

Comment construire un modèle réaliste?

Contribution au séminaire de philosophie des sciences du 19 février 1999

Robert FRANCK

Centre de philosophie des sciences

Methodos, Centre pluridisciplinaire de méthodologie des sciences sociales

INTRODUCTION

Nous avons proposé dans ce Séminaire d'aborder la question du réalisme dans les sciences à partir du rôle des modèles dans la recherche. Tom Dedeurwaerdere, lors de notre première rencontre, il y a quinze jours, nous a présenté la conception dite sémantique des modèles. Je vais, pour ma part, vous présenter une conception inhabituelle des modèles. Inhabituelle tout au moins dans le débat contemporain en philosophie des sciences. Mais elle est proche du sens ordinaire du mot 'modèle' : au sens ordinaire, un modèle c'est ce à quoi il faut se conformer. En ce sens la conception que je vais présenter des modèles est proche aussi de l'idée classique de "loi naturelle" : une loi naturelle, c'est ce à quoi la nature doit se conformer. La nature obéit aux lois de la nature à l'instar des citoyens qui obéissent aux lois civiles.

Une telle conception des modèles scientifiques est difficilement acceptable en philosophie des sciences, de même qu'on accepte difficilement aujourd'hui en philosophie des sciences l'idée classique de 'loi'. Pourquoi? Parce que parler de 'modèle' en ce sens, ou parler de 'loi' scientifique en ce sens, c'est affirmer qu'il existe une *nécessité* dans la nature : il faudrait admettre que les phénomènes naturels DOIVENT se conformer aux 'lois', qu'ils s'y conforment *nécessairement*, et que ces 'lois' sont par conséquent *universelles*. Or la tradition empiriste issue de David Hume récuse tout fondement naturel à l'affirmation de la *nécessité* et de l'*universalité* des lois de la nature. Comme cette tradition empiriste continue de dominer la philosophie des sciences, aussi chez la plupart de ceux qui tournent le dos à l'empirisme logique, il n'est pas facile de plaider en philosophie pour une conception NORMATIVE ou plus exactement NOMOLOGIQUE des modèles scientifiques. C'est pourtant ce que je vais faire maintenant.

Pour mettre plus de chances de mon côté dans ce plaidoyer je vais recourir à un stratagème. Je vais partir de préoccupations méthodologiques, au lieu de partir de préoccupations épistémologiques, autrement dit je vais commencer par

montrer les avantages méthodologiques qu'il y a, en sciences sociales, à accepter cette conception. Les considérations épistémologiques suivront. C'est pour ses avantages méthodologiques que la conception que je vais présenter du rôle des modèles dans l'explication scientifique a été adoptée au Centre Methodos pour orienter une recherche collective menée sur le pouvoir explicatif des modèles dans les sciences sociales. Les résultats de cette recherche collective seront publiés cette année sous forme d'un volume collectif, sous le titre *The explanatory Power of Models in the Social Sciences*.

Voici le plan de mon exposé.

I. 1. Je partirai d'un exemple pour illustrer une difficulté majeure et notoire qu'on rencontre en sciences sociales, et qu'une approche nomologique des modèles pourrait résoudre.

2. Puis nous jetterons un coup d'oeil sur une démarche familière de l'ingénieur, surnommée la *reverse engineering method*, et qu'Armand de Callatay a commentée dans plusieurs de ses publications. Armand de Callatay est un des collaborateurs de Methodos pour la recherche sur les modèles. Il s'appuie lui-même sur Herbert Simon. J'illustrerai cette démarche à l'aide du modèle neuronal de Mc Culloch et Pitts. Ceci va nous permettre de discerner l'existence de DEUX types de modèles auxquels correspondent DEUX types d'explication.

3. Ensuite nous jetterons un autre coup d'oeil sur l'analyse cartésienne qui préfigure merveilleusement la démarche de l'ingénieur.

4. Cela fait nous nous examinerons le désaccord de Spinoza au sujet de l'analyse cartésienne. Ce désaccord entre Spinoza et Descartes nous fournira l'occasion de découvrir comment on peut combiner les deux types de modèles et d'explications évoqués il y a un instant. C'est la différenciation entre deux sortes de modèles et d'explications, et leur combinaison, qui constitue le noyau de la conception NOMOLOGIQUE des modèles que je me propose de vous présenter.

5. Après cela, nous serons en mesure d'apprécier les avantages méthodologiques qu'il y a, en sciences sociales, à adopter cette conception.

Arrivés à ce point j'aurai terminé la partie méthodologique de mon exposé.

II. La deuxième partie sera épistémologique.

Je vais comparer la conception des modèles que je vous aurai présentée avec la conception sémantique des modèles de Frederick Suppe. Je ferai cette comparaison du point de vue de la question du réalisme.

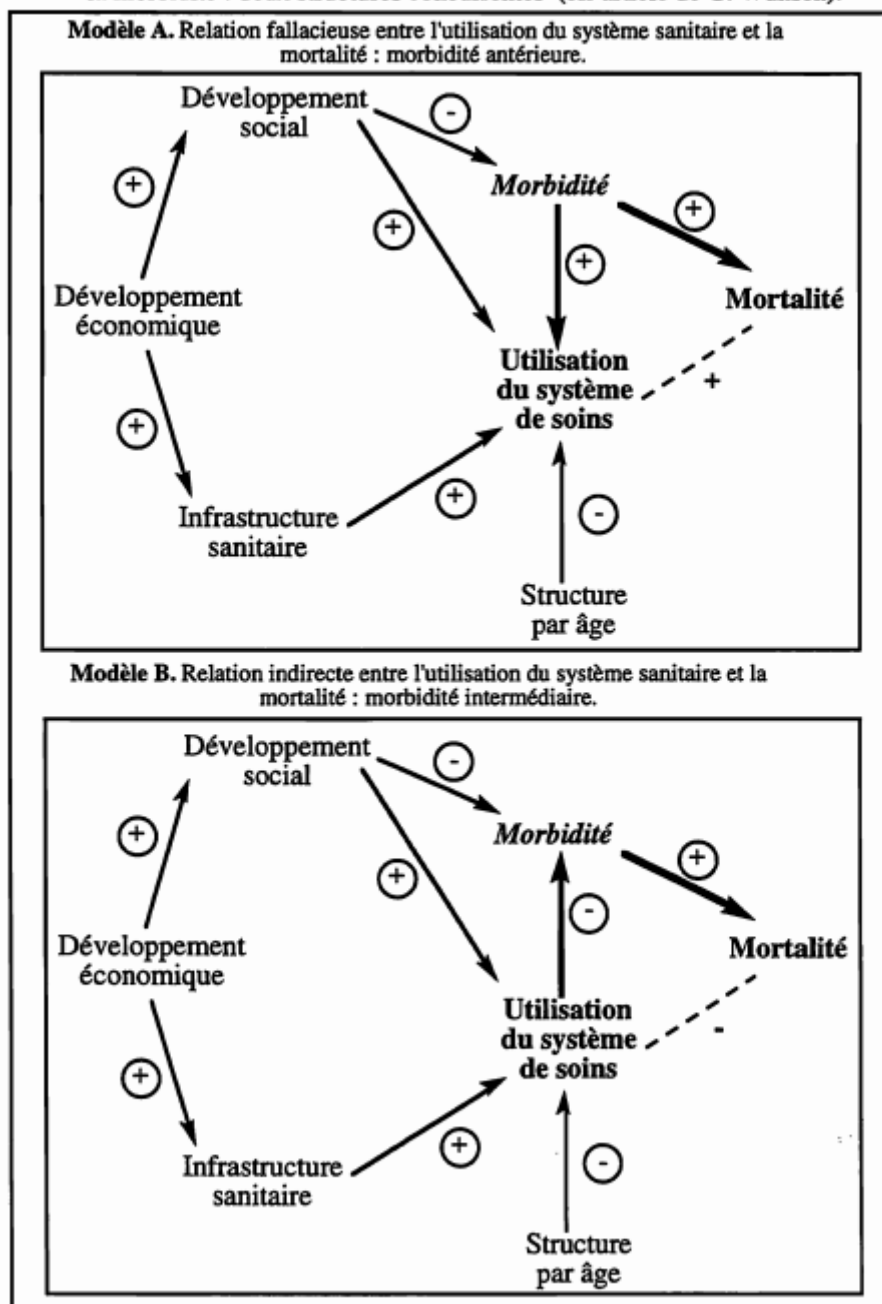
I. APPROCHE METHODOLOGIQUE

1. Une difficulté notoire en sciences sociales

Le modèle de O.Lopez Rios et al., 1992

"Le système de soins est-il une cause possible des différences récentes de mortalité constatées entre les provinces espagnoles?"

Figure 6. Révision du modèle causal de mortalité régionale par introduction de la morbidité : deux structures concurrentes (cfr article de G. Wunsch).



Les concepts qui figurent dans ce modèle, ce sont les variables dites théoriques entre lesquelles l'analyse statistique cherche à repérer des relations de covariation ou d'association; il s'agit ici du *développement économique*, du *développement social*, de l'*infrastructure sanitaire*, et de l'*utilisation du système de soins* qui, exerçant une action causale les uns sur les autres, déterminent -par hypothèse- le taux de mortalité. La théorie se présente, une fois les concepts traduits par des indicateurs mesurables, sous la forme d'un *modèle causal* qu'on va pouvoir tester. Par exemple le concept de développement social est traduit

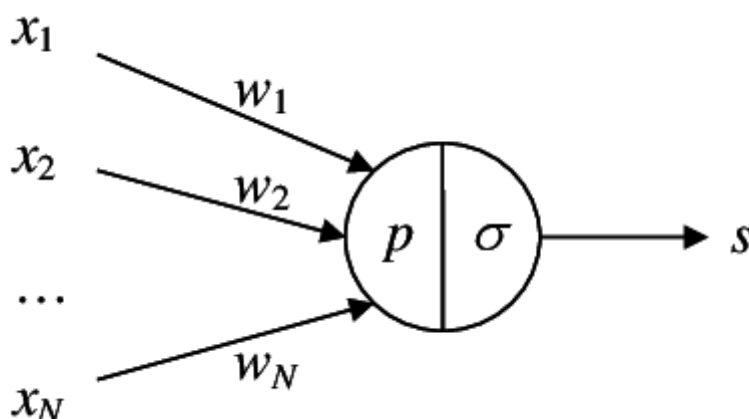
dans cette recherche par six indicateurs, à savoir le degré d'instruction des hommes et celui des femmes, l'activité féminine, le budget des ménages, la concentration urbaine, et le confort du logement. Le test consiste à évaluer l'*ajustement* du modèle aux données statistiques recueillies.

Mais comment guider le choix des *concepts* ou *variables théoriques*, comment définir leur contenu qui peut varier avec le contexte, et comment sélectionner les indicateurs qui représenteront le contenu qu'on attribue à ces concepts? Il n'existe pas jusqu'à présent de réponse satisfaisante à ces questions.

2. La démarche de l'ingénieur

Les modèles cognitifs en Intelligence Artificielle (AI), pour être fonctionnellement équivalents au cerveau, ne doivent pas copier le cerveau. Cette mise en garde est faite par Armand de Callatay dans son remarquable et impressionnant ouvrage *Natural and Artificial Intelligence* (Elsevier Science Publishers 1992, 1ère éd.1986). D'une façon générale, lorsque l'ingénieur conçoit une machine destinée à être fonctionnellement équivalente à un système réel, il ne doit pas chercher à copier ce système réel. Un pacemaker n'offre pas de ressemblance physique avec la partie du système cardiaque qu'il est appelé à remplacer; pourtant il en assure les fonctions de façon équivalente. de Callatay veut attirer l'attention sur la fécondité d'une méthode familière aux ingénieurs, que l'on nomme *reverse engineering method*, et qui consiste à reconstruire l'architecture fonctionnelle d'un objet sans connaître l'architecture matérielle de cet objet. Par exemple, il est possible de reconstruire l'architecture fonctionnelle du cerveau sans ouvrir le cerveau, ou sans faire de neurophysiologie du cerveau. Le cerveau est alors assimilé à une boîte noire. Comment s'y prend-on pour découvrir l'architecture fonctionnelle d'un système réel sans regarder dedans, sans regarder comment il fonctionne? On part de l'observation des propriétés du système, autrement dit on part de l'observation des effets que produit le système en interaction avec son environnement. Et de ces propriétés on induit les fonctions sans lesquelles ces propriétés seraient inconcevables, celles qui doivent nécessairement être présentes dans ce système pour que soient possibles les propriétés observées. Ces fonctions, ce sont des concepts désignant des opérations; l'ingénieur peut ensuite inventer des dispositifs ou des machines capables d'effectuer ces opérations. Les machines elles-mêmes ne nous apprennent rien sur le système réel, par exemple sur le cerveau; mais lorsqu'elles produisent les effets escomptés elles nous confortent dans la pertinence des fonctions conceptuelles que nous avons cru devoir inférer de l'observation des propriétés du système. de Callatay écrit : "l'ingénieur comprend un système existant lorsqu'il peut créer une machine (virtuelle) qui est *fonctionnellement équivalente* à ce système." Que comprend-il alors du système? Son architecture fonctionnelle. Et cette architecture fonctionnelle est, non les mécanismes réels du système, mais une architecture de concepts.

Voici un exemple simple d'architecture fonctionnelle : le modèle de Mc Culloch et Pitts (1943), l'ancêtre des modèles neuronaux. Il représente un neurone.



Mais que représente-t-il du neurone? Ni le *soma*, ni l'*axone*, ni les *dendrites*, ni le noyau des gènes, ni la membrane, ni la forme du neurone, bref rien de ce qu'on peut voir d'un neurone. Que représente-t-il alors? Les auteurs du modèle sont partis de l'observation des propriétés du neurone, et ils ont représenté - non les propriétés observées - mais les fonctions sans lesquelles le neurone ne pourrait pas avoir les propriétés qu'il a. Quelles fonctions?

Il faut un dispositif (X) qui assure la réception des stimulations (les entrées), un autre dispositif (W) qui pondère les stimulations (les coefficients synaptiques W), un troisième dispositif (p) qui calcule la somme des entrées pondérées, un quatrième dispositif (σ) qui fixe le seuil de stimulation en deçà duquel il ne se produit pas de transmission et au-delà duquel il y a transmission ("all- or none") et un cinquième dispositif (S) qui assure la sortie des stimulations.

Mais comme vous le voyez, ce ne sont pas les dispositifs ou mécanismes capables d'assurer les fonctions qui sont représentés dans ce modèle, mais les fonctions elles-mêmes; autrement dit, rien d'observable, rien d'autre que des concepts abstraits. Ces concepts nous apprennent ce sans quoi les propriétés du neurone seraient inexplicables; ils déterminent les conditions *théoriques* nécessaires à l'existence des propriétés observées.

Pour bien souligner la spécificité d'un tel type de modèle, comparons-le au modèle démographique que j'ai commenté tout à l'heure, de Lopez-Rios et al.

Ce modèle-là, au lieu de représenter les *fonctions nécessaires* aux propriétés observées du système et leur combinaison, représente la combinaison des *causes contingentes* hypothétiques d'un effet contingent du système.

Appelons *théorique* le premier type de modèle, celui qui représente les fonctions dont la combinaison est *nécessaire* à l'existence des propriétés du système. Et

appelons *empirique* le second type de modèle, celui qui représente des variables *contingentes* du système et les relations *contingentes* qui les lient entre elles, qualitatives ou quantitatives.

L'existence de ces deux types de modèles attire notre attention sur l'existence de deux types d'*explication* différents. Dans le premier cas l'explication consiste à conceptualiser les fonctions *nécessaires* qui doivent être réunies pour que le système puisse présenter les propriétés qu'on a observées. Dans le second cas l'explication consiste à rechercher les conditions concrètes qui, réunies, ont engendré certains effets du système.

3. L'analyse cartésienne

L'analyse cartésienne préfigure merveilleusement la démarche de l'ingénieur que nous recommandons de Callatay! Rappelez-vous le fameux *Traité de l'Homme* de Descartes. Le modèle qu'il nous propose du corps humain est celui d'une combinaison de machines. L'homme-machine de Descartes a suscité surprise ou indignation et fait couler beaucoup d'encre, comme si l'on refusait de prêter attention à la mise en garde faite par Descartes lui-même dans les premières lignes de son *Traité*. Descartes prend le soin d'écrire que ces machines ne sont PAS des descriptions de parties du corps humain comme celles qu'un anatomiste peut en livrer; ces machines, précise Descartes, représentent les fonctions que remplissent ces parties du corps humain.

Comment Descartes procède-t-il? Quel est le secret de sa méthode analytique qui a eu un tel retentissement? Le premier pas de l'analyse cartésienne (*Règles pour la direction de l'esprit, Oeuvres et Lettres*, Gallimard, Paris, 1952, pp.87-88) est d'effectuer des expériences au sujet de la chose qu'on veut expliquer car c'est seulement par l'expérience qu'on peut arriver à connaître les propriétés d'une chose c'est-à-dire les effets que la chose peut produire sur son environnement. Quand le chercheur a soigneusement collecté ces expériences il "déduit" des effets observés "le mélange de natures simples nécessaire pour produire tous les effets qu'il a constatés". Que sont ces "natures simples"? Ce sont "les idées claires et distinctes" par lesquelles on peut comprendre les choses et expliquer leurs propriétés. Voici des exemples de "natures simples" ou d'"idées claires et distinctes" que donne Descartes : la figure, l'étendue, le mouvement, deux choses sont égales si chacune est égale à une troisième. Donc lorsqu'on veut expliquer quelque chose, par exemple le corps humain, on s'occupera de rechercher les combinaisons d'"idées claires" sans lesquelles il ne serait pas possible d'expliquer ses propriétés. Comme vous le voyez la méthode préconisée par Descartes, à savoir inférer de l'observation des propriétés d'un système la combinaison d'"idées claires" qui peut les expliquer, s'applique parfaitement aux modèles d'intelligence artificielle de l'ingénieur; l'analyse cartésienne procède à la façon de la méthode du *reverse engineering*.

4. Le désaccord de Spinoza au sujet de l'analyse cartésienne

Mais pouvons-nous nous satisfaire de cette méthode? Les modèles d'intelligence artificielle, par exemple, peuvent nous faire découvrir l'architecture fonctionnelle du cerveau, mais ils ne peuvent pas nous faire découvrir comment le cerveau fonctionne concrètement. C'est un vieux débat. Spinoza faisait remarquer que si l'analyse se contente de s'appuyer sur l'observation des propriétés de l'objet, notre connaissance de l'objet se restreint à ce que nous pouvons inférer de ces propriétés et une telle connaissance est condamnée à rester abstraite, une telle méthode ne pourra jamais nous faire accéder à la connaissance de l'objet concret. Que proposait Spinoza? Il proposait de procéder à la manière dont le préconisait Descartes mais de ne pas s'en tenir là; il faut en outre observer la façon dont fonctionne concrètement le système en prenant pour guide l'architecture fonctionnelle que nous fait découvrir l'analyse cartésienne. Revenons à l'exemple de l'intelligence artificielle. Pour la méthode du *reverse engineering*, comme pour l'analyse cartésienne, le cerveau est une boîte noire. Le neurophysiologiste souscrit à Spinoza : il ouvre la boîte et observe ce qui s'y passe; mais s'il ne cherchait pas en même temps à reconstituer l'architecture fonctionnelle à partir des effets observables, il ne comprendrait pas ce qu'il voit dans la boîte et il multiplierait toutes sortes d'hypothèses causales, comme on le fait à perdre haleine dans les sciences sociales.

Bref ce que propose Spinoza, c'est de combiner les deux démarches : celle qui consiste à inférer l'architecture fonctionnelle du système à partir des propriétés observées, et celle qui consiste à décrire les processus causaux à l'oeuvre quand le système fonctionne. En quoi consiste leur combinaison? L'architecture fonctionnelle, d'une part, guide le repérage des processus causaux qui sont à l'origine des effets observés, et la connaissance des processus causaux, d'autre part, permet d'affiner la reconstruction de l'architecture fonctionnelle et de la "remplir" de ses mécanismes concrets.

Que voyons-nous? Il ne s'agit de rien d'autre que de combiner les deux types de modèles, et les deux types d'explication, que nous avons discernés tout à l'heure!

Revenons au modèle causal de Lopez-Rios et alii. Je serais tenté de l'interpréter comme un mixte des deux types de modèle et d'explication. Et je pense que souvent les modèles causaux en sciences sociales présentent de tels mixtes. Mais lorsqu'on ne distingue pas clairement les deux types d'explication, l'amalgame n'est pas fécond. Il faudrait demander jusqu'à quel point les concepts généraux, dans ce modèle, ont l'ambition de refléter ce que nous avons appelé des fonctions, ce sans quoi les variations de mortalité seraient inconcevables, et jusqu'à quel point les indicateurs retenus peuvent être

interprétés comme des mécanismes sociaux assurant les fonctions exprimées par les concepts généraux. C'est un travail délicat; je l'évoque ici seulement pour indiquer dans quel sens je pense qu'il faudrait aller pour combiner les deux types d'explications, et pour améliorer la qualité tant empirique que théorique de l'explication en sciences sociales.

5. Trois avantages méthodologiques majeurs de la conception proposée

- 1) Elle propose un critère de sélection des composantes du modèle (tant des variables empiriques que des concepts théoriques)
- 2) Elle répond à la question : quand un modèle est-il explicatif?
- 3) Elle montre comment on peut articuler la recherche empirique et la recherche théorique

II. APPROCHE EPISTEMOLOGIQUE

J'aborde maintenant la deuxième partie de mon exposé. Je vais comparer la conception NOMOLOGIQUE des modèles que je viens de vous présenter avec la conception SEMANTIQUE des modèles et des théories défendue par Frederick Suppe. Je ferai cette comparaison dans le but d'éclairer la question du réalisme dans les sciences.

A

Frederick Suppe affirme qu'une théorie scientifique ne cherche pas à rendre compte des phénomènes dans toute leur complexité; elle ne retient des phénomènes que les comportements qu'elle peut caractériser au moyen de quelques paramètres obtenus par abstraction. Par exemple s'il s'agit d'étudier la chute des corps physiques, la mécanique classique ne retient du mouvement des corps que ce qui dépend de la masse, de la vitesse, de la distance etc. Mais l'abstraction ne s'arrête pas là : la théorie suppose que la chute des corps se produit dans des conditions idéalisées, par exemple en l'absence de frottement. Autre exemple : les théories du comportement humain ou animal n'étudient pas les comportements dans toute leur complexité ou sous tous leurs aspects, par exemple les théories behavioristes choisissent de ne retenir de ces comportements que ce qui serait fonction de certaines relations Stimulus - Réponse, là aussi dans des conditions idéalisées, par exemple on suppose que

tous les individus sont semblables et ne sont pas affectés par les diverses situations dans lesquelles ils sont placés. (*The Semantic Conception of Theories and Scientific Realism*, University of Illinois Press 1989, pp.65-66) Ce qui résulte de cette abstraction et de cette idéalisation, quelle que soit la discipline scientifique concernée, Suppe l'appelle un *système physique*, ou encore un *mécanisme*. C'est ce *système physique* qui est l'objet proprement dit de la théorie scientifique, selon Suppe, et non les phénomènes dans toute leur complexité; c'est ce *système physique* ou ce mécanisme que la théorie scientifique tente de décrire, c'est de lui qu'elle tente de prédire et d'expliquer les comportements en situation idéalisée.

Suppe précise que la théorie ne se propose pas de décrire comment un *système physique* se comporte *effectivement* ("actually") mais comment il *doit* ("ought") se comporter dans les conditions idéalisées qui ont été définies. Comme le remarque Suppe, ces conditions idéalisées ne sont jamais créées que de façon approximative au moyen de l'expérimentation. La théorie scientifique s'occupe de savoir non comment le monde est réellement mais *comment le monde doit être* ("how the world ought to be"), à l'instar, précise Suppe, de ce qu'on trouve dans les délibérations morales, et tant les assertions théoriques en sciences que les assertions morales sont de nature *contrefactuelle*! (op.cit.p.280)

Selon la conception sémantique une théorie scientifique est avant tout une *structure*. Comment concevoir cette structure? Frederick Suppe répond: c'est un *système relationnel*. Il s'agit du système relationnel des paramètres qui détermine tous les possibles comportements d'un système physique dans des conditions idéalisées. Les lois de la mécanique classique, par exemple, spécifient les formes (*patterns*) "admissibles" des changements d'états d'un système physique. Bref, la structure théorique "modèle" le comportement dynamique des phénomènes ("that structure models the dynamic behavior of the phenomena" op.cit.p.4).

Les trois traits de la conception sémantique de Suppe que je viens d'évoquer peuvent se résumer comme ceci :

- 1° il introduit la notion de *système physique*, distincte de la *théorie*
- 2° il affirme que les assertions théoriques sont contrefactuelles, en ce sens qu'elles expriment comment les phénomènes *doivent* se comporter en situation idéalisée, et non comment ils se comportent en réalité
- 3° une théorie scientifique est pour l'essentiel une *structure*, et cette structure est *normative*.

B

A partir de ces trois traits, il ne sera pas difficile de montrer en quoi la conception de Frederick Suppe conforte la conception nomologique des modèles que je vous ai présentée, et aussi de noter les questions auxquelles elle

ne répond pas et auxquelles la conception que je défends permet de répondre; je dirai ensuite ce qu'il en résulte pour la question du réalisme.

Suppe conforte la conception que j'ai présentée lorsqu'il explique que la science extrait des phénomènes, par abstraction, des paramètres qui permettent de construire un *système physique*. Les paramètres dont il parle sont par exemple la masse, la vitesse, les distances en mécanique classique, ou les stimuli et les réponses en psychologie du comportement, et ne sont donc pas différents des concepts qui figurent dans les modèles théoriques dont nous parlions tout à l'heure, par exemple les concepts qui figurent dans un modèle neuronal, ou les "idées claires" de l'analyse cartésienne. Mais à la différence de la conception que j'ai présentée, celle de Suppe ne nous offre pas de critère permettant de guider la sélection des bons paramètres ou des bons concepts. Pourquoi la masse mais non la couleur, pourquoi les stimuli et non la jalousie?

Suppe conforte aussi la conception présentée tout à l'heure lorsqu'il avance que les assertions théoriques ne représentent pas ce qui se passe en réalité mais définissent les comportements possibles des phénomènes. Mais sa conception ne nous explique pas, contrairement à la conception précédente, pourquoi les phénomènes devraient nécessairement se conformer à ces assertions théoriques. Il se réfère à la logique modale et à la sémantique des mondes possibles : une assertion contrefactuelle sera vraie pour le monde réel, écrit-il (op.cit.p.280), si elle est vraie pour d'autres mondes possibles similaires. Voilà une manière d'interpréter le concept de nécessité, mais cela ne suffit pas à fonder la nécessité des assertions théoriques, il me semble.

Suppe conforte enfin la conception que j'ai présentée en distinguant trois niveaux, celui des phénomènes tels qu'ils sont dans toute leur complexité, celle des "mécanismes physiques" et celle de la structure théorique. Les "mécanismes physiques" correspondent à ce que j'ai appelé les "modèles empiriques", et la "structure" théorique correspond à ce que j'ai appelé les "modèles théoriques". Mais Suppe attribue les mêmes concepts théoriques (masses, distances..., stimuli, réponses,...) aux modèles empiriques ("systèmes physiques") et aux modèles théoriques ("structures théoriques"). La seule chose qui différencie les deux sortes de modèles, dans sa conception, c'est que des valeurs particulières sont affectées aux paramètres des systèmes physiques, qui ne représentent donc qu'une seule possible séquence d'états d'un système (op.cit.p.67) alors que les "structures théoriques" couvrent toutes les séquences d'états possibles du système. Mais le "système physique" n'est pas empirique au sens où il représenterait un système concret, tel qu'une boule de plomb roulant sur un plan incliné, ou tel qu'une machine testée par l'ingénieur, ou tel qu'un segment du cerveau analysé par le neuro-physiologiste.

Ce dernier point de comparaison nous introduit aisément à la question du réalisme scientifique. La science est-elle réaliste? Suppe se qualifie lui-même de "quasi-réaliste", et non de réaliste, parce qu'il pense qu'une théorie scientifique n'a pas pour objet les phénomènes dans toute leur complexité, mais des "systèmes physiques" abstraits dont les comportements sont idéalisés et donc jamais effectifs. Suppe affirme pourtant que la théorie peut AUCSI prédire ou expliquer les phénomènes qui ne se produisent pas dans des conditions idéalisées, lorsqu'elle est couplée à une méthodologie expérimentale appropriée (op.cit.p.67). Pourquoi dès lors ne se déclare-t-il pas réaliste, alors qu'il se sent réaliste de coeur ("the spirit of my enterprise is decidely realistic" op.cit.p.23)? Parce que l'application de la théorie aux phénomènes réels en conditions non idéalisées ne fait pas partie de sa conception de la science, c'est tout au plus un bénéfice que l'on peut tirer de la science, non une partie intégrante de la démarche scientifique. Pour qu'il en soit autrement, il faudrait que le choix des paramètres des systèmes physiques et des structures théoriques soit fondé sur les phénomènes complexes réels en conditions non idéalisées. Faute de quoi, en effet, ces phénomènes n'ont aucun rôle déterminant dans le développement de la recherche scientifique. Une telle conception peut paraître plausible, peut-être, quand on fait de la philosophie de la physique, mais n'est pas tenable, il me semble, quand on fait de la philosophie de la biologie, ou de la philosophie des sciences sociales.

La conception que j'ai présentée tout à l'heure intègre au contraire au sein de la démarche scientifique la confrontation avec les phénomènes réels tels qu'ils se comportent effectivement dans des conditions non idéalisées. Car le choix des concepts des modèles théoriques et le choix des paramètres des modèles empiriques est guidé par l'observation des effets produits par de réelles machines, par de réels cerveaux, par de réelles institutions : le critère étant que, sans les concepts retenus, ces effets ne pourraient pas se concevoir, et que, sans des dispositifs ou des mécanismes concrets assurant les opérations définies par ces mêmes concepts, les effets observés ne pourraient pas se produire.

Cela étant, la conception retenue par le Centre Methodos pour ses avantages méthodologiques majeurs, peut être qualifiée de franchement réaliste, et elle me semble sortir confortée de sa confrontation avec la conception sémantique de Frederick Suppe.
