle et celui de propriété nécessaire. Une propriété P est nécessaire pour une entité x dans le cas où elle est temporellement et modalement permanente pour x (autrement dit, dans le cas où x est P -à- t -a- w pour n'importe quels x et w tels que x existe à t dans w). Une propriété P est *essentielle* pour x dans le cas où elle est temporellement et modalement non qualifiée et x tombe dans son extension. On peut facilement voir que toute propriété essentielle d'une entité est nécessaire pour elle, mais non pas le contraire.

d'elle, soit à t' soit à t'' (en fait, il n'était ni dans cette extension ni au dehors d'elle depuis toujours et il restera là pour toujours). Donc, le principe d'extensionnalité méréologique ne peut pas être faux, même quand il est interprété sur l'univers de nos «continus» ordinaires.

Une fois établi que le principe d'extensionnalité méréologique ne peut pas devenir faux lorsqu'il est appliqué au domaine de nos objets ordinaires, il reste encore à se demander s'il peut devenir *ni vrai ni faux*. Prenons 4', la version formelle du principe, mentionnée dans la note 2: $(\forall x)(\forall y)((x=y) \leftrightarrow (\forall z)(z < x \leftrightarrow z < y))$ – où «<» exprime la relation *absolue* de partie. 4' reste évidemment tout à fait vrai dans tous les cas où chacune de ses parties énonciatives est vraie ou fausse. Mais en fait, puisque plusieurs couples d'objets ne tombent ni à l'intérieur ni à l'extérieur de l'extension de «<», quelques parties énonciatives de 4' ne peuvent être ni vraies ni fausses. Qu'advient-il du principe dans ces cas là? Cela dépend évidemment de la table trivalente de vérité pour « $\rightarrow$ » (et donc pour « $\leftrightarrow$ »), qu'on décide d'adopter. Supposons que l'on adopte celle de Lukasiewicz¹³, selon laquelle une proposition conditionnelle est vraie quand son antécédent est faux, ou son conséquent vrai ou tous les deux ne sont ni vrais ni faux, elle est fausse quand son antécédent est vrai et son conséquent faux, et elle est indéterminée dans tous les autres cas. Cela dit, considérons la moitié droite de 4'. Elle ne peut être ni vraie ni fausse si et seulement si: (i) elle n'a aucune instance fausse; autrement dit, il n'y a aucun z tel qu'il est vrai que z fait partie de x et il est faux que z fait partie de y (ni le contraire); (ii) elle a au moins une instance qui n'est ni vraie ni fausse; autrement dit, il y a au moins un z tel qu'il est vrai que z fait partie de x mais il n'est ni vrai ni faux que z fait partie de y (ou le contraire). Dans un tel cas, toutefois, on pourrait arguer que x et y devraient être différents, et donc la moitié gauche de 4' devrait être fausse¹⁴. Par

13 Voir Lukasiewicz 1920, 1930.

14 Dans pareil cas, en effet, l'entité y aurait la propriété de n'être ni à l'intérieur ni à l'extérieur de la propriété avoir z comme partie, alors que x n'aurait pas cette propriété. Donc, par la loi de Leibniz, x et y devraient être différents. Ce schème argumentatif a été employé pour la première fois par Evans (1978) afin de montrer que, pour n'importe quels objets a et b , il ne peut jamais arriver que a et b ne soient ni identiques ni différents.

conséquent, ayant une moitié droite indéterminée et une moitié gauche fausse, 4' ne serait dans ce cas ni vrai ni faux, du moins selon la table de Lukasiewicz pour « $\leftrightarrow$ ». On obtient le même résultat si on adopte la table de vérité trivalente de Kleene pour « $\rightarrow$ » (et donc pour « $\leftrightarrow$ ») au lieu de celle de Lukasiewicz¹⁵.

J'ai montré que s'il y a des entités x et y telles que (i) et (ii) sont vrais, alors le principe d'extensionnalité méréologique devient ni vrai ni faux. Mais y a-t-il de telles entités? Et *peut-il* y en avoir? Je crois que non, pour une raison assez simple: pour n'importe quels x et y , si x est différent de y , alors il y a au moins une entité z telle qu'il est vrai que z fait partie de x (*simpliciter*) mais il est faux que z fait partie de y (*simpliciter*); il s'agit de x lui-même. Ainsi, il n'y a pas d'entités x et y telles que (i) puisse être vrai d'elles. Donc, ni la moitié droite de 4', ni par conséquent 4' ne peut être indéterminée. En conclusion, le principe d'extensionnalité méréologique reste universellement *vrai*, non pas uniquement pour un univers de sommes méréologiques mais également pour l'univers familier de nos objets ordinaires. On peut arriver à la même conclusion en ce qui concerne ce noyau essentiel de M.E.C. que nous avons appelé «C», et en particulier pour les lois de transitivité et de composition non restreinte.

Toutefois, il se pourrait bien que cela ne soit pas encore suffisant pour le méréologiste extensionnaliste radical. En effet, il pourrait demeurer insatisfait de toute chose qui n'était pas l'assomption du fait qu'il ne peut pas y avoir d'individus x et y tels que le couple ordonné $\langle x, y \rangle$ se trouve dans l'extension du prédicat temporel *être une partie de à t^1* et en dehors de l'extension du prédicat *être une partie de à t^2* . Cette assomption n'a pratiquement rien à voir avec le principe d'extensionnalité méréologique et n'a en tout cas presque rien de sa force intuitive. Essayer de la dériver de ce principe serait un peu comme essayer de dériver de la prémisse que a et b ne peuvent pas être le même chat s'ils n'ont pas les mêmes couleurs, la conclusion qu'ils ne peuvent pas avoir de couleurs différentes dans des parties différentes de leurs

15 Cf. Kleene 1952. En fait, la seule différence entre Lukasiewicz et Kleene à propos la table de vérité de « $\rightarrow$ » c'est que selon Lukasiewicz une conditionnelle est vraie lorsque ses moitiés sont toutes les deux indéterminées, alors que selon Kleene elle est indéterminée comme ses moitiés.

corps¹⁶. De toute façon, cette assumption est tout à fait incompatible avec l'existence d'objets ordinaires tels que des chats, des tables ou des nuages, et doit être abandonnée si on veut avoir une théorie méréologique véritablement universelle.

D'un autre côté – assez étrangement peut-être – deux objets *a* et *b* peuvent être différents même si, pour n'importe quel instant *t*, n'importe quelle entité *x*, *x* est partie de *a* à *t* si et seulement si *x* est partie de *b* à *t*. La raison en est que *a* et *b* pourraient tout de même avoir de différentes relations absolues ou éternelles de partie. C'est le cas de la statue et de la portion de matière recueillie qui la compose lorsque cette portion est générée et détruite en même temps que la statue. Ce type d'approche méréologique peut rendre plus facile le traitement de certains «paradoxes de la composition». Prenons le célèbre paradoxe de «Lumpl et Goliath». En pétrissant de l'argile, un sculpteur modèle séparément deux moitiés du corps de l'infant Goliath, puis il les assemble. En faisant cela, il produit une nouvelle statue et en même temps un nouveau morceau d'argile, mais le lendemain, insatisfait de son œuvre, il détruit la statue et le morceau d'argile. Est-ce que la statue *est* le morceau d'argile? Une réponse affirmative semble tout à fait impossible, mais une réponse négative semble incompatible avec le principe d'extensionnalité méréologique, parce que la statue et le morceau d'argile semblent avoir exactement les mêmes parties. Ce qui se passe, toutefois, c'est seulement que la statue et le morceau d'argile ont les mêmes parties *temporaires*: pour n'importe quelle entité *x*, pour n'importe quel instant *t*, *x* fait partie de la statue à *t* si et seulement si *x* fait partie du morceau d'argile à *t*. Mais la statue et le morceau d'argile n'ont pas les mêmes parties *atemporelles* ou *absolues*. Le morceau d'argile est tel que n'importe quelle molécule *x* est une partie de lui *simpliciter* ou bien ne l'est pas (son identité dépend de celle de chacune des molécules qui le composent, parce qu'un morceau de matière moins une molécule n'est pas le même morceau), tandis que toute molécule particulière n'est ni à l'intérieur ni à l'extérieur de l'extension de la propriété être partie *simpliciter* de la statue (il n'y a aucune molécule *m* telle qu'avoir *m* comme partie est une

16 Tout comme un chat peut avoir des couleurs différentes dans les différentes parties de son corps, il peut avoir des parties différentes à des instants différents de sa vie.

condition nécessaire pour être identique à la statue, parce qu'une statue moins une molécule est encore la même statue).

V. Conclusions

Nos principales conclusions sont les suivantes. C, le noyau essentiel de la méréologie extensionnelle classique (M.E.C.), est universellement valable. L'impression largement répandue qu'il ne l'est pas, est suscitée par une confusion entre la relation absolue ou éternelle de partie et la relation temporellement qualifiée de partie, c'est-à-dire entre la relation dyadique *être partie de* (qui est non bivalente, étant ni vraie ni fausse de plusieurs couples de nos entités ordinaires) et la relation triadique *être partie de... à*. Cette confusion entre propriétés absolues (non bivalentes) et propriétés temporellement qualifiées ne produit pas seulement l'illusion que l'extensionnalité méréologique est incompatible avec notre ontologie ordinaire. En général, elle produit l'illusion que le phénomène du changement est incompatible avec la loi de Leibniz. Notre stratégie a été de démasquer l'illusion particulière comme un cas spécifique de l'illusion générale. En ce qui concerne M.E.C., bien qu'elle traite uniquement de la relation absolue et non bivalente de partie, ses principes essentiels restent tout à fait vrais même pour nos entités ordinaires et, par conséquent, ils s'avèrent ontologiquement neutres. M.E.C. est donc universelle, mais elle n'est pas exhaustive, parce qu'elle ne dit rien des relations triadiques, temporellement qualifiées, de partie. Elle n'explique pas comment le prédicat absolu de partie interagit avec l'appareil des opérateurs temporels et modaux. Étant universelle mais non exhaustive, la M.E.C. n'a pas besoin de réductions mais d'aménagements. Elle doit être renforcée et pas absolument réduite.

Dipartimento di filosofia
Università de Bergamo
Via Salvecchio, 19
I-24129 Bergamo
E-mail: abottani@unibg .it

Bibliographie

- CHISHOLM R. (1973). Parts as Essential to Their Wholes. *Review of Metaphysics* 26, 581-603.
- EVANS (1978). Can There Be Vague Objects? *Analysis* 48/3, 130-4.
- FREY L. (1987). La négation dans la logique d'Aristote (*Pensée naturelle, logique et langage*, à Jean-Blaise Grize). *RESS* XXC/77, 47-60.
- JOHNSTON M. (1987). Is there a Problem about Persistence? *Proceedings of the Aristotelian Society. Supplementary* Vol. 61, 107-35.
- KLEENE S.C. (1952). *Introduction to Metamathematics*. Amsterdam: North Holland.
- LEONARD H.S., GOODMAN N. (1940). The Calculus of Individuals and its Uses. *Journal of Symbolic Logic* 5, 45-55.
- LESNIEWSKI S. (1927-30). O Podstawach Matematyki. *Przegląd Filozoficzny* 30, (1927), 164-206; 31 (1928), 261-91; 32 (1929), 60-101; 33 (1930), 70-105, 142-170.
- LEWIS D. (1986). *On the Plurality of Worlds*. Oxford: Blackwell.
- LEWIS D. (1988). Rearrangement of Particles: Reply to Lowe. *Analysis* 48, 65-72.
- LEWIS D. (1991). *Parts of Classes*. Oxford: Blackwell.
- LUKASIEWICZ J. (1920). On Three-valued Logic. In: S. McCall, 1987, 16-18.
- LUKASIEWICZ J. (1930). Philosophical Remarks on Many valued Systems of propositional Logic. In: S. McCall, 40-65.
- MCCALL S. (ed.) (1987). *Polish Logic: 1920-1939*. Oxford. Clarendon Press.
- MIEVILLE D. (1991). *La négation, une étude logique*. Université de Neuchâtel: Centre de Recherches Sémiologiques (Travaux de logique 6).
- QUINE W. V. O. Worlds Away. *The Journal of Philosophy* LXXIII/22, 859-863.
- SIMONS P. (1987). *Parts. A Study in Ontology*. Oxford: Clarendon Press.
- WIGGINS D. (1980). *Sameness and Substance*. Oxford: Blackwell.

DES TÊTES ET DES HOMMES

Nadine GESSLER

*Donc si j'ai bien compris, dis-je,
les barbicelles sont des composants
des barbules, les barbules des com-
posants des barbes, les barbes des
composants des plumes, et les
plumes des composants des
oiseaux? (Patricia Cornwell)*

1. Difficultés

Force est de reconnaître que la logique extensionnelle moderne, telle qu'elle s'est constituée et a évolué depuis la fin du XIX^e siècle, a joué un mauvais tour à la relation de partie à tout. Pour s'en convaincre, il suffit de considérer les deux arguments suivants.

1) Socrate est un homme.

Donc La tête de Socrate est la tête d'un homme.

2) Tout homme est un animal.

Donc Toute tête d'homme est une tête d'animal.

Le second argument est le célèbre argument que De Morgan cita à titre d'inférence qui n'entre pas dans le cadre de la syllogistique aristotélicienne. (De Morgan 1847: 114; 1966: 29). Souvent mentionné dans les ouvrages de présentation de la logique classique des prédicats, il se traduit comme suit.

Pour a =df être un homme; b =df être un animal; r =df être la tête de.

$$(\forall x)(ax \supset bx)$$

$$(\forall y)((\exists x)(ax \wedge ryx) \supset (\exists x)(bx \wedge ryx))$$

Montrer la validité formelle de cet argument (tout comme du premier) dans le cadre de la logique des prédicats ne soulève aucune difficulté. En revanche, il est impossible d'associer à la procédure de validation syntaxique un modèle susceptible de rendre compte de la relation de partie à tout, *en tant que telle*, qui est en cause dans ces arguments.

En effet, sous une telle interprétation (à supposer qu'elle soit possible), le premier est valide si et seulement si Socrate appartient à la classe de Socrate; le second si et seulement si la classe des têtes des hommes est incluse dans celle des hommes. Or, les individus composant une classe distributive sont des éléments fondamentaux, atomiques, qui ne peuvent pas être abordés à un niveau de complexité ingrédientielle. Aussi Socrate est-il le seul élément de la classe composée de Socrate. De même, l'appartenance d'un objet à la classe des hommes étant univoquement déterminée par le concept «homme», les têtes des hommes n'appartiennent pas à la classe des hommes. De là il est impossible de former, à partir de l'extension composant la classe des hommes, une classe qui serait celle des têtes des hommes et, à ce titre, incluse dans la classe des hommes.

Ces arguments posent donc, fondamentalement, la question d'élargir à la relation de partie à tout le champ de l'extensionnel, délimité dans le cadre de la sémantique extensionnelle classique par les relations usuelles ensemblistes d'appartenance et d'inclusion.

On peut certes valider ces arguments dans le cadre de la sémantique distributive. Seulement, toute interprétation distributive est condamnée à traiter une propriété telle que «être la tête d'un homme» comme une relation entre deux classes, la classe des têtes et la classe des hommes. Mais, en représentant les têtes des hommes, on ne représente en aucune manière les têtes des hommes, au sens de parties ou d'ingrédients constitutifs des tous

que sont les hommes. La classe des têtes des hommes, relativement à l'interprétation méréologique faite de la relation «être la tête de», ne peut pas être représentée.

Mon propos est donc clair. Je ne conteste pas que ces arguments puissent trouver une réalisation ensembliste dans le cadre de la sémantique distributive. Je conteste que cette réalisation soit satisfaisante, compte tenu de la relation de partie à tout qui intervient dans ces démarches déductives. À partir de là, l'objectif à atteindre est celui d'une procédure de validation avec laquelle on puisse réellement interpréter la relation en cause comme une relation méréologique.

Pour ce faire, il est donc nécessaire de renoncer au modèle ensembliste au profit d'un modèle de type méréologique, capable d'offrir une autre définition de la classe et une définition de la relation de partie à tout. J'adopterai ici la méréologie de Lesniewski, appelée aussi théorie des classes collectives. C'est la première théorie formelle de la relation de partie à tout à avoir été construite, durant la deuxième décennie du siècle passé. Elle présente le grand mérite d'être logiquement fondée. Ses fondations logiques sont composées de deux théories: l'ontologie et la protothétique¹. L'axiomatique de la méréologie s'ancre dans celle de l'ontologie, elle-même ancrée dans celle de la protothétique. De la protothétique, qui est un calcul propositionnel quantifié, je ne dirai rien dans le cadre de cet article. Par contre, je serai amenée à dire quelques mots de l'ontologie, pour la simple raison que c'est avec elle que se joue, moyennant les axiomes propres de la méréologie, la possibilité d'articuler ensemble les niveaux distributif et collectif.

Mais dans un premier temps, j'inscrirai les deux arguments dans une problématique plus générale liée à la catégorie du nom. Ce détour nous conduira directement aux deux interprétations distributive et collective des classes qui, jusqu'à Lesniewski, furent mutuellement exclusives l'une de l'autre.

1 Sur la protothétique, cf. Lesniewski 1992: 441-605, Miéville 1984: 43-254).

Sur l'ontologie, cf. Lesniewski 1992: 606-648; 1989: 101-114; Lejewski: 104-119 in Szrednicki 1984; Miéville 1984: 267-373; Joray 1999: 155-172.

2. Distributif *versus* collectif

Tout d'abord, une précision s'impose sur l'emploi que je fais ici du terme «nom». J'entends par «nom» toute expression nominale, simple ou complexe, visant à désigner, référer à un individu particulier. Autrement dit tout terme singulier, mais au sens grammatical. Les divergences de point de vue des logiciens ou philosophes qui ont accompagné les définitions du nom logique n'entrent pas dans mon propos. Seul le critère linguistique est retenu. Par conséquent, parmi les noms, on trouve aussi bien «Orion», «Socrate», «le verrou de la porte de l'entrée de l'aile ouest du château» que «la collection de mes timbres» ou encore «la classe des hommes».

Pourquoi parler du nom singulier? Tout simplement parce que derrière la question du nom, comme désignateur référentiel, se projette celle de l'objet *en tant que tel*, au cœur même de la problématique méréologique. Le constat qui va suivre est banal, mais c'est bien le nœud à défaire, si l'on veut résoudre le problème sémantique posé par les arguments. De par sa fonction référentielle, le nom nous met face à deux modes d'analyse sémantique des objets. En ce qui concerne le premier, c'est l'acte de prédication distributive qui en est le dépositaire. En tant que capturé par un nom, l'objet se présente comme une entité individuelle et peut, à ce titre, être rangé parmi d'autres sous un concept distinctif. Quant au second mode d'analyse, c'est celui de l'objet *en tant que tel*, nommé par un nom. L'objet est un tout collectif, l'unité effective des parties qui le constituent, au sens littéral du terme «constituer».

Ainsi les noms «Socrate» et «Orion» désignent-ils chacun un objet particulier qui se révèle être, sur le plan purement référentiel, une entité collective de parties ou d'ingrédients particuliers. Que l'un soit un continuum spatio-temporel et l'autre pas est sans importance du point de vue nominal adopté.

J'ai placé, parmi les exemples de noms singuliers, deux expressions nominales de classe: «la collection de mes timbres» et «la classe des hommes»². Du point de vue grammatical adopté, ce

2 Le nom complexe «le verrou de la porte de l'entrée de l'aile ouest du château» est un simple clin d'œil à la préposition «de», modalité omniprésente et souveraine de la langue

sont des noms singuliers. Se pose donc, à leur propos, la question du statut et de la nature de l'objet qu'ils dénotent. Deux réponses à cette question sont possibles, selon la lecture distributive ou collective qui est faite des expressions de classe. Ces lectures ou interprétations ne sont que le prolongement des deux modes d'analyse sémantique évoqués précédemment avec Orion et Socrate, exemples choisis parmi d'autres d'objets singuliers de ce monde.

La première lecture est celle qui, dans la lignée de Cantor, a été exclusivement privilégiée par la logique de tradition fregéo-russellienne. Comme nous l'avons vu précédemment, les limites, par rapport à une problématique de type méréologique, en découlent immédiatement. Les individus composant une classe, en tant que collection d'objets distincts, en sont les éléments fondamentaux. Une telle appréhension de la notion de classe (ou d'ensemble) fixe immédiatement la signification de la relation d'appartenance d'un individu à une classe. Si un objet est élément de la classe des *a*, autrement dit s'il appartient à l'extension des *a*, cet objet est nécessairement *a*. Dire par exemple que Socrate est élément de la classe des hommes revient à dire que Socrate est un parmi l'extension du concept «homme», c'est-à-dire que Socrate est un homme.

Deux points sont à souligner, en ce qui concerne la théorisation de cette notion de classe dans les théories classiques. Le premier concerne la classe vide et la classe singulière. Le second, quant à lui, touche à la question de la nature et de l'existence des classes.

Tout d'abord, les classes ne peuvent pas être abordées de manière purement extensionnelle. En effet, si on considère que ce sont les objets qui font la classe, conformément à la conception «naïve» qu'une classe en tant que collection d'objets est composée de ces objets, on ne peut avoir qu'une conception agréga-

française par laquelle s'établit la relation de partie à tout. Ce nom pose la question de l'analyse logique des expressions nominales dénotantes construites sur la base de relations sémantiques de type méréologique. Cette question n'est pas traitée dans le cadre de cet exposé où la priorité est accordée au rôle joué par la relation de partie à tout dans certaines démarches déductives. Je renvoie le lecteur à l'étude de P Joray, *La subordination logique, une étude du nom complexe dans l'Ontologie de S. Lesniewski*, dans laquelle on trouve une analyse du relatif «dont».

tive des classes. Ainsi, en faisant fond sur une notion d'extension définie comme un tout formé de parties, on ne peut disposer légitimement de la classe vide qui n'a pas d'éléments du tout. De plus, cette appréhension de la classe conduit à identifier les classes à un seul élément avec leur élément. Je cite à ce sujet Russell:

Nous ne pouvons pas prendre les classes de manière *purement* extensionnelle comme étant simplement des tas, ou conglomerats. [*heaps or conglomerations*] Si nous voulions essayer de le faire, nous trouverions impossible de comprendre où peut bien se trouver une classe comme la classe vide, laquelle n'a pas d'éléments du tout et ne peut pas être tenue pour un «tas», il nous serait aussi très difficile de comprendre qu'une classe qui n'a qu'un élément ne soit pas identique à celui-ci. Je n'ai pas l'intention d'affirmer ou de nier qu'il y ait des entités telles que les "tas". En tant que logicien mathématique, je ne suis pas appelé à avoir une opinion à ce sujet. Tout ce que je maintiens, c'est que, s'il existe quelque chose comme des tas, alors nous ne pouvons pas les identifier aux classes composées de leurs constituants. (Russell 1920: 183; cité in Lesniewski 1989: 65)

Ce commentaire fait à propos du rejet d'une conception purement extensionnelle des classes, par les théories frégeol-russelliennes, j'en viens au second point. On aborde avec lui ce que fut le problème fondamental des théories des ensembles ou des classes héritières de la définition cantorienne. Je rappelle donc cette définition, sous deux formulations:

Par ensemble, nous entendons n'importe quel rassemblement en un tout M d'objets bien définis et distincts m de notre intuition ou de notre pensée; ces objets étant appelés les «éléments» de M . (Cantor 1932: 282)

Ou encore:

Chaque ensemble de choses distinctes peut être considéré *en lui-même comme une seule chose* dont les choses en question sont les parties constituantes ou les éléments constitutifs. (Cantor 1887: 83; cité in Lesniewski 1989: 53)

Le problème en cause, c'est le vieux problème de l'un et du multiple. Comment parler d'unité *et* de multiplicité au sujet des classes? Mais sur cette question, comme on le sait, le paradoxe de Russell a tranché. Une théorie dans laquelle les classes seraient multiples, en tant qu'ayant plusieurs éléments, et une, en tant qu'entités, serait soumise au paradoxe.

Dans les *Principia Mathematica*, Russell et Whitehead résolvent le problème en opérant une réduction des classes aux fonctions propositionnelles qui les déterminent. Cette réduction, condition nécessaire de la théorie des types, est opérée par la «no class theory». Les symboles de classe sont rangés parmi les symboles incomplets: «leurs usages sont définis, mais eux-mêmes sont supposés ne rien vouloir dire du tout». Ainsi, grâce à la définition contextuelle de la classe, les symboles de classe sont éliminables de la théorie. Par conséquent, la question du statut ontologique de l'objet qu'ils semblaient dénoter – objet manifestement paradoxal – n'a tout simplement pas à être posée. Les classes sont devenues des «fictions logiques», de simples «façons de parler».

Aussi les classes, dans la mesure où elles sont introduites, ne le sont que comme des commodités purement symboliques ou linguistiques, et non comme des objets authentiques tels que le sont leurs membres lorsque ce sont des individus. (Whitehead & Russell 1910: 72-73; 1989: 317)

Ainsi fait-on, dans les *Principia Mathematica* le deuil total des classes comme entités de même type que les individus les composant. Comme l'écrivit Russell, les classes ne font pas partie de l'ameublement dernier du monde.

J'insiste sur le nominalisme instrumentaliste élaboré par Whitehead et Russell pour dompter l'antinomie et concilier les dimensions unitaire et multiple des classes pour deux raisons. La première est qu'il est aux antipodes du nominalisme caractérisant la conception collective des classes, telle qu'elle se trouve systématisée dans la méréologie de Lesniewski. La deuxième est que la naissance de la méréologie s'inscrit dans le contexte bien précis des travaux des fondements des mathématiques, marqué par

l'antinomie de Russell. Lesniewski montra en effet qu'une résolution, *stricto sensu*, de l'antinomie nécessitait l'abandon de la définition cantorienne de l'ensemble ou de la classe au profit d'une définition méréologique.

Toute la réflexion de Lesniewski fut basée sur la présupposition que les expressions de classes sont des expressions dénotantes. C'est ce que révèle, écrit-il:

la manière habituelle d'employer les mots "classe" et "ensemble" dans le langage courant des hommes qui ne se sont occupés d'aucune "théorie des classes" ni d'aucune "théorie des ensembles". (Lesniewski 1992: 207; 1989: 53)

Parler de classe a donc une signification opérationnelle évidente en tant qu'acte de désignation d'un objet de même nature que les objets le composant. L'objet dénoté par une expression de classe doit être compté comme un, au sens de l'agrégat ou du tas des entités qui tombent sous le concept distinctif.

À partir de là, la rupture est totale avec les théories classiques. Tout d'abord, sous une approche collective des classes, la condition indispensable de l'existence d'un objet «classe des a» est qu'il existe au moins un a. Autrement dit, la classe vide est récusée.

Étant d'avis que si un objet est la classe de tels et tels a (des hommes pas exemple, des points, des cercles carrés), alors il se compose de ces a, j'ai toujours rejeté [...] l'existence de monstres théoriques dans le genre de la classe des cercles carrés, comprenant bien que rien ne peut être composé de ce qui n'existe pas. (Lesniewski 1992: 214; 1989: 58)

Ensuite, tout objet est identique à la classe de lui-même et, *a fortiori*, avec la classe de la classe de lui-même. Le paradoxe de Russell se trouve donc résolu. Puisque toute classe est élément d'elle-même, aucune classe ne peut être la classe des classes qui ne sont pas éléments d'elles-mêmes. Autrement dit, la classe collective de Russell n'existe pas.

Lesniewski renoue donc avec la conception purement extensionnelle des classes, celle-là même que combattaient les tenants

du logicisme. Cependant, interprétant les expressions de classe d'un point de vue *réellement* collectif, il fait un pas de plus, tout à fait décisif. Lorsque Russell, dans le passage rapporté ci-dessus, parle de la «classe des a » comme «le tas des a », il ne voit comme éléments composant la classe que les a , conformément à une lecture cantorienne de la classe. Par contre, pour Lesniewski, si une classe comme agrégat ou tas est conçue comme une unité constituée d'éléments ou de parties, c'est au sens littéral du terme «constituer». L'agrégat ou le tas des a n'est pas simplement la juxtaposition des a , il en est la totalité effective, conformément cette fois à l'usage ordinaire de ces termes. Interprétant ainsi les expressions de classes, Lesniewski peut alors soutenir que sa position est fidèle à la définition de Cantor.

Ainsi chacun des sons, dont une pièce de musique est l'ensemble, est, conformément à la position de Cantor, une partie constitutive de cet ensemble et la pièce de musique elle-même est composée des sons dont elle est l'ensemble, *tout comme le tableau est composé de telles ou telles parties bien choisies dont il constitue l'ensemble*. (Lesniewski 1992: 207-208; 1989: 53)

La conséquence majeure de cette lecture méréologique concerne la notion d'élément. Une classe collective de a , en tant que tout composé des a , peut être analysée *ad infinitum* et ne nécessite pas d'éléments atomiques. Par exemple, si A est la classe collective des hommes et B est un élément de A , B peut certes être un homme, mais il peut également être la tête de Socrate ou toute collection arbitraire d'hommes ou de leurs parties. En un mot, si un objet est un élément de la classe collective des a , cet objet n'est pas nécessairement a . Il peut être un ingrédient quelconque de la classe.

Par conséquent, le principe d'extensionnalité caractéristique de la sémantique distributive est faux dans la méréologie. Si un objet est la classe collective des a et également la classe collective des b , la classe des a est identique à la classe des b , mais il n'en résulte pas que les objets a soient les mêmes que les objets b . Par exemple, bien qu'une paire de chaussure ne soit pas une chaussure, la classe de mes paires de chaussure est le même objet que la classe

de mes chaussures. Dans les deux cas, on a en effet affaire au même «tas». Ce qui, soit dit en passant, est assez conforme à une certaine réalité.

Je ferai encore un dernier commentaire concernant la terminologie de «classe singulière» dans le cadre de la méréologie. Tout objet est identique à la classe de lui-même. Néanmoins, cela n'entraîne pas qu'il en est l'unique élément. Socrate, par exemple, est un objet singulier valant comme élément ou partie de lui-même et il est également l'ensemble de ses parties constitutives. Souvenons-nous de l'exemple du tableau donné par Lesniewski. On ne pourra donc parler de classe singulière que dans le cas où l'élément en question est un élément atomique, au sens méréologique, c'est-à-dire non composé de parties.

Pour conclure cette partie, je rapporte le commentaire de Lesniewski sur les propos de Russell au sujet de la classe en tant que tas.

[...] conformément à l'emploi que je fais des termes "ensemble" et "classe" ainsi que compte tenu de la manière de se servir du terme "tas" dans le langage courant (M. Russell ne détermine pas le sens du terme "heap" de façon explicite, mais il emprunte ce terme au langage courant dans l'état où il se trouve, donc tout "cru"), je peux toujours dire du "tas" de n'importe quel a qu'il est un ensemble de a , et du "tas" des a composé de tous les a , qu'il est la classe des a . (Lesniewski 1992: 225; 1989: 65-66)

Cet «échange» sur la notion de tas éclaire magnifiquement le conflit de points de vue. D'un côté, le rejet d'une conception purement extensionnelle des classes que l'on refuse pour n'être pas fondée logiquement. De l'autre, la primauté accordée à cette même conception et dont on déroule les conséquences méréologiques jusqu'au bout, rompant ainsi avec toute la tradition ensembliste.

Cette confrontation faite des interprétations distributive et collective des classes, j'en viens maintenant à une brève présentation de l'ontologie.

3. L'ontologie

L'ontologie règle le niveau purement extensionnel, mais sans classes. Il s'agit d'un calcul extensionnel des noms basé sur un seul axiome. Cet axiome formalise l'usage du «est», noté « ϵ », dans la proposition singulière du type «a est b», soit symboliquement « $a \epsilon b$ ».

L'épsilon, qui connecte deux noms, est un foncteur propositionnel à deux arguments nominaux, de la catégorie S/NN. Contrairement à la logique standard où les constantes et les variables doivent être singulières dans toutes les interprétations, les noms dans l'ontologie peuvent être vides, singuliers ou pluriels. Tout nom, qu'il soit vide, pluriel ou singulier joue le même rôle grammatical. L'ontologie ne nécessite en effet pas de distinction de catégorie entre un sujet et un prédicat.

La proposition « $a \epsilon b$ » se lit «a est un parmi les b». Autrement dit «a est un parmi l'extension du nom b». Mais il n'y a ici aucune interprétation de l'expression «l'extension du nom b» en termes de classe distributive. Les significations nominales ne sont pas des classes. En d'autres termes, le nom pluriel b possède une extension, mais il n'est pas le nom de cette extension.

L'axiome est le suivant:

Axiome:

$$\begin{aligned} \lfloor ab \rfloor \lceil a \epsilon b \rceil \equiv . \lfloor \exists c \rfloor \lceil c \epsilon a \rceil \wedge \\ \lfloor d \epsilon c \rfloor \lceil (c \epsilon a \wedge d \epsilon a) \supset d \epsilon c \rceil \wedge \\ \lfloor c \rfloor \lceil c \epsilon a \supset c \epsilon b \rceil \end{aligned}$$

Ainsi, trois conditions sont associées à la vérité d'une proposition singulière « $a \epsilon b$ ».

- | | |
|---|--------------------------|
| 1) a n'est pas un nom vide | Condition d'existence |
| 2) a n'est pas un nom pluriel | Condition d'unicité |
| 3) Ce que désigne le nom a est également désigné par le nom b | Condition de convergence |

Si l'on applique cet axiome, par exemple, à «Orion est une constellation», on obtient:

Orion est une constellation si et seulement si:

- 1) il existe un *c* qui est Orion;
- 2) pour tout objet *c* et *d*, si *c* est Orion et *d* est Orion alors *c* est *d*;
- 3) pour tout *c*, si *c* est Orion alors *c* est une constellation.

Relevons que ces conditions de vérité étaient présentes dans la résolution de l'antinomie de Russell. Puisque l'expression «la classe des classes qui ne sont pas éléments d'elles-mêmes» ne dénote aucun objet, c'est-à-dire est un nom vide, les deux propositions conduisant à l'antinomie, à savoir «la classe des classes qui ne sont pas éléments d'elles-mêmes est élément d'elle-même» et «la classe des classes qui ne sont pas éléments d'elles-mêmes n'est pas élément d'elle-même» sont fausses. Par conséquent l'antinomie ne pouvant pas être construite, il n'y a tout simplement plus d'antinomie (Lesniewski 1992: 202-205; 1989: 49-51).

Malheureusement, il n'est pas possible dans les limites de cet article d'insister davantage sur l'ontologie ni sur le caractère propre des systèmes déductifs de Lesniewski. Je me contenterai d'insister sur le fait que l'ontologie est un calcul extensionnel des noms structurellement adéquat pour parler des objets concrets que sont les classes.

4. Une axiomatique de la méréologie

Je présenterai la première axiomatique construite par Lesniewski. Elle date de 1916 (Lesniewski 1992: 230-232; 1989: 79-80)³.

Le terme primitif est «partie de», symboliquement «pt». Deux axiomes établissent les propriétés d'antisymétrie et de transitivité de cette relation de partie à tout. Ensuite, deux définitions introduisent les foncteurs «élément de», soit «el» et «classe de» soit «kl». Deux autres axiomes énoncent l'unicité de la classe et que s'il existe au moins un *a*, alors la classe des *a* existe.

3 Cette axiomatique fut revue par Lesniewski en 1918, 1920 et 1921. D'autres axiomatiques furent par la suite proposées par Sobocinski 1949-50; Lejewski 1954, et Clay.

Axiome 1:

$$\lfloor AB \rfloor \lceil A \in \text{pt}(B) \supset A \in \sim(\text{pt}(B)) \rceil$$

Quels que soient A et B, si A est partie de B alors B n'est pas partie de A.

Axiome 2:

$$\lfloor ABC \rfloor \lceil A \in \text{pt}(B) \wedge B \in \text{pt}(C) \supset A \in \text{pt}(C) \rceil$$

Quels que soient A, B et C, si A est partie de B et B est partie de C, alors A est partie de C.

Définition 1:

$$\lfloor AB \rfloor \lceil A \in \text{el}(B) \equiv A \in \text{pt}(B) \vee A = B \rceil$$

Quels que soient A et B, A est un élément de B si et seulement si A est le même objet que B ou A est partie de B.

Définition 2:

$$\begin{aligned} \lfloor Aa \rfloor \lceil A \in \text{Kl}(a) \equiv A \in A \wedge \lfloor \exists B \rfloor \lceil B \in a \rceil \wedge \\ \lfloor B \rfloor \lceil B \in a \supset B \in \text{el}(A) \rceil \wedge \\ \lfloor B \rfloor \lceil B \in \text{el}(A) \supset \lfloor \exists CD \rfloor \lceil C \in a \wedge D \in \text{el}(C) \wedge D \in \text{el}(B) \rceil \rceil \end{aligned}$$

Quels que soient A et a, A est la classe des a si et seulement si les trois conditions suivantes sont remplies:

- 1) A est un objet (c'est-à-dire A existe).
- 2) Chaque a est élément de A.
- 3) Si B est élément de A, alors il existe un élément de B qui est élément d'un a.

Axiome 3:

$$\lfloor ABa \rfloor \lceil A \in \text{Kl}(a) \wedge B \in \text{Kl}(a) \supset A = B \rceil$$

Quels que soient A, B et a, si A est la classe des a et B est la classe des a alors A est identique à B.

Axiome 4:

$$\lfloor Aa \rfloor \lceil A \in a \supset \lfloor \exists B \rfloor \lceil B \in \text{Kl}(a) \rceil \rceil$$

Quels que soient A et a, si A est un des a alors il existe B qui est la classe des a.

Quelques rapides commentaires sur cette axiomatique. Il en existe d'autres, mais celle-ci est la plus explicite, même si on peut lui reprocher l'utilisation des définitions dans les axiomes. Cette faiblesse est cependant sans conséquences puisqu'il a été démontré que les définitions ne sont pas créatives.

En premier lieu soulignons que les expressions de classe tombent sous la catégorie des noms. Les foncteurs «partie de», «élément de» et «classe de» sont des foncteurs nominaux de catégorie N/N.

On a donc l'analyse catégorielle suivante pour les énoncés ci-dessous:

A est la classe des a
N S/NN N/N N.

B est élément de la classe des a
N S/NN N/N N/N N.

Quant à la définition de la classe collective, je soulignerai les points suivants. «Kl» fonctionne comme un foncteur unificateur de rassemblement en un tout d'une multiplicité d'objets singuliers. C'est ce que révèle l'analyse catégorielle effectuée ci-dessus. S'il existe des objets quelconques, par exemple des baleines, «la classe des baleines» est un nom singulier qui désigne un objet concret, le tout se composant de toutes les baleines.

Le nom a n'est qu'un nom générique possible parmi d'autres. Souvenons-nous à ce propos des chaussures.

La clause 3 de la définition est celle qui permet de procéder à une analyse ingrédientielle des objets. Considérons par exemple un bol d'olives. Les olives contenues dans le bol peuvent être appréhendées comme une classe collective d'olives, soit A. Le noyau d'une olive quelconque est un élément de A. En effet, il existe un élément du noyau, par exemple son amande, qui est élément d'une olive, en l'occurrence l'olive à laquelle appartient le noyau. Cet élément pourrait être également le noyau lui-même, puisque la relation d'appartenance «être élément de» est réflexive. De même, la classe collective des noyaux des olives est un élément de A. Il existe bien un élément de cette classe, par exemple une amande quelconque, qui est élément d'une olive, soit

l'olive à laquelle appartient l'amande en question. Cette illustration concrète suffit, je crois, à montrer de quelle manière la clause 3 donne à la classe collective sa dimension pluri-extensionnelle.

Pour terminer, je ferai en bref commentaire sur l'axiome 4. En substance, cet axiome dit que l'on peut engendrer, à partir de tout nom pluriel, autrement dit à partir de toute extension, l'entité collective constituée de cette extension. C'est dans la possibilité de ce passage du niveau purement extensionnel au niveau collectif que se résout le problème de l'un et du multiple pour les classes dans les systèmes de Lesniewski. En effet, l'objet dénoté par le terme «classe des a», c'est-à-dire engendré à partir de l'extension des a, peut être appréhendé comme un tout, l'agrégat des entités a qui le compose. L'entité collective et chacun de ses ingrédients est un objet de même nature dont les expressions qui les nomment appartiennent à la catégorie des noms singuliers.

J'en viens à la présentation des thèses de la méréologie validant les deux arguments.

5. Une solution formelle pour les arguments

La résolution du problème sémantique posé par les arguments se joue très précisément sur la distinction et l'articulation entre les niveaux distributif et collectif. Considérons tout d'abord le premier.

Socrate est un homme.

Donc la tête de Socrate est la tête d'un homme.

La thèse suivante est une thèse de la méréologie:

$$1) \lfloor Aa \rfloor \lceil A \varepsilon a \rceil \supset \lfloor B \rfloor \lceil B \varepsilon el(A) \rceil \supset B \varepsilon el(Kl(a)) \rfloor$$

Elle peut se traduire ainsi:

Quels que soient A et a, si A est un des a, alors tout élément de A est élément de la classe collective générée par les a.

Cette thèse valide donc le premier argument. Elle illustre clairement que tout objet singulier restreint, dans les théories classiques

au statut d'objet atomique, peut ici être appréhendé comme la totalité des ingrédients ou parties qui le constituent.

Soit maintenant le deuxième argument:

Tout homme est animal

Donc toute tête d'un homme est une tête d'un animal.

La prémisse est une proposition universellement quantifiée. On considérera qu'elle exprime une quantification forte (inclusion forte) impliquant l'existence d'au moins un objet. Conformément aux conventions en usage et au langage informel qu'est celui de Lesniewski, la quantification forte se traduit par le terme «chaque». La quantification faible (inclusion faible), elle, n'impose aucun engagement existentiel et se traduit par «tous les». La prémisse «Tout homme est un animal» est donc de la forme «chaque a est b ».

En premier lieu, il s'agit de la traduire dans la syntaxe de l'ontologie. Pour ce faire, il est nécessaire d'inscrire une définition pour le relateur de l'inclusion forte. Cette définition est la suivante:

$$\lfloor ab \rfloor \lceil a \subseteq b \equiv . \lfloor \exists C \rfloor \lceil C \varepsilon a \rfloor \wedge \lfloor D \rfloor \lceil D \varepsilon a \supset D \varepsilon b \rfloor \rfloor^4$$

« $a \subseteq b$ » traduit donc la forme verbale «chaque a est b ».

Notons que l'inclusion, et tout relateur extensionnel, est définie pour des noms. On lira donc cette définition de la manière suivante: «Le nom a est fortement inclus dans le nom b c'est-à-dire l'extension du nom a est fortement incluse dans l'extension du nom b , si et seulement si il existe un objet qui est a et ce qui est dénoté par a est dénoté par b ».

Les thèses suivantes sont des thèses de la méréologie.

- 2) $\lfloor ab \rfloor \lceil a \subseteq b \supset \lfloor C \rfloor \lceil C \varepsilon a \supset \lfloor D \rfloor \lceil D \varepsilon \text{el}(C) \supset D \varepsilon \text{el}(\text{Kl}(b)) \rfloor \rfloor$
- 3) $\lfloor Aa \rfloor \lceil a \subseteq b \supset \lfloor C \rfloor \lceil C \varepsilon \text{el}(\text{Kl}(a)) \supset C \varepsilon \text{el}(\text{Kl}(b)) \rfloor \rfloor$
- 4) $\lfloor Aa \rfloor \lceil a \subseteq b \supset . \text{el}(\text{Kl}(a)) \subseteq \text{el}(\text{Kl}(b)) \rfloor \rfloor$

4 Les lettres majuscules sont utilisées pour désigner des noms singuliers. Cependant, les systèmes de Lesniewski ne nécessitent pas de marquer par la notation les noms singuliers.

Soient:

- 2) Si l'extension des objets a est fortement incluse dans l'extension des objets b, alors pour tout objet C qui est a, si D est élément de C, alors D est élément de la classe collective générée par les b.
- 3) Si l'extension des objets a est fortement incluse dans l'extension des objets b, alors tout élément de la classe collective générée par les a est élément de la classe collective générée par les b.
- 4) Si l'extension des objets a est fortement incluse dans l'extension des objets b, alors l'extension des éléments de la classe collective générée par les a est fortement incluse dans l'extension des éléments de la classe collective générée par les b.

Chacune de ces thèses valident l'argument. Néanmoins, on accordera une attention particulière à la thèse 4, eu égard aux commentaires tenus en début de cet article au sujet de cet argument. J'ai écrit que pour valider cet argument dans le cadre de la sémantique ensembliste sous l'angle d'analyse de la relation de partie à tout, ce qui bien entendu se révèle impossible, il faudrait que la classe des têtes des hommes fût incluse dans la classe des hommes. C'est ce rapport d'inclusion que traduit la thèse 4, entre les noms pluriels $el(Kl(a))$ et $el(Kl(b))$, traités comme des noms pluriels. Avec cette thèse, on lit parfaitement le passage du niveau extensionnel au niveau pluri-extensionnel qu'autorise la méréologie, ancrée syntaxiquement dans l'ontologie.

Pour clore cet article, j'évoquerai une remarque formulée par P. Simons lors du colloque, à la fin de mon exposé. M. Simons fit remarquer que l'on pouvait valider les arguments dans l'ontologie, signifiant ainsi que l'on pouvait se passer de la solution méréologique. Ma réponse à cette objection est claire. Que l'on puisse valider les arguments dans l'ontologie est une chose. Mais le cadre d'analyse demeure un cadre d'analyse distributif. Or l'objectif poursuivi est celui d'une solution qui traverse l'épreuve d'une sémantique descriptive et valide ces arguments en traitant les têtes, non pas comme des éléments atomiques tombant sous le nom générique tête, mais comme des parties des tous que sont les hommes. Par conséquent, l'objection de M. Simons n'affecte en

rien la solution que je propose puisque, du point de vue adopté, elle est la seule satisfaisante.

6. Conclusion

Il s'agissait de réparer le mauvais tour joué à la relation de partie à tout par la logique moderne. Adeptes ou non de la sémantique extensionnelle représentée par la théorie des ensembles, je crois que l'on peut admettre que l'objectif est atteint. Pour le reste c'est une affaire de point de vue ou de motivations. Mais il n'en demeure pas moins que lorsque l'on regarde passer un vol d'oiseaux, c'est bien le tout composé des oiseaux que l'on voit passer, avec ailes, plumes et barbicelles.

Institut de logique
Université de Neuchâtel
Espace Louis-Agassiz 1
CH 2000 NEUCHÂTEL
 e-mail: nadine.gessler@unine.ch

Références bibliographiques

- DE MORGAN A. (1847). *Formal Logic or the Calculus of Inference, Necessary and Probable*. London: Tailor & Walton.
- DE MORGAN A. (1966). *On the Syllogism and other Logical Writings*. London: Routledge & Kegan.
- LUSCHEI E.C. (1962). *The Logical Systems of Lesniewski*. Amsterdam: North Holland.
- LESNIEWSKI S. (1989). *Sur les fondements de la mathématique. Fragments*. Paris: Hermès, trad. de G. Kalinowski.
- LESNIEWSKI S. (1992). *Collected Works*. S.J. Surma, J.T. Srezednicki, D.I. Barnett (eds). Dordrecht: Kluwer, 2 vols.

- MIÉVILLE D. (1984). *Un développement des systèmes logiques de S. Lesniewski. Protothétique, Ontologie, Méréologie*. Berne: Lang.
- SIMONS P.M. (1987). *Parts. A Study in Ontology*. Oxford: Clarendon.
- SRZEDNICKI T.J., RICKEY V.F. & CZELAKOWSKI J. (1984). *Lesniewski's Systems. Ontology and Mereology*. The Hague: M. Nijhoff.
- SOBOCINSKI B. (1949-50). L'analyse de l'antinomie russellienne par Lesniewski. *Methodos* 1.2, 220-228.
- WHITEHEAD A. N. & RUSSELL B. (1910-1013). *Principia Mathematica*. Cambridge: CUP. 3 vols.

DÉPENDANCE EXISTENTIELLE, FONDATION ET OBJETS COMPOSÉS

Fabrice CORREIA

Il existe principalement deux théories de la *dépendance existentielle individuelle* – si l'on ne compte pas la troisième théorie obtenue en «fusionnant» les deux premières. Dans cet exposé, je vais montrer que ces deux théories, ainsi que leur fusion, font face à de graves difficultés. Parmi ces difficultés, la plupart de celles que je vais présenter ici ont trait à des considérations méréologiques. Plus précisément, la plupart des objections que je vais présenter contre les théories en question mentionnent des objets composés et leurs composants.

En réaction à ces objections, je vais proposer une autre théorie de la dépendance existentielle qui, comme je vais essayer de le montrer, échappe aux difficultés rencontrées par les autres. Selon cette théorie, la notion de dépendance existentielle est définie en termes d'une relation plus fondamentale, la relation de *fondation*.

Nous verrons ensuite qu'en retour, une classification intéressante des entités composées peut être établie en termes de la notion de fondation et de la notion de dépendance existentielle associée.

1. Les deux théories de la dépendance, leur fusion, et leurs problèmes

Un objet est dit existentiellement dépendant d'un autre objet lorsque l'existence du premier nécessite celle du second, ou encore lorsque le premier a besoin du second pour exister.

Il y a deux thèses que je vais admettre concernant la relation de dépendance existentielle. La première thèse est, je suppose, non controversée. Elle affirme que:

(T1) Si x dépend existentiellement de y , alors il est métaphysiquement impossible que x existe et pas y .

En outre, la notion de dépendance existentielle est supposée renfermer celle de *priorité*: lorsqu'un objet dépend d'un autre, son existence est en un sens «moins basique» que celle de ce dernier. Certains ne seraient pas d'accord avec cette affirmation, car une relation de priorité est asymétrique, et ils admettent comme possible les cas de dépendance mutuelle entre deux objets. Pour eux, l'idée de priorité est renfermée dans celle de dépendance *unilatérale*:

(T2) Si x dépend existentiellement unilatéralement de y , l'existence de x est «moins basique» que celle de y .

Par souci de neutralité, c'est cette dernière thèse que je retiendrai. Il s'agit d'un choix neutre en effet, car pour un ennemi de la dépendance mutuelle, cette deuxième thèse est équivalente à la première (pour lui la dépendance est toujours unilatérale).

Il existe plusieurs manières de comprendre l'idiome «l'existence de x nécessite celle de y ». Sans doute celle qui vient le plus facilement à l'esprit consiste à le comprendre comme «si x existe, alors nécessairement y existe» – ou de manière plus stylisée, comme «l'existence de x implique nécessairement celle de y ». Il s'agit là de la première théorie de la dépendance que je voulais mentionner. Si $\Rightarrow$ représente l'opérateur «si..., alors nécessairement...», la théorie se formule donc ainsi:

(dep1) x dépend existentiellement de y :

x existe $\Rightarrow y$ existe.

Évidemment, cette théorie reste vague tant que l'on n'a pas spécifié plus précisément le sens de l'opérateur $\Rightarrow$. Par «si x existe, alors nécessairement y existe», veut-on dire «nécessairement, si x existe alors y existe», c'est-à-dire est-on en train d'utiliser l'opérateur monadique «nécessairement...» et l'opérateur binaire «si... alors...»? Ou bien est-on en train d'utiliser un opérateur d'implication nécessaire non analysable «si... alors-nécessairement...»? D'autre part, de quel type de nécessité (logique, métaphysique, physique, biologique,...) s'agit-il ici? Ces questions peuvent recevoir plusieurs réponses, et

par conséquent (**dep1**) ne décrit pas une théorie de la dépendance, mais plutôt une classe de telles théories.

Une réponse à la première de ces deux questions ne sera pas requise pour la suite de la discussion. En ce qui concerne la seconde question, on peut admettre que la notion de nécessité en jeu est celle de nécessité métaphysique. En particulier, on peut admettre que si x existe $\Rightarrow y$ existe, alors il est métaphysiquement impossible que x existe et pas y . Cette décision a l'avantage d'avoir pour conséquence immédiate que le principe (**T1**) est automatiquement satisfait si la dépendance est définie selon une théorie du type (**dep1**).

Une des théories du type (**dep1**) que nous allons mentionner prend le principe (**T1**) très au sérieux, tellement qu'elle accepte également sa converse et fait même un pas supplémentaire en admettant le conséquent de (**T1**) comme *definiens* de la dépendance:

(**dep1'**) x dépend existentiellement de y : Il est métaphysiquement impossible que x existe et pas y .

Il s'agit là en effet bien d'une théorie du type (**dep1**): l'opérateur $\Rightarrow$ est ici compris comme l'opérateur d'implication stricte, où la modalité est comprise comme modalité métaphysique¹.

Les théories du type (**dep1**) sont sujettes à des difficultés importantes. Certaines de ces difficultés atteignent seulement une partie de ces théories. Cependant, il existe des problèmes qui les concernent toutes. Pour exploiter un exemple que Kit Fine donne pour critiquer une toute autre position, considérons Socrate et le singleton {Socrate}². Il est plausible de dire que l'existence de Socrate implique nécessairement celle du singleton: pour ainsi dire, l'existence du singleton est automatiquement assurée par celle de son élément. D'autre part, il est également plausible de dire que l'existence du singleton implique nécessairement celle de Socrate: l'existence du singleton requiert celle de son membre. Mais si nous acceptons ces deux thèses en même temps, nous

1 Cette théorie de la dépendance existentielle se trouve dans de nombreux textes. Cf. par exemple Simons 1987: 294 ff.

2 Fine utilise cet exemple pour critiquer l'idée selon laquelle la notion de dépendance d'identité (cf. *infra*) doit être analysée selon le *definiens* dans (**dep1'**). Cf. Fine 1995.

devons dire d'après (**dep1**) que Socrate et le singleton {Socrate} sont existentiellement dépendant l'un de l'autre. Or, cette conclusion pose des problèmes. En effet, l'existence de Socrate est intuitivement plus basique que celle du singleton; et cette priorité existentielle ne correspond pas à une relation de dépendance existentielle unilatérale du singleton à son membre si la notion de dépendance est celle que définit (**dep1**). Bien entendu, au lieu du couple Socrate – {Socrate}, nous serions arrivés à la même conclusion en utilisant n'importe quel couple $a-b$ où b est un ensemble construit à partir de a seulement (comme par exemple $\{a, \{a\}\}$). Ainsi, plus généralement, les théories du type (**dep1**) ne permettent pas de rendre compte de l'asymétrie existentielle qu'il y a entre un objet et un ensemble construit à partir de cet objet seulement.

On pourrait penser qu'il n'y a là aucun problème, que l'asymétrie en question n'a rien à voir avec la dépendance. Mais intuitivement, ça n'est pas le cas. Car en effet, il semble que le singleton soit unilatéralement dépendant du philosophe: intuitivement, le singleton a besoin du philosophe pour exister, mais l'inverse est faux.

Un autre type d'exemple méréologique pose le même genre de problèmes à certaines des théories du type (**dep1**), plus précisément à (**dep1'**). Considérons en effet un individu comme Socrate et le nombre 2, et supposons que ce dernier existe nécessairement. Supposons de plus que la somme méréologique de deux objets quelconques peut toujours être formée, et que dans tous les mondes possibles, une somme existe ssi chacun des objets sommés existe. Alors dans tous les mondes possibles où Socrate existe, la somme méréologique Socrate + 2 existe; et de même, dans tous les mondes possibles où Socrate + 2 existe, Socrate existe. Aussi d'après (**dep1'**), Socrate et Socrate + 2 sont dans une relation de dépendance existentielle mutuelle. Mais là encore, intuitivement Socrate est plus basique que Socrate + 2, et (**dep1'**) ne nous permet pas de rendre compte de cette priorité. Et encore une fois, il n'y a pas là une asymétrie qui n'a rien à voir avec la dépendance: intuitivement, la somme a besoin du philosophe pour exister, mais l'inverse est faux.

La deuxième théorie que je voudrais présenter prétend échapper aux précédentes difficultés. D'après cette théorie, la dépendance existentielle doit être comprise comme *dépendance d'identité*:

(dep2) x dépend existentiellement de y : l'identité de x dépend de y^3 .

Un objet est dit dépendre pour son identité d'un autre objet lorsque le premier est essentiellement relié au second, i.e. lorsqu'une propriété essentielle du premier est d'entretenir une certaine relation au second.

Cette nouvelle théorie de la dépendance existentielle évite en effet de manière remarquable les difficultés précédentes. Considérons à nouveau le cas de Socrate et de son singleton. Intuitivement, le singleton est essentiellement relié à son membre: c'est une propriété essentielle de $\{\text{Socrate}\}$ que de contenir Socrate. Par contre, il est non moins intuitivement clair que le philosophe ne dépend pas pour son identité du singleton. Ainsi selon **(dep2)**, il est intuitivement le cas que $\{\text{Socrate}\}$ dépend existentiellement *unilatéralement* de Socrate. Bien entendu, le même argument peut être reproduit à propos de n'importe quel couple $a - b$ où b est un ensemble construit à partir de a seulement. Le cas de Socrate et du nombre 2 est similaire. Intuitivement, l'identité de la somme Socrate + 2 dépend de Socrate, mais l'identité de Socrate ne dépend pas de la somme. Ainsi, **(dep2)** permet de rendre compte en termes de dépendance existentielle les asymétries discutées jusqu'à présent.

Cependant, **(dep2)** rencontre elle aussi des difficultés. Considérons en effet deux événements ou états de choses a et b , tels que a peut exister (se produire, obtenir,...) sans b . Considérons maintenant l'événement ou état de choses complexe a -ou- b . Alors intuitivement, a -ou- b dépend pour son identité de b : il est essentiellement en partie composé par b . Mais par hypothèse, a -ou- b ne dépend pas existentiellement de b , i.e. l'existence du

3 La théorie est proposée par E. J. Lowe (1998: 147-151), à ceci près qu'il propose une conception légèrement différente de la notion de dépendance d'identité de celle que je propose ici.

premier ne nécessite pas celle du second. Car nous avons admis que a peut exister sans b , et dans tous les cas où a existe, a -ou- b existe également. Le problème avec **(dep2)**, c'est donc que la relation de dépendance existentielle comprise comme dépendance d'identité ne satisfait pas le principe **(T1)**.

Le cas des événements ou états de chose disjonctifs peut être généralisé. Il est assez commun en méréologie d'admettre l'existence de fusions d'objets qui n'existent pas en même temps, par exemple la fusion de Socrate et de mon ami Sam. Et il est également assez commun d'admettre que l'objet résultant – appelons-le Socram – existe exactement aux instants où Socrate *ou* Sam existe. Maintenant, remplaçons le temps par les mondes possibles, et assumons que Sam existe seulement dans un monde possible non actuel. Alors Socram existe exactement dans les mondes possibles où Socrate ou Sam existe. Or, il est plausible de dire que l'identité de Socram dépend de Sam; mais par hypothèse, le premier peut exister sans le second.

Une manière simple de faire face à toutes les difficultés précédentes consiste à «fusionner» les deux théories précédentes, pour obtenir:

(dep3) x dépend existentiellement de y : (x existe $\Rightarrow y$ existe) & (l'identité de x dépend de y).

Malheureusement, **(dep3)** rencontre encore une difficulté, qui d'ailleurs atteint également **(dep1)** et **(dep2)**⁴. Considérons un objet c qui a la propriété essentielle de causer l'effet e (on peut penser à c comme à un dieu leibnizien et à e comme au monde actuel). Alors l'identité de c dépend de e , et il est plausible de dire que l'existence de c implique nécessairement celle de e . Alors d'après **(dep3)**, c dépend existentiellement de e . Or, intuitivement c , la cause, n'a pas besoin de son effet pour exister.

2. Les relations de fondation

Reconsidérons à nouveau Socrate et le singleton {Socrate}. L'existence de Socrate, avons-nous affirmé, est plus fondamen-

4 Je laisse de côté l'objection selon laquelle la proposition est *ad hoc*.

tale ou basique que celle du singleton. Nous avons également vu que cette priorité ne pouvait pas être caractérisée en termes d'implication nécessaire: intuitivement, l'existence de Socrate implique nécessairement celle du singleton, et de la même manière l'existence du singleton implique nécessairement celle de Socrate. Finalement, nous avons vu que cette priorité pouvait être formulée en termes essentialistes: le singleton est essentiellement relié au philosophe, alors que l'inverse est faux. Mais cette asymétrie n'est au fond pas celle qui nous intéressait au début: il y a intuitivement une asymétrie au niveau des existences, pas seulement au niveau des essences. En quoi consiste cette asymétrie?

Une réponse de sens commun est la suivante: le singleton *doit son existence* à son membre, alors que le philosophe ne doit pas son existence au singleton.

La notion de *devoir son existence à quelque chose* est, me semble-t-il, une notion qui mérite une place centrale en métaphysique: elle permet de formuler des thèses importantes concernant le statut existentiel des choses. Par exemple, il est plausible de dire que les ensembles non vides existants doivent leur existence à leurs membres; les changements aux objets changeants; les frontières de corps matériels à ces mêmes corps; peut-être les événements et états mentaux à certains événements et états physiques du cerveau; une perception aux objets perçus; etc. Cette relation de devoir son existence à quelque chose, je l'appellerai la relation d'*être fondé sur*. Être fondé sur quelque chose, c'est en quelque sorte être un parasite de cette chose. Au lieu de fonder, on peut de manière alternative utiliser l'expression *faire exister* ou *rendre existant* («making exist»).

On peut en fait distinguer deux relations de fondation, irréductibles l'une à l'autre. Un objet peut être *partiellement fondé* sur certains objets, ou il peut être *totalement ou complètement fondé* sur ces objets. Ainsi, les ensembles non vides sont plausiblement totalement fondés sur leurs membres: ils doivent complètement leur existence à ces derniers. Par contre, un changement ne doit qu'en partie son existence aux objets changeants. Nous pouvons exprimer la relation de fondation partielle à l'aide d'un prédicat binaire «... est partiellement fondé sur...»; car en effet, être partiellement fondé sur plusieurs objets, c'est être partiellement fondé

sur chacun d'entre eux. Par contre, nous ne pouvons pas faire de même avec la fondation totale; en effet, être totalement fondé sur certains objets n'implique pas être totalement fondé sur chacun d'entre eux. Les relations de fondation partielle et totale sont sous cet aspect comme les relations de causalité partielle et de causalité totale.

Lorsque je dis que les deux relations de fondation sont irréductibles l'une à l'autre, je pense en particulier au fait que l'on ne peut pas définir la fondation partielle en termes de la fondation totale, en posant que x est partiellement fondé sur y ssi il existe des objets $z, t, \dots$ tels que x est totalement fondé sur $y, z, t, \dots$. Le cas des changements est clair sur ce point: il est faux de dire qu'un changement doit complètement son existence aux objets qui changent, et, du moins en général, on ne peut pas trouver de liste plus compréhensive d'objets qui fonderaient totalement le changement. Si les deux relations de fondation sont irréductibles, elles n'en sont cependant pas moins fortement reliées: si x est totalement fondé sur certains objets, alors x est partiellement fondé sur chacun d'entre eux.

Par la suite, par « x est fondé sur y » (sans qualificatif), j'entendrai: « x est partiellement ou totalement fondé sur y ». C'est en termes de cette relation que je définirai la relation de dépendance existentielle.

La notion de fondation nous permet de rendre compte de manière très simple des asymétries rencontrées plus haut. En effet, comme nous l'avons déjà vu, les singletons sont intuitivement fondés sur leur membre, alors que la réciproque est fausse. De même, la somme méréologique Socrate + 2 est fondée unilatéralement sur les deux éléments fusionnés.

3. La dépendance existentielle comme exigence d'un fondement

La relation d'être fondé sur ne peut pas être identifiée à la relation de dépendance existentielle: être fondé sur n'implique pas être existentiellement dépendant de. En effet, contrairement à la relation de dépendance, la relation de fondation est *active*: néces-

sairement, si un objet x est fondé sur un objet y , alors x et y existent. D'autre part, la relation de dépendance est telle que si x dépend existentiellement de y , alors il est impossible que x existe et pas y : il s'agit du principe (T1). Mais la relation de fondation ne satisfait pas ce principe. En effet, considérons l'état de choses disjonctif a -ou- b , et supposons que a existe actuellement sans b , et qu'il est possible que b existe sans a . Alors a -ou- b est (totalement) fondé sur a ; mais il est possible que a -ou- b existe sans que a n'existe.

La dépendance existentielle n'est donc pas la relation d'être fondé sur. En fait, je considère cette dernière comme plus fondamentale que la première: je propose de définir la dépendance existentielle en terme de la notion de fondation. Être causalement dépendant de x , c'est, en un sens de l'expression, avoir besoin de x comme d'une cause (au moins partielle) pour exister. De manière similaire, nous pouvons définir «être existentiellement dépendant de x » par «avoir besoin de x comme d'un fondement au moins partiel pour exister». Plus précisément, je propose:

(dep4) x dépend existentiellement de y : il est métaphysiquement impossible que x existe sans être fondé sur y .

Cette nouvelle théorie de la dépendance échappe à tous les problèmes rencontrés par les théories précédentes. Je les examine tour à tour.

[Socrate et {Socrate}].

Intuitivement, selon (dep4), {Socrate} est unilatéralement dépendant de Socrate: dans tous les mondes possibles où le singleton {Socrate} existe, il doit son existence à Socrate; mais l'inverse est faux: il n'y a pas de monde possible où Socrate doit son existence au singleton, et en particulier dans le monde actuel, Socrate existe et n'est pas fondé sur {Socrate}. Le même verdict vaut pour n'importe quel couple $a - b$ où b est un ensemble construit à partir de a seulement. (dep4) échappe ainsi à la difficulté rencontrée par (dep1).

[Socrate et Socrate + 2].

Même chose, en remplaçant le singleton par Socrate + 2. *Contra (dep1')*.

[L'événement ou état de choses disjonctif; les fusions disjonctives].

Du fait que la relation de fondation est factive, il suit de (**dep4**) que le principe (**T1**) est satisfait. Ainsi, (**dep4**) échappe au problème rencontré par (**dep2**). En particulier, selon (**dep4**), et contrairement aux prédictions de (**dep2**), l'événement disjonctif mentionné *a-ou-b* plus haut n'est pas dépendant de *b*. Le cas des fusions disjonctives est similaire.

[La cause essentielle et son effet].

Les causes ne sont jamais fondées sur leurs effets. Donc d'après notre théorie, et contrairement à la théorie (**dep3**), la cause qui produit essentiellement son effet ne dépend pas de ce dernier.

Nous avons remarqué que la relation de dépendance telle que caractérisée par (**dep4**) satisfait le principe (**T1**). Intuitivement, elle satisfait également le principe de priorité (**T2**). En effet, si nous supposons qu'un objet *x* ne peut pas exister sans devoir (au moins en partie) son existence à un autre objet *y*, et si nous supposons d'autre part que l'inverse est faux – i.e. que *y* peut exister sans devoir son existence à *x* – alors il est suffisamment clair que *y* jouit d'une forme de priorité sur *x*.

4. Une classification des entités complexes

Nous nous sommes servi d'exemples méréologiques pour critiquer certaines théories de la dépendance existentielle. Maintenant que nous avons à disposition les relations de fondation et une nouvelle relation de dépendance, nous allons pour ainsi dire renvoyer l'ascenseur à la méréologie, et essayer de caractériser les entités composées en termes de fondation et de dépendance.

Il est plausible de dire que toute entité complexe existante est fondée sur chacun de ses composants: les composants d'une entité font exister (au moins en partie) cette entité. Ceci étant dit, il

peut y avoir *a priori* quatre catégories d'entités complexes existantes, selon qu'elles sont *totale*ment fondées sur leurs parties ou pas, et selon qu'elles sont dépendantes de leurs parties ou pas:

- (1) Entités ni totalement fondées ni dépendantes;
- (2) Entités totalement fondées mais non dépendantes;
- (3) Entités non totalement fondées mais dépendantes;
- (4) Entités à la fois totalement fondées et dépendantes.

À la catégorie (1) appartiennent peut-être certains tous «naturels» comme les tas de sable: un tas de sable n'est pas totalement fondé sur les grains qui le composent, car il doit son existence à plus qu'à l'existence des grains; et il n'est pas dépendant des grains car la destruction d'un grain n'entraîne pas la destruction du tas. À la catégorie (2) appartiennent certains⁵ événements et états de chose disjonctifs, et plus généralement certaines fusions disjonctives. À la catégorie (3) certains⁶ états de choses et événements concernant des individus substantiels. Enfin, à la catégorie (4) les ensembles non vides et les fusions conjonctives.

5. Quelques vérités sur la fondation, la fondation complète et la dépendance

Il ne s'agit pas là de proposer une théorie compréhensive de ces notions, mais simplement d'en donner la saveur.

Utilisons comme primitives les notions de fondation et de fondation complète. La première sera exprimée par le prédicat binaire «... est fondé sur...» – symbolisé par *F*. Pour la deuxième, nous devons tenir compte du fait qu'un nombre arbitraire d'objets peuvent totalement fonder un autre objet. La notion de fondation complète sera donc exprimée par un prédicat binaire dont le deuxième membre doit être saturé par un terme d'ensemble

5 L'événement *a-ou-b* est dépendant de *a* si *b* dépend de *a*. Et donc *a-ou-b* dépend de *a* et de *b* si *a* et *b* sont mutuellement dépendants.

6 L'existence de *a* est un état de chose qui est totalement fondé sur *a*.

«... est totalement fondé sur les membres de...». Ce prédicat sera symbolisé par TF.

Les principes que nous allons formuler sont modaux, et nous devons donc choisir une logique modale quantifiée pour le faire. On pourra supposer qu'il s'agit de la plus simple des logiques modales quantifiées, S5 avec domaine constant. L'opérateur de nécessité (métaphysique) sera symbolisé par L. Chacun des axiomes mentionnés peut être considéré comme nécessairement vrai, et l'on pourra donc prefixer un opérateur de nécessité à chacun d'eux.

Définitions:

$$x \text{ PTF } y := \exists \alpha (y \in \alpha \wedge x \text{ TF } \alpha)$$

x est partiellement totalement fondé sur $y := x$ est totalement fondé sur des objets dont y fait partie.

$$Fx := \exists y x F y$$

x est fondé := quelque chose fonde x .

$$\text{TF}x := \exists \alpha x \text{ TF } \alpha$$

x est totalement fondé := quelque chose fonde totalement x .

$$x D y := L (E!x \supset x F y)$$

x dépend existentiellement de $y :=$ nécessairement, x existe seulement s'il est fondé sur y .

Axiomes:

$$(1) x F y \supset (E!x \wedge E!y)$$

La relation de fondation est factive.

$$(2) x \text{ PTF } y \supset x F y$$

Être partiellement totalement fondé sur quelque chose implique être fondé sur cette chose.

$$(3) Fx \supset L (E!x \supset Fx)$$

Tout ce qui est fondé l'est nécessairement s'il existe.

$$(4) \quad TFx \supset L (E!x \supset TFx)$$

Tout ce qui est totalement fondé l'est nécessairement s'il existe.

$$(5) \quad (E!x \wedge \neg Fx) \supset L \neg Fx$$

Rien de ce qui existe non fondé ne peut être fondé.

$$(6) \quad (E!x \wedge \neg TFx) \supset L \neg TFx$$

Rien de ce qui existe non totalement fondé ne peut être totalement fondé.

Conséquences:

$$(1) \quad x TF \alpha \supset (E!x \wedge \forall y \in \alpha E!y)$$

La relation de fondation totale est factive.

$$(2) \quad (E!x \wedge x D y) \supset x F y$$

Tout ce qui existe et est dépendant d'une chose est fondé sur cette chose.

$$(3) \quad x D y \supset L x D y$$

Tout ce qui est dépendant d'un objet l'est nécessairement.

$$(4) \quad \neg x D y \supset L \neg x D y$$

Un objet ne peut pas être dépendant d'un autre s'il ne l'est pas actuellement.

$$(5) \quad (E!x \wedge \neg Fx) \supset L \neg x D y$$

Un objet existant non fondé ne peut pas être dépendant d'un autre.

$$(6) \quad x D y \supset L (E!x \supset E!y)$$

Un objet dépendant ne peut pas exister sans ce dont il dépend.

Bibliographie

- FINE K. (1995). Ontological Dependence. *Proceedings of the Aristotelian Society* XCV, Part 3, 269-290.
- LOWE E. J. (1998) *The Possibility of Metaphysics*. Oxford: Clarendon Press.
- SIMONS P. (1987) *Parts. A Study in Ontology*. Oxford: Clarendon Press.

ARE ALL ESSENTIAL PARTS ANALYTICALLY ESSENTIAL?

Peter SIMONS

1. Essentiality of parthood

What is an essential part? We need to disambiguate. To be or to have a part essentially, that is the question. From the point of view of the whole, a part is essential if that object, that whole, could not have existed unless it that part existed and was its part. In symbols:

$$N(Ew \rightarrow p < w).$$

This is the relation upon which I shall concentrate in this paper. It is to be distinguished from the other essential part-whole relation, where from the point of view of the part, it is essentially part of a whole, if it could not exist unless that whole existed and had that part as its part. In symbols

$$N(Ep \rightarrow p < w).$$

Call an object essentially dependent on another object if the first cannot exist without the second.

$$a \text{ dep } b := N(Ea \rightarrow Eb).$$

Then

$$N(Ew \rightarrow p < w) \rightarrow w \text{ dep } p$$

and

$$N(Ep \rightarrow p < w) \rightarrow p \text{ dep } w$$

So there are two notions of essential part and they reverse the order of dependence. Of course it is not ruled out that a part and a whole be essential to each other, in which case they are mutually dependent. In each case the dependence, and the essentiality of parthood, is rigid with respect to the second object. A whole w cannot exist without this individual p as its part; alternatively, the part p cannot exist except as part of this whole w .

It is quite a different matter if a whole needs a part of a certain kind but when it is not important which individual it is which supplements or satisfies its need. This is generic dependence. An adult human being needs some heart, but as we now know from transplant surgery, it is not essential that it be and remain the original one, so the dependence of human on heart is not modally rigid. Similarly with wholes: such *generic* dependence of part on whole or vice versa is important in its own right but is not what we are interested in here.

Let us then return to rigid dependence and look out for examples of essentiality of a part to its whole. Whatever else it has, a table has to have a top. If it didn't have a flat top, it wouldn't be a table. Further, a particular table T cannot exist unless its top P is its part, and indeed its top. Likewise a knife has to have a blade, as does a sword or a spade, as a chair needs a seat, as a boat needs a hull, and so on.

We have to be careful in using the expression «An A needs a B » because sometimes this describes mere generic dependence, as in «A human being needs a heart». A house needs a roof, but houses can be re-roofed. So the part in question has to be in some sense individually irreplaceable as part of this whole. Its loss or destruction must spell the destruction of the whole of which it is part. Also, if the whole in question had not had this individual as its part, it would not have existed.

The examples to date have been of artefacts, which are ontologically special. Let us consider some cases of essential parts among natural objects. A given atom has one or more nucleons, positrons or neutrons. Each positron or neutron is an

individual and if it is part of an atom is an essential part of that atom. That is

$$\forall x \forall y ((x \text{ is an atom and } y \text{ is a nucleon and } y < x) \rightarrow N(Ex \rightarrow y < x))$$

The loss of a given nucleon spells the end of that atom, and had the nucleon been combined with other nucleons than the ones with which it is in fact combined, then not that atom but some other atom would have existed.

Let us take a biological example. Like nearly all sexually reproducing eucaryotic organisms, we each originate from single zygotic cell produced by the fusion of an egg and a sperm¹ The zygote from which a human being comes makes up that human being completely, though not for long, as it soon ceases to exist by fission. Later, no particular cell is essential to a human being, but for that brief period, they coincide. The human being is coincident with but of course were not identical with the zygote². Then the zygote splits. Could either of the two cells into which it splits continue to exist without the other? The answer is «Yes». When such paired cells split but continue to exist the result is identical twins (or n -tuplets for some $n > 2$). Of the split cells, ones could die and the other continue to live. Our *original* cell is necessarily a part, indeed the whole of us, for that short period.

Biology is a slippery science, replete with counterexamples to every putatively true generalization. Brain, heart and lungs are parts which every normal exemplar of the human species has, and the loss of which normally spells death. But as we learn more and as medical science advances, the limits of the mereological changes that we know a human being to be able to sustain are extended.

Nevertheless, it seems that your brain is definitely an essential part of you. But then, who or what are you? Are you essentially an organism? Are you a Lockean person? Could you survive a brain transplant, in either capacity? I suggest we don't know clear

-
- 1 The exception would be if one were cloned, either naturally (as an identical twin) or artificially.
 - 2 The sense of 'coincident' is that they share all parts: see Simons 1987: 180. It follows that they are when coincident in the same place.

answers to these questions. So I conclude for the moment that we are not so sure about finding individual essential parts in the realm of biology, except for those cases where an organism is initially constituted by a single cell, which are not directly germane to our study since they make up the whole organism and not a part of it.

This lack of examples is illuminating. It shows that when we go humbly to nature looking for essential parts, she has many surprises for us, but few essential parts.

2. Analytically essential parts

Let us revisit a couple of these examples. Take the knife and its blade on the one hand and the helium atom and its proton on the other. In what does the essentiality of the parthood consist? If the knife K were to permanently lose its blade B , either because B were destroyed or because B were permanently separated from the knife handle H , K would cease to exist. If a knife K^* is made of a different blade $B^* \neq B$ then $K^* \neq K$. Similarly for the helium atom.

Recalling the uncertainty in the biological examples, how can we be so sure that a proton in a helium atom is an essential part of it? Why can we not envisage a counterfactual situation in which this very knife K comes into existence with a different blade B^* ? Why can we not imagine the helium atom surviving the loss of one of its protons, even if the proton were replaced by another one?

The reason we can be so sure is, I claim, that the concepts in question decide the issue for us, and since we are the authors of the concepts, we decide. That is why nature cannot surprise us or trick us into being mistaken. It is not that nature is very docile and fails to come up with a counterexample. On the contrary, we have so styled our language in this area that nothing nature can throw at us will count as a possible counterexample. Take a familiar example from elsewhere. The reason we can be so sure that no vixen is male is that the term 'vixen' has been determined to *mean* 'female fox'. Even if nature throws monsters and sexless or hermaphrodite foxes at us, we would withhold them the

appellation 'vixen' *because* they are not female. Likewise, it is our concept *knife* which incorporates the data that a knife's blade is essential to it. If a helium atom were to lose one of its protons, the remaining object would not be a helium atom but a tritium atom. If it were to exchange one proton for another one we'd still have a helium atom, but a *new* one. That is what we understand by a helium atom: those identity and persistence conditions, while they may be suggested by regularities of nature, are *imposed* by us and regulate our concept of a helium atom.

I stress that this does not mean that helium atoms *per se* are constructions of the human mind, in whole or in part. Rather, those things which exist independently and do what they do and change as they change, are delimited, individuated and identified by us according to a conceptual schema which is ours. What the schema reveals is not constructed but revealed, but if we did not use that schema, *those* things would not be revealed to us. Helium atoms do not depend on us or our conceptual schemata or individuating practices for their existence, but by employing those practices we reveal *them*.

The knife's blade is essential to it because our concept of knife stipulates that a knife's blade is a part of it and loss of the blade or creation with another blade means the end of the knife or the existence of another knife. We do allow a blade to be detached from the handle, especially temporarily for cleaning or repair, say, without detriment to the knife's continued existence. We might have had a slightly different conceptual scheme and different concepts, taking the separation of the blade from the handle to spell the end of that knife, and any later reassembly of the same blade and handle to mean a new knife came into existence, no matter how like the old one it is. On that alternative scheme, the knife could not survive dismantling. Our actual conceptual scheme is relatively liberal about dismantling. Many artefacts such as frame tents, Christmas Chimes and clarinets spend most of their lives dismantled. I think we often allow that a knife may receive a new handle and still be the same knife. In that case we have repaired the knife with a new part. The same would not be allowed of the blade: the blade is essential in a way the handle is not. But again one could imagine a concept of knife

similar to ours except that blade and handle are both equally essential.

Now in fact I am not so sure about our concept of *knife*: it is not so determinate on this point. There are two possible conceptual clarifications. According to one, the handle is definitely essential to the knife, according to the other it definitely is not. Which of the two more unequivocal concepts we prefer may depend on our interests and concerns. A butcher may not scruple to call a knife the same although it has been given a new handle, but a museum curator would be most alarmed if a new handle were given to a Bronze Age blade: she would deny that the original knife is preserved by such a change. So one concept is more practical, the other more antiquarian. The actual and merely statistically determined concept is a general working compromise between these two positions. Nevertheless both concepts agree on the essentiality of the blade as well as the inessentiality of small particles or incidental attachments such as say a carrying ring. For the extant and average concept let us call the handle a *penumbrally* essential part. It is neither clearly essential nor clearly not essential. Our concepts of particular sorts of artefacts often, perhaps always, carry such penumbra or areas of indeterminacy as to what is essential and what is not. Usually the concepts are averaged out over a range of more or less similar variants. For some people the engine might count as essential to a car and the body not, for others both, for yet others neither. «The» concept of car is an average across such variations of time, place, context and person, not a completely firm disposition. This is logically rather distracting but it is on further thought one of the most useful features of such everyday concepts, their very looseness lending them much of their utility. But the nature of the case may vary in quite detailed and small-scale ways, from one artefactual kind to another yet closely related one. For example there are multi-bladed knives like the famous Swiss Army knives. For such knives, it is quite likely that none of the several blades itself counts as individually essential: the knife would be a damaged and defective one without that blade, or could be repaired and continue to exist with a different blade. But again there is room

for manoeuvre: if one blade were much larger or more important than the others it might be taken to be essential and the others not.

The fact that our sortal concepts for many artefactual kinds have these areas of indeterminacy in what they count as mereological essential is a source of paradoxes of which the most notorious is Plutarch's and Hobbes's Ship of Theseus. In the face of such examples, two extreme positions may be adopted. One is to rule that any part is essential to any whole of which it is part. This is *mereological essentialism* (ME), and is most frequently associated with Roderick M. Chisholm (1976: 148-158), although the view can be found earlier³. According to this view, which typically has both a modal and a temporal component, nothing could have been made of anything other than the parts of which it is actually made, and nothing can either lose or gain a part and remain in existence. This view is so far out of alignment with our actual concepts of artefacts that it can only be considered worth adopting if these regular concepts lead us into such dire difficulties that the revision is worth the effort. It is precisely the office of Ship of Theseus style paradoxes to attempt to provide such strong grounds for revising our everyday conceptual practice regarding the essentiality of parts. However, if alternative plausible accounts of the paradoxes are forthcoming, then the stringent revisions required by ME may be avoided. I have argued elsewhere (1987: 198-204; 272-280). that the Ship of Theseus cases result precisely from the penumbral slack in our everyday concepts of artefacts, and that nothing ontologically untoward results, so that the difficulties of being forced to adopt ME may be avoided.

The other extreme is to say that nothing has any essential part. This is not to allow that anything might be composed of anything. A knife must still have *some* blade and a broom *some* brush, but which ones these are is not fixed. We retreat in other words to generic rather than rigid dependence. This is also too extreme, because examining many of the artefact concepts we actually employ shows that some parts are essential to some artefacts of some kinds: its blade to a knife, its top to a table, their lenses to a pair of spectacles. That is just how it is. We might have had other

3 Leibniz seems to have upheld mereological essentialism earlier.

concepts but in fact don't. Nevertheless such a reform away from accepting any essential parts among artefacts would be less radical than adopting ME.

In the area of natural kinds too we notice that some of the concepts we actually employ accept some individual parts as essential, as in the helium-proton case, or a water molecule and its three constituent atoms. But here too the three different cases we noticed for the knife may occur. There are several isotopes of helium, including one, Helium-3, which has two protons but only one neutron. Normally helium arises from tritium (a hydrogen isotope with two neutrons) by beta decay. But although it does not physically happen this way, were an atom of Helium 4 to lose a neutron, and still be the same Helium atom, then that neutron would not be an essential part of the helium atom, nor by parity would the other one be. Only the protons would be essential to the helium atom. We could leave it indeterminate whether each neutron is essential to a helium atom. We could take them to be essential (so if a helium-4 atom lost a neutron then *that helium atom* would cease to exist and be replaced by another, a helium-3 atom. Or we could definitely take the neutrons to be inessential. Nothing that happens in nature is made any different by our choosing among these alternatives. We are simply deciding where to draw lines around individuals, *which* individuals to recognise out of the ones we could recognise. In theory we could recognise both helium atoms which could survive the loss of a neutron alongside ones which could not: we would need different expressions for them to avoid confusion but that is all. The in-between case of penumbally essential neutrons would be a half-way house between these two alternatives, and might be accepted for practical reasons, or suggested by the weight of observed phenomena. Since as far as I know He-4 never decays into He-3, there is no practical need for a decision. Such penumbra of essential parthood are more frequent higher up the periodic table where matters become more complicated.

There are scientific phenomena which are relevant to our choice of concepts, for example what isotopes exist in nature and the extent to which chemical properties mark a clear distinction between distinct kinds (such as the distinction between He-3 and

tritium, H-3). We tailor our concepts of natural kinds according to what nature puts before us in the way of regularities, and over these regularities we have no power of decision, but which particular sortal concepts we fit within these constraints is not always uniquely prescribed.

3. Biological kinds

The examples I have given of essential parts among natural kinds are chemical. Philosophers often use the term 'natural kind' to cover also biological species such as the lion and higher taxa such as the mammals. This usage is unfortunate and ought to be discontinued. Biological species, with the limited exception of asexually cloning or agamospecies do not form natural kinds in the sense of Putnam and Kripke, classes of things sharing a microstructure. This is true despite the genetic similarity between members of the same species and the genetic affinities between members of the same higher taxa, genus, family, class, order, phylum and so on. The reason is that biological kinds are too fluid and mobile in their characters to be captured by necessary and sufficient membership conditions. Biological kinds are delimited over time genealogically, by historical facts of origin and descent, and they are delimited at a time by relations of co-reproductivity, and these are vague relationships. They have no stable necessary essence which is guaranteed to remain unchanged, and in this respect they are quite different from the natural kinds dealt with in physics and chemistry. For example *Homo sapiens* is factually distinguished from other animals by a large brain, bipedal gait, opposable thumbs, and the use of language. None of these characters is essential. Not only are there humans lacking one or other of these characters, the whole species could develop by genetic drift under environmental pressure in the direction of losing any of them, and provided this development were not accompanied by extinction or speciation, *Homo sapiens* would continue to exist in the absence of the trait in question. So despite its name, *Homo sapiens* is not necessarily brainy or rational.

Returning to mereology, it is rather clear that no part of any organism is permanently essential to it. We know of course that the flux of small parts in a large multicellular organism such as a human being means that none of these small parts is essential. We know from numerous cases of accident, illness and defect that human beings can accidentally lack larger parts. Finally we know from advances in medical technology that few indeed of our actual parts are individually essential, and there is no telling what the future may bring in this regard. That is why a question mark hangs over even the most likely putative essential part of a human being, the brain. But whatever the status of the brain for an adult or reasonably grown human, we know for certain that the brain is not essential to any individual human being, because in the earliest stages of their existence as an embryo no human has a brain.

Up to now I have not distinguished, as Locke advised us to do, between human beings and human beings who are persons⁴. Retention of mental abilities despite a complete change of body and brain would incline us to talk of the same *persons* rather than the same *human being*, and there would be very good grounds for keeping those two cases distinct, as they would have distinct persistence conditions. At present most of the persons we know about are humans. However I stress that these thought-experiments are speculative. Given our present state of relative ignorance it would be unwise to offer strong predictions on what will happen should such experiments eventually become practical. All I would stress is that the term 'person', unlike the term 'human being', is very open-textured and not obviously a substance sortal at all.

4. No non-analytic essential parts

This very uncertainty that we have at present underlines the sceptical conclusion towards which I am moving. My hypothesis is that there are no clear cases of essential parts where the nature

4 I say 'person who is a human being' rather than just 'person' because there is no reason why humans should have any monopoly on personhood.

of the necessity is not analytical, that is, where it does not rest in the concepts we employ or equivalently in the particular lexical meanings of certain uses of words. That this or that individual is an essential part of something is so of analytic necessity and not because of natural, metaphysical or logical necessity. It is analytically true that this table must have this top, or that this helium atom must have this proton as part; in just the same way as it is analytically necessary of this electron that it has a negative charge.

The comparison is chosen deliberately to highlight what is in our power and what is not. The electron, a light particle with negative electric charge, was discovered by J.J. Thompson in 1897. Its antiparticle the positron was predicted theoretically by Paul Dirac in 1931 and observed (uninfluenced by Dirac's prediction) by Carl Anderson in 1932. The positron differs from the electron only in having positive rather than negative charge, of the same magnitude. In the early papers positrons were frequently conceived, as in Dirac's original idea, as not particles but as «holes» in a sea of negative electricity. But even where they were regarded as particles, they were called 'positive electrons'. During this stage, anyone who stated that some electrons were positively charged would have been speaking correctly. By 1933, Anderson had coined the term 'positron' and the term 'electron' was thenceforth restricted in the physics community to the negatively charged particles. After 1933, it became analytically true to say that electrons have negative charge: before that time, the lexicon had not fixed this property, even before the discovery of positrons. The term 'electron' underwent a change in meaning.

The point is that neither electrons nor positrons were in any way changed in their own nature by the terminological convention adopted by physicists in the 1930s. They are as they are and behave as they behave irrespective of our names for them or how we draw the boundaries around kinds. It would have made perfect sense to continue to use 'electron' for both positively charged and negatively charged particles and the terminology of particle physics would have been slightly different from what it now is. So the nature of the necessity involved in saying that electrons are essentially negatively charged is analytic: it is not

logical or metaphysical or natural, though of course regularities motivate our terminological choices. The difference between ourselves and a physicist of 1925 on the question whether electrons are essentially negative reflects that it is within our power to legislate where some word's extension runs, and therewith to include some properties as essential. We *discover* that electrons *are* negative, but *legislate* that they *must* be.

It is, I suggest, the same with essential parts. There are many cases of kinds of things where we see them preserving certain parts through changes and this is a basis for but not yet sufficient for us to be sure that all such objects must have such parts. That helium atoms have their protons essentially lies in the concept *Helium*, or equivalently, in the way we use the word 'helium' and its synonyms. This usage is strongly guided by natural regularities, but these never confer the certainty of necessity. The reason we can be so sure that the loss of a proton would spell the end of a helium atom is not because we have some intuitive insight into a law of nature or metaphysics, but because we know that is how we use the word 'helium': we are parties to that usage.

Before moving on I wish to discuss two concerns about my examples and thesis. The first is that we might be accused of naivety in treating electrons, protons etc. as if they were little material bodies when in fact their truer and more elusive nature is revealed by the wave-particle duality of quantum theory. If we could peer closer into a helium nucleus we should not see discrete little points but smeared somethings which are the result of the superpositions of eigenstates of various fields. Does this not mean that talk of part and whole is inappropriate when we get down to physical fundamentals? Of course such a possibility cannot be ruled out, but when field quantities can be summed as the resultant of quantities which have a straightforward interpretation in isolation, we are still entitled to treat the contributory phenomena as parts of the whole, as when two superposed water waves or sound waves may be seen as the resultant of two distinct simpler waves. So even something as unparticulate as a wave may be considered – *sensu strictu*, not by analogy or metaphor – as having parts.

This brings me to a second concern, which is related to the first. The trope ontology within which I consider it prudent to interpret and situate the theory of field quantities is one where the distinction between a part and a quality instance breaks down. If, as I believe, a wave-particle like an electron is a nexus or bundle of tropes unified by formal foundation relations, then the negative charge of an electron is a part of it, but a dependent, not detachable part. So the electron case and the helium case are not so different after all: a helium atom is simply a more complex bundle, with qualitatively identifiable sub-bundles, whereas the electron appears not to exhibit internal sub-constituents apart from its actual tropes, taken as a single unified nexus.

5. Essential parts of occurrents

Of course not just continuants, that is, things which endure but lack temporal parts, have parts. Occurrents, things which endure through time by the accretion over time of temporal parts or phases, also have parts, and not just temporal ones. An explosion or a collision is an event with different temporal phases as well as spatial parts and spatio-temporal segments. Similar remarks apply to all manner of occurrents: smiles, processes of growth or digestion, all have manifold parts, including temporal phases with no part in common. What of their essential parts?

It seems at first sight that our intuitions about what parts occurrents have essentially are much hazier than for the case of continuants, not least because we speak about occurrents less frequently and with less intense interest than continuants. One position which at first seems extreme is that all parts of occurrents are essential to them⁵. This seems less in conflict with our views about what occurrents are than the corresponding view about continuants.

Let us illustrate the point with an example. Suppose Lucy takes a walk to her local shop to buy a newspaper. She can walk faster or slower as she wishes, she can stop to admire her neighbour's roses or new car, talk to the children next door, cross

5 This position is canvassed in Simons 1987: 281-3.

the road to read a notice, and so on. Suppose ME applied to all such events. Then no part of the walk which actually took place could not have taken place. So if a counterfactually mereologically distinct walk took place, say one lasting 10 seconds less, this would have to be a distinct walk. But if Lucy simply walks more slowly to look at the roses because she wants to, but all other features of the walk are the same, are we really prepared to say it is a wholly different walk rather than the same walk with a few different parts in the counterfactual situation? Both positions seem quite reasonable in fact: it is not clear whether we have any solid intuitions here.

Take the continuant Lucy herself. She happens to have a non-essential part that some people lack: an appendix. She might have had appendicitis at some earlier age and have had her appendix removed, but she did not. But no one would doubt that in the situation in which Lucy had lost her appendix, that it would have been the same woman. All the events that befall Lucy's appendix are parts of her life, and if she had lost her appendix earlier those parts would simply not have taken place. So if we consider the large complex occurrent *Lucy's life* to be governed by mereological essentialism then it would seem that since the parts of her actual life involving her appendix could not have been absent, so Lucy could not have lost her appendix, so her appendix is after all an essential part of her. It seems then that ME for occurrents is incompatible with the rejection of ME for continuants, since the lives of continuants are occurrents. Even worse, ME for occurrents seems to commit us to determinism for continuants. Surely this is too high a price to pay for ME for occurrents.

However, there is a reply to the objection and it goes like this. Many, in fact nearly all of our names for events are definite descriptions or contain other reference-determining descriptive content. There are apparent exceptions such as names of battles, usually borrowed from the names of the places they occurred, but even a name like 'Waterloo' should be construed as meaning 'the Battle of Waterloo'. Definite descriptions are not modally rigid, but modally *flaccid*⁶. That means in the idiom of possible worlds

6 Honderich 1995: 282 (entry by Wayne Davis).

that they may designate different individuals in different worlds. 'The 34th President of the United States' refers to Dwight D. Eisenhower, but it might have referred to Thomas E. Dewey had the 1948 election gone his way. So the expressions 'Lucy's walk to the shop on 17 October 2000' and 'Lucy's life' are also modally flaccid and may designate different events in different worlds. Lucy's life might have been quite different. She might have been born in Australia rather than England if her parents had decided to emigrate in 1955, she might have become a mathematics teacher in Melbourne rather than a bank manager in Leeds, might have met and married a different man and had different children. None of those events described in the alternative possible life of Lucy are events in her actual life, yet it would still be her. So determinism and incompatibility with the denial of ME for continuants are avoided.

There remains a more subtle issue. We cannot counterfactually vary the events of a life arbitrarily and it still remain the life of that thing. Lucy could not have been born to Chinese parents in Shanghai in 1930 if she was born to English parents in Leeds in 1955, any more than a helium atom might have been a water molecule or a tree might have been a mouse. Are there any occurrents in Lucy's life in which she *must* participate? I suggest there is *at most* one, and that is her actual creation. That actual fusion of those two cells was her origin. But again 'Lucy's origin' might be a modally flaccid term: had those two cells fused earlier or later or in a different place, would it have been numerically the same event? We can agree that fusion of different cells would not have been Lucy's origin. But provided it was the very same cells which fused, I think we are happy to identify the resulting human being as Lucy. We can therefore imagine Lucy's life starting an hour earlier or later than it in fact did, and it still being her life, but a numerically different one, with in fact no single occurrent in common with her actual life, despite whatever descriptive similarities might reign.

Here is a third objection⁷. Take a very complex event like the Battle of Waterloo, and consider a small and insignificant part of it, such as a soldier's firing a certain bullet which let's say misses

7 Due to Christopher Hughes.

its target. Perhaps the soldier brought that bullet rather than another one quite accidentally. Could the soldier, thinking about this possibility, reasonably think to himself, «If I had brought the other bullet rather than this one, I could have prevented this battle from taking place»? Clearly there is something odd about this idea: certain things could have prevented the battle, such as Wellington failing to get his troops to Waterloo, or Napoleon deciding to avoid a fight, but the trivial action of a minor participant seem not to be in this league.

The answer is in this case that the soldier can and indeed if I am right correctly *should* think that thought, but that there is a huge difference between preventing *this* (particular, individual) battle from taking place, and preventing *any* battle from taking place. It is a trivial matter to as it were divert history so that a different battle of Waterloo takes place: any one of the many participants could do that by any one of many alternative possible actions. Numerically it is then no longer the same battle: we must remain strict. But in the midst of a raging battle no action or alternative action can bring it about that *no* battle takes place. Once something denominable as a battle has started (and it takes enough of the right kind of events coming together for it to be a battle and not preparations or skirmishing) not even divine intervention can cause no battle to have happened. The objection turns on not taking the individuality of the actual battle seriously.

I conclude that there is no obvious intuitive obstacle to ME for occurrents, and that this position fits in well enough with our somewhat shaky intuitions on the subject.

6. Ontological primacy of occurrents no obstacle

I have argued elsewhere (2000a & b) that occurrents are ontologically prior to continuants because with continuants alone we are unable to find suitable truth-makers for propositions of the form *c exists at t* where '*c*' names a continuant, but occurrents which are vital to the existence of *c* and which have temporal parts at *t* provide us with such truth-makers. Continuants however do not fail to exist, as eliminativist and four-dimensionalist

ontologists would have us believe: rather they are invariants over collections of occurrents under certain equivalence relations. This holds for all continuants without exception: it holds for more complex continuants like Lucy the lady as well as simpler continuants like Eric the electron. But does not the dependence of continuants on their occurrent basis mean that ME for occurrents again forces us to accept ME for continuants after all? No it does not, for the dependence of a continuant on its vital occurrents is generic. *It*, that very same continuant, could have had other vital occurrents, provided only they were of the right kind. Lucy could have had a completely different collection of heartbeats sustaining her than those that in fact sustain her, and it still be Lucy. We see here the value of a clear distinction between rigid dependence and generic dependence in ontological analysis. Thus the dependence of a continuant on its vital and other life-occurrents is a rather light and liberal affair, and it is this relatively relaxed relationship which allows continuants to be counterfactually identified across variant occurrent bases. In this respect therefore, though occurrents are generically necessary and basic to continuants, the dependence is not so stringent that one is justified in claiming continuants could be *reduced* to occurrents. If they could, surely continuants' identities would not be so robust. That continuants are invariants does not mean they are created by us or discerned only as abstracta, indeed as Strawson has argued, recognition of continuants is more stable to our cognitive and conceptual system than recognition of occurrents, even though I hold there can be no continuants without occurrents. Epistemological or cognitive primacy is no obstacle to ontological secundariness. There are deep issues involved here about the relevance of cognitive choices to ontology, and I have only touched on some of them in connection with the question of our freedom of delimitation, but I cannot penetrate the issues more deeply here in the course of another topic and must leave them to another occasion.

7. Mereological essentialism given its rightful place

So here is my somewhat ironic conclusion with regard to Chisholm's mereological essentialism. ME is correct for occurrents, that is, objects with temporal parts, and it is so because our whole conceptual scheme accords with it. The irony is that Chisholm himself rejects ontologies based on occurrents, such as Rescher's process ontology⁸. That occurrents have all their parts essentially is however itself an analytic truth, concerning the concept of an occurrent, so my main thesis that all essentiality of parts is analytic is not undermined. Now it is too facile to say that this is because of the way we have decided to use terms like 'occurrent', 'event' etc. It is decidedly *not* like the helium case. Expressions like 'event', 'process', 'occurrent' are highly abstract and to some extent philosophically artificial, so to simply seek to pass the responsibility onto some local or limited decision or convention is very implausible. Rather we should look at global features of our use of event and other occurrent terms and indeed our use of verbs, adverbs etc. in sentences which don't name events and othe occurrents. This is indeed deep water. But in principle there is nothing except attachment to simple and simple-minded cases to be said against the idea that some things can be analytically true and yet very unobviously analytically true⁹. That its blade is essential to this knife is, let us be honest, a very boring and trivial example, a superficial aspect of our concepts regarding this type of artefact. It cuts little ice, so to speak. On the other hand, if it is analytically true of any occurrent and a part it has that it has this part, then this is almost certainly deeply inherent of our conception of occurrents as entities spread out over time, and not a superficial feature of the lexeme 'occurrent'. Concepts and their features are manifested in regularities of linguistic and cognitive practice as well as in the lexicon we command, so the concepts and their features can therefore be relatively opaque to us. Part of our task as philosophers is to bring such latent complexity to patency.

8 Chisholm, 'Reply to Rescher', in Hahn (ed.) 1997: 200.

9 Cf. Bennett 1966: 42.

If, as I have argued, ME applies to occurrents, then we have a clear way of seeing how it is that continuants may nevertheless participate in the events of their lives accidentally. Any two possible lives of Lucy must either have some initial part in common, or the materials of which Lucy is originally made must have some initial parts of their lives in common. The former case is the more straightforward and usual. Had Lucy stopped to talk to her neighbour rather than walking past on the way to the shop, her later life would have been numerically and in all its parts different, whatever the similarities of the later stages. We can envisage Lucy's possible lives after a given time as like a branching tree, and the events which determine how things turn out in large or small ways as pruning the tree so that only one path remains up to the present¹⁰.

The more complex case allows that the actual time and place of creation of an individual be not essential to it, but that all of its initial materials be essential to it. It would be conceptually neater if the originating event of a continuant be essential to it in all cases, but our concepts as a matter of fact do not operate this way.

8. Conclusion

I have used our fairly robust intuitions about what is or is not essential as a part to a whole to argue that all essential parts are so analytically, either for local lexical reasons or because of the character of the general notion of an occurrent in our thinking about the world. The thesis put forward strongly suggests a more radical general position, that all essentiality, whether of part or property or origin or material composition, or whatever else, is analytic. I have not argued for this more general thesis but I suggest that arguments can be put forward along the lines given for essentiality of parts. This in turn suggests a more radical position still, namely that there are only two kinds of real necessity. The first is lexical necessity, as in 'Necessarily, all

10 This idea is found in great clarity in McCall 1994. Unlike McCall I regard the tree idea as a *façon de parler* and not as designating anything metaphysically real.

vixens are female', or 'Necessarily, no handkerchief is a colour'. These are analytic truths, but trifling ones, to use Locke's derogatory description. The second and much more interesting kind of necessity attaches to groups of concepts which are very general and deeply embedded in our way of thinking about the world. The mereological essentialism of occurrents or the essentiality of origin of artefacts belong here, as does the special kind of necessity attaching to that slightly fuzzy group of concepts called the logical constants, and which gives us logical necessity. Whether we choose to call such broadly conceptual necessity 'analytic' or 'synthetic' is probably a terminological matter, but my preference would be to call it 'analytic', because it allows us to formulate in simple terms the deflationary statement that all necessity is analytic. I have not argued for this, but I am inclined to believe it and hope to have given some grounds for supporting it¹¹.

School of Philosophy
University of Leeds
LEEDS LS2 9JT England
e-mail: p.m.simons@leeds.ac.uk

11 Thanks to participants at the Neuchâtel conference as well as members of the Philosophy Departments of Sheffield, Tufts and Dortmund Universities for their comments.

References

- BENNETT J. (1966). *Kant's Analytic*. Cambridge: Cambridge University Press.
- CHISHOLM R.M. (1976). *Person and Object*. London: Allen & Unwin.
- HAHN L.E. (ed.) (1997). *The Philosophy of Roderick M. Chisholm*. La Salle: Open Court.
- HONDERICH T. (ed.) (1995). *The Oxford Companion to Philosophy*. Oxford: Oxford University Press.
- MCCALL S. (1994). *A Model of the Universe*. Oxford: Oxford University Press.
- SIMONS P.M. (1987). *Parts*. Oxford: Oxford University Press.
- SIMONS P.M. (2000a). Continuants and Occurrents. *The Aristotelian Society*, Supplementary Vol. LXXIV, 78-101.
- SIMONS P.M. (2000b). How to exist at a time when you have no temporal parts. *The Monist* 83, 419-436.

PARTS, COUNTERPARTS, AND MODAL OCCURRENTS

Achille C. VARZI

A modal occurrent is an individual entity that stretches across possible worlds. It is a trans-world individual. It consists of parts that are located in different worlds just as ordinary objects consist of parts that are located at different places, or processes and other events (so-called temporal occurrents) consist of parts that are located at different times. If we do not think that quantification over possible worlds is intelligible, then the notion of a modal occurrent will be unintelligible to us. (Some philosophers feel that way about the notion of a temporal occurrent.) If we think that quantification over possible worlds is intelligible but that mereological fusions should be restricted to entities existing in the same world, then to us the notion of a modal occurrent will be intelligible but empty. (Three-dimensionalists feel that way about temporal occurrents, at least insofar as these are meant to include objects besides events.) However, if we think that quantification over possible worlds is intelligible and that mereological fusions should not be so restricted, then we are well off. For then we can make sense of modal talk even if we cannot make sense of trans-world identity.

The purpose of this paper is to illustrate how this can be done. First I will lay out the basic apparatus of the theory of modal occurrents. Next I will show how to interpret modal talk extensionally within the apparatus regardless of whether trans-world identity is allowed, i.e., regardless of whether distinct worlds can overlap. Then I will show that ruling out trans-world identity by requiring possible worlds to be pairwise discrete is all that is needed to reconstruct the bulk of counterpart theory (i.e., to define the concept of a counterpart and to derive the standard postulates governing that concept). This will throw some new light on

counterpart theory itself. Finally, I will argue that the account allows us to see the indeterminacy of our modal intuitions as being part and parcel with the indeterminacy of our criteria for individuating modal occurrents, and that this indeterminacy is best explained in terms of semantic (as opposed to theoretical or ontic) vagueness.

Mereological preliminaries

I will assume a standard mereological background constructed around the binary relational predicate '*x* is part of *y*', written '*Pxy*'.¹ I take this predicate to be true when *x* is any sort of part of *y*, including an improper part, so that *Pxy* will be consistent with *x* and *y* being the same. The predicate for proper parthood, along with other familiar mereological predicates, can be introduced by definition in the usual way:

- (1) $PPxy =_{df} Pxy \wedge \neg x=y$
x is a proper part of *y* iff *x* is a part of *y* but is distinct from *y*.
- (2) $Oxy =_{df} \exists z(Pzx \wedge Pzy)$
x overlaps *y* iff something is part of both.
- (3) $POxy =_{df} Oxy \wedge \neg x=y$
x properly overlaps *y* iff *x* overlaps *y* but is distinct from *y*.

As axioms I will assume those of standard extensional mereology, except for the principle of unrestricted fusion. Thus, I will assume parthood to be a reflexive, antisymmetric, and transitive relation—a partial ordering—subject to a strong supplementation principle:

- (4) $\forall xPxx$
 Everything is part (viz., an improper part) of itself.

1 For an overview of mereology see Simons 1987.

- (5) $\forall x \forall y (Pxy \wedge Pyx \rightarrow x=y)$
No two things are part of each other.
- (6) $\forall x \forall y (Pxy \wedge Pyz \rightarrow Pxz)$
The parts of a thing's parts are themselves part of that thing.
- (7) $\forall x \forall y (\forall z (Pzx \rightarrow Ozy) \rightarrow Pxy)$
If every single part of x overlaps y , then x itself is part of y .

I shall not assume the axiom of unrestricted fusion but I shall not impose any special conditions pro or against any kind of mereological fusion, either. In the present context I will leave the issue open. If, given a condition ϕ , there is a mereological fusion of all the ϕ ers, then it is unique by (7) and we can refer to it by a definite description²:

- (8) $\sigma x \phi =_{\text{df}} \iota z \forall y (Ozy \leftrightarrow \exists x (\phi \wedge Oyx))$
The fusion of all x such that ϕ (if it exists) is that thing which overlaps all and only those things which overlap some x such that ϕ .

The notation for finitary sums is introduced accordingly:

- (9) $x+y =_{\text{df}} \sigma z (z=x \vee z=y)$
The sum of x and y is the fusion of x and y (if it exists).

Modal occurents

The theory of modal occurents can be built on the basis of mereology by adding two more primitives, namely, a one-place predicate 'Wx' (read: ' x is a possible world') and an individual constant α (for 'the actual world'). These primitives are governed by the following axioms:

2 I will assume descriptions to be handled in standard, Russellian fashion.

- (10) $W\alpha$
 α is a world—the actual world.
- (11) $\forall x\exists y(Wy \wedge Oxy)$
 Everything overlaps some world.
- (12) $\forall x(Wx \wedge \exists yPPyx)$
 Every world includes something.

The first axiom is to fix the intended role of ' α ' and the second axiom is to fix the intended meaning of 'W' on the understanding that a world is the mereological fusion of everything in it. The third axiom is only needed to go along with the standard assumption that there are no «empty worlds», an assumption which generalizes—modally—the assumption that the domain of discourse is non-empty. Strictly speaking this assumption is not essential and one could allow for empty worlds just as one can allow for the empty domain; but this would introduce complications into the logical machinery which would only becloud the basic picture³. On the other hand, there is no need to require that all worlds overlap, or that no two worlds overlap, for such requirements do not belong to a general formal theory of worlds. For the same reason we need not require that worlds be maximal individuals, i.e., that no world is part of another.

With this apparatus in hand, a modal occurrent is defined, quite simply, as any object in the domain of quantification whose mereological make up does not include any worlds:

- (13) $Mx =_{df} \forall y(Wy \rightarrow \neg Pyx)$
 x is a modal occurrent iff no world is part of x .

Typically, a modal occurrent is a trans-world individual—i.e., it has fragments that are too large to be mereologically included in a single world. This is how the concept has been originally introduced in the literature (though a preferred idiom is David Lewis's 'modal continuant'⁴). However, our definition includes world-

3 On these matters I refer to Garson 1984.

4 The terminology is introduced in Lewis 1983: 41. I prefer my idiom because it parallels the distinction between occurrents and continuants familiar from the literature on

bound individuals as a limit case and is therefore to be understood very broadly. We can always introduce the limit case by definition:

- (14) $PMx =_{df} \forall y(Wy \rightarrow \neg Pxy)$
 x is a proper modal occurrent iff x is (a modal occurrent which is) not part of any world y .

Given our modest assumptions concerning possible worlds, it is also possible for a trans-world modal occurrent to be world-bound, provided there is a world large enough to contain it all. But whether or not a modal occurrent is world-bound, it certainly has world-bound parts, the largest of which may be thought of as the «actualizations» (or stages) of that occurrent in the relevant worlds. We can make this more precise through the following auxiliary definitions:

- (15) $WBxy =_{df} Mx \wedge Wy \wedge Pxy$
 x is world-bound in y iff x is a modal occurrent and y is a world wholly containing x .
- (16) $WFxyz =_{df} Pxy \wedge WBxz$
 x is a world fragment of y in z iff x is a part of y which is world-bound in z .
- (17) $WSxyz =_{df} WFxyz \wedge \forall u(WFuyz \rightarrow \neg PPxu)$
 x is a world stage of y in z iff x is a maximal world fragment of y in z .

Some immediate corollaries of these definitions:

- (18) $\forall x \forall y (WBxy \rightarrow WSxxy)$
 Every world-bound occurrent is a world stage of itself (in the relevant world).

temporal persistence: a temporal occurrent is an entity that has temporal parts whereas a temporal continuant is an entity which is present in its entirety at any time at which it exists. By analogy, then, I would use 'modal continuant' to refer to those items which are present in their entirety in any worlds in which they exist—exactly the opposite of a (proper) modal occurrent.

- (19) $\forall x \forall y \forall z (WS_{xyz} \wedge WByz \rightarrow x=y)$
 Every world-bound occurrent coincides with its own world stage (in the relevant world).
- (20) $\forall x \forall y \forall z \forall u (WS_{xyz} \wedge WS_{uyz} \rightarrow x=u)$
 Every occurrent has at most one world stage in any world.

Given (20), we can introduce a definite description to talk about world stages (where they exist):

- (21) $x^y =_{\text{df}} \iota z WS_{zxy}$
 The unique stage (if it exists) of occurrent x in world y .

Modalities

To interpret modal discourse against the background of the formal theory of modal occurrents, we interpret names and predicates as usual, i.e., as objects and relations defined on a universe of discourse. In particular, ' α ' will pick out an object in the extension of ' W ', i.e., a world. However, the intuition is that ordinary names such as 'Pavarotti' correspond to proper modal occurrents and ordinary predicates such as 'is a tenor' correspond to sets of world stages. (Ordinary relational predicates such as 'is a better tenor than' will correspond to relations among world stages.) A statement such as

- (22) Pavarotti is a tenor

will then count as true iff the actual world stage of Pavarotti is a tenor. And a statement such as

- (23) Pavarotti might have been a ballerina

will count as true iff some world stage of Pavarotti is a ballerina⁵. This intuition is familiar from the literature on four-dimensionalism, according to which a ordinary proper name such as

5 This is the intuition put forward in Lewis 1983: 40–42. See also Lewis 1986: 213–217.

'Pavarotti' denotes a temporally extended individual and a statement of the form

(24) Pavarotti will be a ballerina

is true iff there is a future temporal part of Pavarotti which is a ballerina⁶. In effect, if we construed possible worlds intuitively as temporal stages of the actual world, then the theory of modal occurrents outlined here would yield a theory of temporal occurrents which is compatible with much literature on the topic. (This would be especially evident if we assumed all worlds to be pairwise discrete, for typically a four-dimensional ontology has no room for continuants, i.e., things wholly existing at different times.)

Keeping with this basic framework, the intuition can be generalized as follows. Let L be the given modal (first-order) language and let L^* be the non-modal language that results from L by replacing the modal operators with the primitive vocabulary of mereology and the theory of modal occurrents. For every formula ϕ of L we define a corresponding L^* -formula ϕ^y (ϕ holds at y) by recursion⁷,

- (25) if $\phi = Rt_1 \dots t_n$, then $\phi^y = Rt_1^y \dots t_n^y$
 if $\phi = \neg\psi$, then $\phi^y = \neg\psi^y$
 if $\phi = \psi \cdot \chi$, then $\phi^y = \psi^y \cdot \chi^y$
 if $\phi = \forall x\psi$, then $\phi^y = \forall x(E!x^y \rightarrow \psi^y)$
 if $\phi = \exists x\psi$, then $\phi^y = \exists x(E!x^y \wedge \psi^y)$
 if $\phi = \Box\psi$, then $\phi^y = \forall x(Wx \rightarrow \psi^y)$
 if $\phi = \Diamond\psi$, then $\phi^y = \exists x(Wx \wedge \psi^y)$

where ' $\cdot$ ' is any binary connective and ' $E!$ ' is the existence predicate defined by:

-
- 6 Lewis himself is a proponent of a four-dimensionalist ontology in which tensed statements are interpreted according to this schema. See e.g. Lewis 1986: 204.
 7 The first conjunct in the clause for ' $\exists$ ' is actually redundant if descriptions are understood *à la* Russell. We leave it in to ensure immediate duality between ' $\exists$ ' and ' $\forall$ '.

$$(26) \quad E!x =_{\text{df}} \exists y(y = x).$$

Then the translation of any given L-formula ϕ into L* is given by the function τ defined by setting:

$$(27) \quad \tau(\phi) = \phi^\alpha.$$

Here are some illustrative examples:

$$(28) \quad \tau(Fb) = Fb^\alpha$$

' b is F' means 'The actual world stage of b is F'. (Thus, if b has no world fragments in the actual world, then the statement is false.)

$$(29) \quad \tau(\forall x Fx) = \forall x(E!x^\alpha \rightarrow Fx^\alpha)$$

'Everything is F' means 'Every actual world stage is F'.

$$(30) \quad \tau(\Box Fb) = \forall x(Wx \rightarrow Fb^\alpha)$$

'Necessarily b is F' means 'In every world, b 's world stage is F'. (Thus, if there is a world in which b has no world fragments, then the statement is false.)

$$(31) \quad \tau(\Diamond \exists x Fx) = \exists y(Wy \wedge \exists x(E!x^y \wedge Fx^y))$$

'Possibly, something is F' means 'There is a world in which something has a stage that is F'.

$$(32) \quad \tau(\exists x \Diamond Fx) = \exists x(E!x^\alpha \wedge \exists y(Wy \wedge Fx^y))$$

'Something is possibly F' means 'Something which has an actual world stage has a possible world stage that is F'.

$$(33) \quad \tau(\forall x(Fx \rightarrow \Box Gx)) = \forall x(E!x^\alpha \rightarrow (Fx^\alpha \rightarrow \forall z(Wz \rightarrow Gx^z)))$$

'Every F is necessarily G' means 'Everything that has an actual world stage that is F has in every world a stage that is G'. (Thus, if something x has an actual world stage that is F, then the statement is false if there are worlds in which x lacks a world stage altogether even if every possible world stage of x is indeed G.)

It is worth observing that our definition of 'modal occurrent' (13) makes it impossible to use this machinery to express modal

discourse that involves counterfactualizing about entire worlds. One might want to say, for example, that this world of ours could have been better than it is, or that it might have been something else than a world, or that something which is not a world (e.g., Switzerland) might have been a world. These are not statements that can be expressed in L^* unless we relax our definitions so as to allow for modal occurents which include worlds among their parts. As a matter of fact, there would be nothing wrong with such a revision (except that 'M' would then have to be assumed as a primitive predicate). However, the definition given above reflects a limitation that is common to all familiar ways of expressing the semantics of modal discourse, so I will stick to it for comparative convenience.

Counterparts

At this point it is natural to establish a link between all the things that are world stages of the same occurrent. And the natural way of doing so is to think of such things as *counterparts* of one another, so that our translation scheme would effectively amount to a scheme for representing counterfactual statements about world stages in terms of factual statements about their counterparts. For example, on this understanding a statement such as

(23) Pavarotti might have been a ballerina

would be true just in case there is at least one counterpart of the actual world stage of Pavarotti which is a ballerina. More generally, let us define:

(34) $Ax =_{df} PPx\alpha$
 x is actual iff x is part of the actual world.

(35) $Ixy =_{df} PPxy \wedge Wy$
 x is in y iff y is a world of which x is a proper part.

- (36) $Cxy =_{df} \exists z \exists u \exists v (M_z \wedge WS_{xzu} \wedge WS_{yzv})$
 x is a counterpart of y iff x and y are world stages of a common occurrent z . (Here the fact that we are not assuming the unrestricted fusion axiom is relevant, for otherwise anything x in any world u would count as a counterpart of anything y in any world v for the simple reason that x and y would count as world stages of their sum $x+y$. Without that axiom, however, we are free to think of the occurrents that exist as being counterpart-interrelated, in fact maximally counterpart-interrelated⁸.)

Together with 'W', these three predicates form the primitive vocabulary of Lewis's counterpart theory⁹. And the appropriateness of these definitions is shown by the fact that the bulk of counterpart theory can now be derived from our three axioms as long as we make one extra assumption—viz., that there are no trans-world identities:

- (37) $\forall x \forall y (Wx \wedge Wy \rightarrow \neg POxy)$
 No two worlds overlap.

Given this additional assumption, all the postulates of counterpart theory can be proved as theorems:

- (38) $\forall x \forall y (Lxy \rightarrow Wy)$
 Nothing is in anything except a world. [This follows directly from definition (35)]
- (39) $\forall x \forall y \forall z (Lxy \wedge Lxz \rightarrow y=z)$
 Nothing is in two worlds. [From axiom (37) via definitions (1), (2), (3), and (35)]
- (40) $\forall x \forall y (Cxy \rightarrow \exists z Lxz)$
 Whatever is a counterpart is in a world. [Directly from definition (36), second conjunct, via definitions (15), (16), (17), and (35)]

8 Such occurrents correspond to what Lewis (1983: 41) calls 'modal continuants' and Lewis 1986: 214 calls '*-possible individuals'.

9 See Lewis 1968.

- (41) $\forall x \forall y (Cxy \rightarrow \exists z Iyz)$
 Whatever has a counterpart is in a world. [Directly from definition (36), third conjunct, via definitions (15), (16), (17), and (35)]
- (42) $\forall x \forall x \forall z (Ixy \wedge Izy \wedge Cxz \rightarrow x=z)$
 Nothing is a counterpart of anything else in its world. [From corollary (20) via definitions (17), (35), and (36)]
- (43) $\forall x \forall y (Ixy \rightarrow Cxx)$
 Anything in a world is a counterpart of itself. [From corollary (18) via definitions (15), (35), and (36)]
- (44) $\exists x (Wx \wedge \forall y (Ixy \leftrightarrow Ay))$
 Some world contains all and only actual things. [From axiom (10) via definitions (34) and (35)]
- (45) $\exists x Ay$
 Something is actual. [From axioms (10) and (12) via definition (34)]

There is but one notable difference between this reconstruction and Lewis's original formulation, namely, our axioms have the following consequence which is not provable in Lewis's counterpart theory:

- (46) $\forall x \forall y (Cxy \leftrightarrow Cyx)$
 Counterparthood is symmetric. [From definition (36)]

To rule this out we would have to introduce *C* as a primitive relation and replace definition (36) by an axiom which preserves only the implication from left to right:

- (47) $\forall x \forall y (Cxy \rightarrow \exists z \exists u \exists v (Mz \wedge WSxzu \wedge WSyzv))$
 A thing and its counterparts are world stages of a common occurrent.

Then every postulate of counterpart theory would still be provable except for (43), which would have to be assumed as an independent axiom.

Comparisons

We have derived counterpart theory from the theory of modal occurrents, and this may be viewed as an exercise in theory formulation. However, there is a difference between the way we would now express modal discourse in counterpart theory and the way modal discourse is expressed in Lewis's original version. The difference concerns the way ordinary names are interpreted, and consequently the definition of the translation function from L to L^* . For us, as we have seen, a name such as 'Pavarotti' denotes a modal occurrent and a statement such as

(23) Pavarotti might have been a ballerina

is true iff there is some world stage of Pavarotti which is a ballerina. For Lewis 'Pavarotti' denotes only the *actual* world stage of a counterpart-interrelated modal occurrent, and (23) is true iff there is some counterpart of Pavarotti, i.e., some world stage of that occurrent, which is a ballerina. Formally, this means that for Lewis the translation function τ is based on a different recursive definition of ϕ^y , namely:

- (25') if $\phi = Rt_1 \dots t_n$, then $\phi^y = \phi$
 if $\phi = \neg\psi$, then $\phi^y = \neg\psi^y$
 if $\phi = \psi \cdot \chi$, then $\phi^y = \psi^y \cdot \chi^y$
 if $\phi = \forall x\psi$, then $\phi^y = \forall x(Ixy \rightarrow \psi^y)$
 if $\phi = \exists x\psi$, then $\phi^y = \exists x(Ixy \wedge \psi^y)$
 if $\phi = \Box\psi x_1 \dots x_n$, then $\phi^y = \forall z \forall z_1 \dots \forall z_n (Wz \wedge Iz_1 z \wedge Cz_1 x_1 \wedge \dots \wedge Iz_n z \wedge Cz_n x_n \rightarrow \psi^z z_1 \dots z_n)$
 if $\phi = \Diamond\psi x_1 \dots x_n$, then $\phi^y = \exists z \exists z_1 \dots \exists z_n (Wz \wedge Iz_1 z \wedge Cz_1 x_1 \wedge \dots \wedge Iz_n z \wedge Cz_n x_n \wedge \psi^z z_1 \dots z_n)$

With reference to our earlier examples, this yields the following alternative translations:

- (28') $\tau(Fb) = Fb$
 'b is F' means 'b is F'.
 (29') $\tau(\forall xFx) = \forall x(Ix \alpha \rightarrow Fx)$
 'Everything is F' means 'Everything actual is F'.

- (30') $\tau(\Box Fb) = \forall x \forall y (Wy \wedge Ixy \wedge Cxb \rightarrow Fx)$
 'Necessarily b is F ' means 'Every counterpart of b , in any world, is F '.
- (31') $\tau(\Diamond \exists x Fx) = \exists y (Wy \wedge \exists x (Ixy \wedge Fx))$
 'Possibly, something is F ' means 'There is a world in which something is F '.
- (32') $\tau(\exists x \Diamond Fx) = \exists x (Ix\alpha \wedge \exists y \exists z (Wz \wedge Iyz \wedge Cyx \wedge Fy))$
 'Something is possibly F ' means 'Something actual has a counterpart that is F '.
- (33') $\tau(\forall x (Fx \rightarrow \Box Gx)) = \forall x (Ix\alpha \rightarrow (Fx \rightarrow \forall y \forall z (Wz \wedge Iyz \wedge Cyx \rightarrow Gy)))$
 'Every F is necessarily G ' means 'If something actual is F , then its counterparts, in any world, are G '.

There are several differences between these translations and our earlier translations in (28)–(33). How important are they?

One difference, for example, is that the truth of the translation in (30) depends on whether b has a world stage in every world, while the translation in (30') does not depend on that—that is, it does not depend on the corresponding question of whether b has a counterpart in every world. A similar point could be made about (33) and (33'): the former, but not the latter, depends on whether x has a world stage (respectively: a counterpart) in every world. However these are differences that reflect only a minor disagreement concerning the scope of modal statements: the translation function defined in (25) takes a modal statement to range over all possible worlds whereas Lewis's translation function (25') takes a modal statement to range exclusively over those worlds in which the relevant counterparts exist. This is a minor difference because it can easily be accommodated by changing the relevant clauses in one translation function or the other. For example, we could replace the last two clauses of (25) with the following clauses and the difference would disappear¹⁰:

10 The amendment of the clause for ' $\exists$ ' is actually redundant if descriptions are understood *à la* Russell. Again, we include it here for the sake of duality.

$$\begin{aligned}
 (25'') \text{ if } \phi &= \Box \psi x_1 \dots x_n, \text{ then } \phi^y = \forall z (Wz \wedge E!x_1^z \wedge \dots \wedge E!x_n^z \\
 &\quad \rightarrow (\psi x_1 \dots x_n)^z) \\
 \text{if } \phi &= \Diamond \psi x_1 \dots x_n, \text{ then } \phi^y = \exists z (Wz \wedge E!x_1^z \wedge \dots \wedge E!x_n^z \\
 &\quad \wedge (\psi x_1 \dots x_n)^z)
 \end{aligned}$$

Another difference is that Lewis's translation function (25') exploits the possibility that an object (a world-bound individual) has more than one counterpart in some worlds, whereas our translation function (25) ignores that possibility altogether. More precisely, definition (36) allows an object to have more than one counterpart per world, but (25) as well as its revised version (25'') say that the truth-value of a modal statement about the actual world stages of certain occurrents $x_1, \dots, x_n$ depends exclusively on the properties of the possible world stages of these occurrents—and each occurrent has at most one world stage in every world. So, for example, according to both (25) and (25'') a statement such as

(23) Pavarotti might have been a ballerina

is true iff Pavarotti's world stages comprise a ballerina. These world stages are counterparts of Pavarotti^α (i.e., of that actual thing that Lewis calls 'Pavarotti'). They are those counterparts of Pavarotti^α which fall under the same name: 'Pavarotti'. But Pavarotti^α may have other counterparts besides those. He will have as many counterparts as there are modal occurrents that share the same actual world stage as the occurrent Pavarotti. And those are not taken into account by our clauses for expressing modal statements in terms of modal occurrents (whereas they are taken into account by Lewis's clauses).

If desired, however, this difference between the two translation functions can also be eliminated. All we have to do is to emend our clauses so that a statement such as (23) will be true iff at least one counterpart of Pavarotti^α is a ballerina, regardless of whether that counterpart is a world stage of Pavarotti or of some other modal occurrent whose actual world stage is Pavarotti^α. In general, this amounts to revising (25) or (25'') along the following lines:

$$\begin{aligned}
 (25''') \text{ if } \phi = \Box \psi x_1 \dots x_n, \text{ then } \phi^y &= \forall z \forall z_1 \dots \forall z_n (Wz \wedge \\
 &Cz_1^z x_1^\alpha \wedge \dots \wedge Cz_n^z x_n^\alpha \rightarrow (\psi z_1 \dots z_n)^z) \\
 \text{if } \phi = \Diamond \psi x_1 \dots x_n, \text{ then } \phi^y &= \exists z \exists z_1 \dots \exists z_n (Wz \wedge \\
 &Cz_1^z x_1^\alpha \wedge \dots \wedge Cz_n^z x_n^\alpha \wedge (\psi z_1 \dots z_n)^z)
 \end{aligned}$$

These conditions are now pretty close to those of (25').

So it does not look as though the choice between taking C as a primitive relation, as in Lewis's original formulation of counterpart theory, or defining it in terms of modal occurrents, as in our formulation, will make much difference at all. As Lewis would say, the two theories appear to be verbal variants of one and the same technique for expressing modal talk extensionally. New terminology is not a new theory¹¹.

However that is not true. One *important* difference between the two theories concerns the way they handle indeterminacy – the sort of indeterminacy that gets into counterfactuals about named individuals¹². And this is an important difference insofar as this sort of indeterminacy – has sometimes been claimed to cause troubles for counterpart theory.

Suppose that we have doubts about (23). Suppose, that is, that there is indeterminacy about whether or not Pavarotti might have been a ballerina. Assuming that there is no indeterminacy as to who Pavarotti is in this world, this means that there is indeterminacy as to what Pavarotti's counterparts are. What sort of indeterminacy is this? In the original formulation of counterpart theory it certainly need not be a form of ontological indeterminacy. If the counterpart relation is understood as any old relation of comparative overall similarity, then the indeterminacy is merely semantic (or pragmatic): the predicate 'C' can suffer from the vagueness of our stipulations, and such vagueness can be dealt with in semantic (or pragmatic) terms. It can be differently resolved in different contexts. This is how Lewis himself sees the matter, and he insists that vagueness is therefore not a problem. Still, C is the primitive relation of counterpart theory. So to say that C is vague (albeit only semantically or pragmatically vague)

11 This is claimed both in Lewis 1983 and in Lewis 1986.

12 The point is made informally in Bennett 1988: 63–64, though with reference to a slightly different way of understanding the theory of modal occurrents.

is to concede that the theory is vague—it is a theory built around a vague primitive predicate. And to some people this is enough to say that the theory itself—and the reductive account of modality that the theory is supposed to provide—is seriously defective.

Consider now the theory of modal occurrents. In fact, let us distinguish two cases depending on whether we go along with the initial formulation (25) (or its revision based on (25'')) or with the later formulation based on (25'''). If we go along with the latter, then the indeterminacy is not at all semantic (or pragmatic). The indeterminacy exhibited by a statement such as (23) is truly ontological. This is because C itself is construed as an ontological relation, in terms of P (which is an ontological relation). If it is indetermined whether there is a counterpart of the actual Pavarotti which is a ballerina, then it must be indetermined whether there exists a modal occurrent x such that $x^\alpha = \text{Pavarotti}^\alpha$ and x^y is a ballerina for some world y ¹³. So it is indetermined whether certain occurrents exist. So there is indeterminacy as to what there is. The result is that now counterpart theory is no longer a vague theory, but it calls for a vague ontology.

By contrast, if we go along with the initial formulation (25) (or its revision based on (25'')), taking ordinary names to denote modal occurrents, then the indeterminacy of a statement such as (23) need not be a source of worry: it does not make the theory vague and it does not call for a vague ontology. Simply, the statement is *prima facie* indeterminate because it is indeterminate which modal occurrent is picked out by the name 'Pavarotti', hence which world stages count as relevant counterparts of Pavarotti's actual stage. And this *prima facie* indeterminacy is real if some of the admissible candidates contain a world stage that is a ballerina while other candidates do not. In fact, a supervaluational semantics suggests itself naturally here. A term is indeterminate insofar as there appear to be many ways of assigning it a referent, all of them compatible with the way the term is used. Hence the truth-value of a statement containing indeterminate terms is naturally construed as a function of its truth-values under the various admissible ways of assigning a unique referent to those indeterminate terms. If the statement is true under all such assignments,

13 Let us assume that the predicate 'is a ballerina' involves no vagueness.

then we may take it to be true *simpliciter*; the unmade semantic stipulations do not matter because what the statement says is true regardless (or *super-true*, as Kit Fine has it¹⁴). This is the case of a non-modal statement such as (22), for example:

(22) Pavarotti is a tenor.

Likewise, if the statement comes out false under every admissible way of assigning a unique referent to its indeterminate terms, then we may regard it as false (or *super-false*) in spite of the indeterminacy. The negation of (22) would be a case in point. It is only when the statement is true under some assignments and false under others that there is trouble. In such cases, nothing will settle the question for us and the statement will fall into a truth-value gap. This seems to be precisely the situation exhibited by a statement such as (23).

To sum up, then, we have three different accounts of the indeterminacy that gets into counterfactuals, depending on which theory we work with. On Lewis's original formulation of counterpart theory, it is the *theory* itself that is vague. On its closest mereological reconstruction in terms of modal occurrents, it is the *ontology* that is vague. And on our initial formulation of the theory of modal occurrents, it is just the names that we use that are vague: the indeterminacy exhibited by a statement such as (23) is just a case of *semantic* vagueness on a par with many others, and can be handled like any other case¹⁵. In this sense, 'Pavarotti' is modally vague just as 'Everest' is spatially vague or 'the industrial revolution' is temporally (and spatially) vague. I think that these differences between the three accounts are important. And I think that they speak in favor of the theory of modal occurrents insofar as ontological vagueness and theoretical vagueness can be objected to on independent grounds¹⁶.

14 See Fine (1975), where a supervaluational semantics for vague predicates is developed. The notion of a supervaluation goes back to van Fraassen (1966), who used it to provide a semantics for free logic, i.e., a logic in which singular terms may lack a reference. Our situation here is perfectly dual.

15 Heller (1990) has a similar account of the indeterminacy involved in our intuitions about persistence and identity through time.

16 I summarize my views on ontological vagueness in Varzi 2001.

A coda on existence and possibility

What are the consequences of this way of putting things for the prospects of modal extensionalism, i.e., the view that one can provide an extensional reductive analysis of modal discourse? Clearly, one important consequence is that by placing the vagueness of our modal discourse in the semantics of the names we use, rather than in the underlying ontology or in the semantics of the primitive notions around which the analysis is formulated, the account is not directly affected by the phenomenon of vagueness—neither more nor less than a physicalist account of the mind-body problem, for example, is affected by the vagueness of ordinary proper names. (Is the candy that Pavarotti is presently chewing part of his body? Will it be part of Pavarotti's body only after he has swallowed it? Only after he has digested it? On closer look, our earlier assumption concerning the determinacy of the actual Pavarotti will have to be discharged too.) There is, however, also a seemingly negative consequence of our way of explaining modality in terms of modal occurrents. For the appeal to modal occurrents may be said to yield a paradox.

Briefly put, the paradox is that on the proposed reductive analysis of modality one is forced to say that all possible things exist (modal realism), but not all existing things are possible. This follows from the fact that possibility is explained in terms of what goes on *in* different worlds, hence what goes on *across* worlds is not possible. The possible individuals are the world-bound stages, so the trans-world individuals (the proper modal occurrents) are impossible. They exist, but they cannot possibly exist. And what goes for existence goes for truth. Let Pavarotti₀ be one of the various modal occurrents that can plausibly be associated with the vague name 'Pavarotti'. Then the statement

(48) Pavarotti₀ exists

is a necessary falsehood. It *is* true that there is such a thing as Pavarotti₀, but it *cannot be* true.

This puzzle arises in the theory of modal occurrents but not in the original theory of counterparts in which names designate what we have called world stages, i.e., things that bear the relation I

(in) to the things which are W (possible worlds). At least, the puzzle does not arise in the original theory as long as we do not formulate it in mereological terms. Accordingly, we would have another difference between counterpart theory and the theory of modal occurents—and a remarkable difference to say the least¹⁷.

Now, one could fix things by changing the relevant notion of existence. To exist in a world is not to *be* part of that world (i.e., to be in that world, as per definition (36)). Rather, to exist in a world is to *have* parts in that world. If we accept this, then Pavarotti₀ would exist and the paradox would dissolve. Not only, but the difference between such terms as 'the winged horse' and 'the round square' (say) would be restored too: the former, not the latter, would denote an object that has parts in some possible worlds, so the former, not the latter, would denote (perhaps vaguely) a possibly existing thing. And what goes for existence goes for truth. The proposition

(49) The-winged-horse exists

would be possibly true, but

(50) The-round-square exists

would be necessarily false¹⁸.

If this way out is accepted, then the initial paradox dissolves, and so does the relevant difference between the theory of modal occurents and the theory of counterparts. Moreover, modal realism is vindicated and modality explained away: to be possible is just to exist in some world, i.e., to have parts that exist in some world.

On the other hand, one could object that a change in the notion of existence is a change in the strength of the reductive analysis. For the reductive analysis presupposes what Lewis calls the principle of Plenitude: every way that a world could be is a way that some world is¹⁹. Roughly speaking, this means that anything can

17 As a matter of fact Lewis (1983) does interpret I as P and he does accept unrestricted composition, so the puzzle does affect his understanding of the theory; but it need not.

18 The point is made in Hudson 1999.

19 See Lewis 1986: 86–92.

coexist with anything, and anything can fail to coexist with anything. *Strictly speaking*, the no-overlap principle (37) prevents this from being true in counterpart theory, so what the principle really says is this: any number of possible individuals has duplicates which coexist in some world as well as duplicates that do not coexist—where two things are duplicates iff they have exactly the same natural properties and their parts can be mapped onto one another in such a way that corresponding parts have the same natural properties too. (Never mind the details of defining «natural property».) The principle of Plenitude is crucially important if the account is to deliver a complete analysis of the notion of possibility in extensional terms, for it is this principle that guarantees that the analysis leaves «no gaps in logical space». But now there is trouble. For there is no way that modal occurrents can have duplicates which coexist. There is no way, that is, that modal occurrents can «recombine» because such a recombination would violate the no-overlap axiom. If there is a world that contains the whole of Pavarotti₀, then that world will also include among its parts the actual world stage of Pavarotti, contrary to the no-overlap principle²⁰. So if Pavarotti₀ is a possible thing, then it cannot recombine with other occurrents. And if it can recombine then it is not a trans-world occurrent.

The only way out, it seems to me, is to give up the no-overlap principle itself—or else we must go back to a vague theory of counterparts²¹.

Department of Philosophy

Columbia University

New York, NY 10027

USA

e-mail: achille.varzi@columbia.edu

20 This point was brought to my attention by Borghini 2000: §3.2.1.

21 Thanks to Andrea Borghini, Berit Brogaard, and Chris Partridge for helpful comments on an early draft of this paper.

References

- BENNETT J. (1988). *Events and Their Names*. Oxford: Clarendon Press.
- BORGHINI A. (2000). *La teoria della possibilità di David K. Lewis*. Tesi di Laurea, University of Florence.
- FINE K. (1975). Vagueness, Truth and Logic. *Synthese* 30, 265-300.
- GALLOIS A. (1998). *Occasions of Identity: The Metaphysics of Persistence, Change and Sameness*. Oxford: Clarendon Press.
- GARSON J.W. (1984). Quantification in Modal Logic. In: D.M. Gabbay & F. Guenther (eds.), *Handbook of Philosophical Logic*. Vol. II. Dordrecht: Reidel, 249-307.
- HELLER M. (1990). *The Ontology of Physical Objects: Four-dimensional Hunks of Matter*. Cambridge: Cambridge University Press.
- HUDSON H. (1999). A True, Necessary Falsehood. *Australasian Journal of Philosophy* 77, 89-91.
- LEWIS D.K. (1968). Counterpart Theory and Quantified Modal Logic. *Journal of Philosophy* 65, 113-126; reprinted in Lewis (1983), 26-39.
- LEWIS D.K. (1983). Postscripts to «Counterpart Theory and Quantified Modal Logic». In: *Philosophical Papers*, Vol. 1. New York: Oxford University Press, 39-46.
- LEWIS D.K. (1986). *On the Plurality of Worlds*. Oxford: Blackwell.
- SIMONS P. (1987). *Parts. A Study in Ontology*. Oxford: Clarendon Press.
- VAN FRAASSEN B.C. (1966). Singular Terms, Truth-Value Gaps, and Free Logic. *Journal of Philosophy* 63, 481-495.
- VARZI A.C. (2001). Vagueness, Logic, and Ontology. *The Dialogue*, to appear.

111

111

**PUBLICATIONS DE
L'INSTITUT DE LOGIQUE ET DU
CENTRE DE RECHERCHES SEMIOLOGIQUES**

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: Études 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.
11. D. Miéville (éd.): Introduction aux logiques non classiques. Octobre 1997.
12. F. Vuissoz: La conception sémantique de la vérité. Logique et philosophie chez Alfred Tarski. Décembre 1998.
13. D. Miéville, P. Joray, F. Nef, M. Bourdeau, D. Bourquin, A. Lecomte, J. Lambek, B. Godart-Wendling: Rôle et enjeux de la notion de catégorie en logique. Actes du colloque organisé à Neuchâtel, les 16 et 17 octobre 1998. Septembre 1999.

Ces publications peuvent être obtenues auprès du Centre de Recherches Sémiologiques au prix de **Fr.s. 15.-; Fr. 20.-** dès le n° 14 (TVA comprise).

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: Raisonement et calcul. Actes du colloque, Neuchâtel, 24-25 juin 1994. Septembre 1995.
64. D. Apothéloz, U. Bähler, M. Schulz (éds), 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, J.-L. Gardies, J.-B. Grize, O. Houdé, J.-P. Bronckart: Temps, logique et langage. Actes du Symposium tenu lors du colloque international «Penser le temps», Neuchâtel, 8-10 septembre 1996. Avril 1997.
00. A. Roulet Juan: Benno Besson en mouvement. Notes sur une mise en scène de «Lapin Lapin», comédie de Coline Serreau. Numéro spécial septembre 1998.
66. C. Salavastru: Identité et altérité. Les avatars de la rhétorique contemporaine. Novembre 1998.
67. D. Miéville, P. Joray, N. Gessler, B. Godart-Wendling, A. Bottani: Essais sur le nom et la nominalisation. Novembre 2000.

Les titres précédés d'un astérisque sont épuisés.

Les publications disponibles peuvent être obtenues auprès du Centre de Recherches Sémiologiques au **Fr.s. 15.-** , dès le n° 67 **Fr.s. 20.-** (TVA comprise).