

Améliorer les théories de la fiabilité

Robert Hudson

Volume 34, numéro 2, automne 2007

URI : <https://id.erudit.org/iderudit/016999ar>

DOI : <https://doi.org/10.7202/016999ar>

[Aller au sommaire du numéro](#)

Éditeur(s)

Société de philosophie du Québec

ISSN

0316-2923 (imprimé)

1492-1391 (numérique)

[Découvrir la revue](#)

Citer ce document

Hudson, R. (2007). Améliorer les théories de la fiabilité. *Philosophiques*, 34(2), 363–366. <https://doi.org/10.7202/016999ar>

Améliorer les théories de la fiabilité

ROBERT HUDSON

Université de Saskatchewan

Dans *Reconstructing Reason and Representation*, Murray Clarke explore la pertinence de la recherche empirique pour la philosophie de l'esprit et l'épistémologie. Mon intérêt porte, dans ce commentaire, sur la solution naturaliste qu'il cherche à fournir pour un programme particulièrement dense en faveur d'une épistémologie fiabiliste, à savoir, le problème de la généralité. Le cœur de sa solution se trouve dans son déploiement d'une « théorie de la modularité massive » qui décrit l'esprit comme étant composé d'un ensemble étendu de « processeurs computationnels spécialisés, ou modules darwiniens » (p. 6 ; toutes les références se rapportent au livre de Clarke) forgés dans l'histoire évolutive de l'*homo sapiens*. Clarke pense que sa théorie de la modularité procure une solution effective au problème de la généralité. Je rejeterai cette affirmation et offrirai une solution de rechange qui satisfait mieux un des objectifs avoués de Clarke dans ce livre — incorporer une perspective externaliste (dans les termes de Clarke : non méliorative) à une perspective internaliste (dans les termes de Clarke : méliorative) aux théories de la justification épistémique.

Clarke est, sans équivoque, un fiabiliste quant à la justification épistémique et il appuie clairement son fiabilisme sur un naturalisme, fondé sur l'évolution : pour l'essentiel, les organismes qui survivent et vivent pour transmettre leur bagage génétique aux générations futures sont celles qui entretiennent des croyances vraies sur le monde ; et la façon de garantir qu'on entretient des croyances vraies sur le monde est de posséder des mécanismes cognitifs qui aboutissent à des croyances vraies, c'est-à-dire des mécanismes cognitifs qui sont épistémologiquement fiables (voir p. 88). À cet égard, il offre une histoire naturaliste en désaccord avec les histoires présentées par Ruth Millikan (p. 65-66) et Stephen Stich (p. 77), qui affirment tous deux que la production fiable de croyances vraies n'est pas essentielle à la survivance (en bref, posséder des croyances faussement positives peut être plus adaptatif). Le débat opposant Clarke à Millikan et Stich sur ce point n'a pas à être clarifié ici. Pour les besoins de l'argumentation, nous accordons à Clarke l'avantage évolutionnaire de la fiabilité, des processus engendrant les croyances, et nous nous concentrerons sur son autre affirmation selon laquelle lorsque ces processus impliquent des raisonnements, ceux-ci doivent être implantés dans des modules darwiniens spécifiques. L'existence de tels modules a été originellement montrée, selon Clarke, par un ensemble d'expériences sur la tâche de sélection de Wason faites par Leda Cosmides et John Tooby. En supposant toujours, pour les besoins de l'argumentation, que de tels modules existent et qu'ils permettent la fiabilité de la cognition, l'affirmation de Clarke qui nous importe est que l'existence de tels

modules résout l'épineux problème de la généralité qui empoisonne toutes les théories fiabilistes de la justification.

Essentiellement, le problème de la généralité se présente comme suit : il n'est pas évident de savoir à quel niveau de description nous devrions caractériser les processus (de formation de croyances) afin de parvenir à déterminer s'ils sont fiables ou non. Si nous décrivons ces processus trop précisément, disons, avec seulement une instance (pour l'esprit, le particulier en question), alors tout processus qui culmine en une croyance vraie sera considéré comme fiable. De l'autre côté, si nous caractérisons un processus trop largement, alors nous pourrions manquer le détail des informations cruciales au processus en question — qui auraient pu, potentiellement, nous fournir une évaluation plus précise de sa fiabilité. Pour l'internalisme épistémologique, l'intractabilité du problème de la généralité annonce l'incohérence fondamentale des théories fiabilistes, et, par conséquent, il est important que les fiabilistes aient une réponse. La réponse de Clarke est la suivante :

Ma perspective ne souffre pas du problème de la généralité puisque la méthode utilisée pour déterminer la fiabilité modulaire ne repose pas sur des types qui sont pris a priori. Les types de modules sont plutôt découverts empiriquement. Nous ne créons pas les modules, nous les découvrons. La fiabilité de tels types de modules est expliquée par le fait qu'ils ont été sélectionnés précisément parce qu'ils étaient fiables dans l'EAE [environnement de l'adaptation évolutive]. Dans le EAE, ces modules sont une réussite parce qu'ils sont fiables. Par conséquent, nous ne sommes pas en position de devoir fournir une explication des types de modules qui soit trop étroite ou trop large. La sélection naturelle a fait le travail à notre place. Tout ce que nous avons à faire est de découvrir les modules en question. Il n'existe tout simplement pas de problème de la généralité pour déterminer les types de modules (p. 65).

Cette réponse aurait été excellente s'il y avait eu une voie empirique toute tracée pour isoler les modules — mais il n'y en a pas. Pour identifier les modules, il faut identifier les comportements propices à l'évolution qui sont observables chez les organismes et inférer alors l'existence d'un module qui permet, de par sa structure, ce comportement. À cet égard, Clarke cite l'exemple des castors qui claquent leur queue à la surface de l'eau pour avertir d'un danger (p. 60-61 et p. 66). Millikan note la tendance des castors à faire des claquements faussement positifs, et incrimine donc la fiabilité de cette tendance. Ici, Clarke affirme la stupidité évolutionnaire de tels claquements erronés s'ils ont été développés dans le EAE comme étant un dispositif fiable de défense par avertissement. Et il cite d'autres exemples de processus fiables, comme la tendance humaine de se recroqueviller en présence de serpents (même inoffensifs), et aussi cette étude minutieusement examinée de la réussite des sujets à la tâche de sélection de Wason, dans les cas où cette tâche introduit des conditions d'échanges sociaux tout « comme ceux [...] que nos chasseurs-cueilleurs du Pléistocène ont dû partager eux-mêmes à l'époque » (p. 98). En bref, il est convaincu d'identifier les processus fiables engendrant les croyances (ou, plutôt, il est convaincu

d'hypostasier de tels processus dans notre passé évolutionnaire) et d'affirmer alors l'existence des modules sous-jacents à ces processus. De cette façon, les modules sont présumés empiriquement discernables, dépassant ainsi le bournier analytique du problème de la généralité.

Cependant, je voudrais soutenir que cela revient à présumer une solution au problème de la généralité, non pas à y répondre. Parce que le problème précis qui nous intéresse est la mesure par laquelle nous pouvons déterminer quand un processus engendrant des croyances est fiable. Le problème de la généralité suggère que cette évaluation de la fiabilité doit attendre une détermination antérieure du niveau approprié de description pour ce processus — et rien de ce qu'a dit Clarke n'exclut ce réquisit. Il est vrai que Clarke s'intéresse aux avantages évolutifs d'un processus pour déterminer sa fiabilité ; mais il existe aussi un problème de la généralité analogue avec les avantages évolutifs — la question de savoir si un processus (engendrant des croyances) est évolutivement avantageux dépend elle-même de la façon dont il est décrit. Par exemple, quant aux claquements des castors déjà mentionnés, la question de savoir s'ils ont un avantage évolutif dépend de la façon dont nous décrivons la situation dans laquelle ils s'effectuent. Donc, pour formuler brièvement l'objection, les modules empiriquement isolés (de la façon suggérée par Clarke) ne peuvent résoudre le problème de la généralité puisque des tels isollements supposent que nous sommes capables d'identifier les processus engendrant des croyances qui sont fiables (ce qui signifie évolutivement avantageux), et une telle identification antérieure est problématique précisément en raison du problème de la généralité.

Comme solution de rechange, laissez-moi explorer comment il est possible de résoudre le problème de la généralité d'une manière cohérente avec un des buts avoués par Clarke dans son livre, qui consiste à montrer comment une perspective externaliste (non méliorative) de la justification épistémique peut et doit informer la perspective internaliste (méliorative) (voir Clarke, p. 101-114). Essentiellement, en déterminant le niveau de description approprié pour un processus engendrant des croyances, je soutiens qu'il faut solliciter la description de l'agent qui croit et qui décrira ce processus d'une certaine façon. Alors, sa croyance en tant qu'engendrée par certains processus — tels que décrits — est-elle justifiée ? Seulement si, en premier lieu, sa description est précise : il ne peut être *méliorativement* (de façon internaliste) justifié d'accepter une croyance s'il l'énonce de façon confuse. De plus, le processus engendrant les croyances telles que décrites doit être fiable en engendrant les croyances vraies. Par exemple, il doit être *suffisamment* fiable, ce qui peut ne pas être le cas si la description de l'agent est trop large. Donc, dans l'ensemble, la croyance d'une personne ne sera pas *épistémiquement* justifiée si elle décrit erronément la façon dont elle arrive (du point de vue externaliste) à sa croyance, ou si elle décrit précisément le processus par lequel elle arrive à sa croyance (peu importe le niveau de description) mais que le processus résultant, tel que décrit, n'engendre pas, de manière fiable, de croyances vraies. Dans le cas où les deux conditions suggérées sont

satisfaites, je soumets que nous avons une forme de justification qui combine tout à la fois la perspective méliorative (puisque la description même de l'agent joue un rôle central) et la perspective non méliorative (puisque la fiabilité externe est requise), une forme qui résout tout aussi bien le problème de la généralité. Elle est en mesure de résoudre le problème de la généralité puisque le niveau de description est établi de façon déterminée par la manière dont l'agent décrit le processus par lequel il est arrivé à sa croyance.

Pour donner un bref exemple afin d'illustrer mon propos, supposez un agent affirmant qu'il sait qu'un certain objet observé est un chêne. Dans le fiabilisme, nous demandons si le processus par lequel il est arrivé à sa croyance est fiable. Mais comment devrions-nous décrire ce processus ? La description est déterminée par la façon dont l'agent lui-même décrit le processus par lequel il est arrivé (du point de vue externaliste) à sa croyance. S'il affirme être arrivé à sa croyance qu'un objet observé est un chêne en donnant un certain crédit à des conseillers extra-terrestres, alors sa description erronée du processus d'acquisition des croyances exclut la possibilité qu'il soit justifié de croire que c'est un chêne. Inversement, s'il décrit le processus comme « une observation faite lors d'un après-midi clair, dans un état psychologique normal », si le processus par lequel il est arrivé à sa croyance peut être décrit (précisément) de la même façon, et si nous sommes assurés que cela est vraiment une méthode fiable à l'acquisition de croyances vraies à propos des chênes, alors sa croyance est épistémiquement justifiée. Je soutiens que ce type de perspective combine effectivement tout à la fois les aspects mélioratifs et non mélioratifs de la justification, une tâche que Clarke approuve. De plus, nous réussissons à faire cela tout en évitant le problème de la généralité, et tout en évitant également (et en ce qu'elle considère le problème de la généralité comme improductif) un détour controversé par les contributions psychologiques de la modularité.