

When Philosophers Encounter Artificial Intelligence

HOW IS IT POSSIBLE for a physical thing—a person, an animal, a robot—to extract knowledge of the world from perception and then exploit that knowledge in the guidance of successful action? That is a question with which philosophers have grappled for generations, but it could also be taken to be one of the defining questions of artificial intelligence. AI is, in large measure, philosophy. It is often directly concerned with instantly recognizable philosophical questions: What is mind? What is meaning? What is reasoning and rationality? What are the necessary conditions for the recognition of objects in perception? How are decisions made and justified?

Some philosophers have appreciated this aspect of AI, and a few have even cheerfully switched fields to pursue their philosophical quarries through thickets of LISP.* In general, however, philosophers have not welcomed this new style of philosophy with much enthusiasm. One might suppose that this is because they have seen through it. Some philosophers have indeed concluded, after cursory inspection of the field, that in spite of the breathtaking pretension of some of its publicists, artificial intelligence has nothing new to offer philosophers beyond the spectacle of ancient, well-drubbed errors replayed in a glitzy new medium. And other philosophers are so sure

Daniel C. Dennett is distinguished arts and sciences professor and professor of philosophy at Tufts University. He is director of the Center for Cognitive Studies and codirector of the Curricular Software Studio, both at Tufts.

*The programming language LISP, created by John McCarthy, is the lingua franca of AI.

this must be so that they haven't bothered conducting the cursory inspection. They are sure the field is dismissable on "general principles."

Philosophers have been dreaming about AI for centuries. Hobbes and Leibniz, in very different ways, tried to explore the implications of the idea of breaking down the mind into small, ultimately mechanical, operations. Descartes even anticipated the Turing test (Alan Turing's much-discussed proposal of an audition of sorts for computers, in which the computer's task is to convince the judges that they are conversing with a human being¹) and did not hesitate to issue a confident prediction of its inevitable result:

It is indeed conceivable that a machine could be made so that it would utter words, and even words appropriate to the presence of physical acts or objects which cause some change in its organs; as, for example, if it was touched in some spot that it would ask what you wanted to say to it; if in another, that it would cry that it was hurt, and so on for similar things. But it could never modify its phrases to reply to the sense of whatever was said in its presence, as even the most stupid men can do.²

The appreciation Descartes had for the powers of mechanism was colored by his acquaintance with the marvelous clockwork automata of his day. He could see very clearly and distinctly, no doubt, the limitations of that technology. Not even a thousand tiny gears—not even ten thousand—would permit an automaton to respond gracefully and rationally! Perhaps Hobbes or Leibniz would have been less confident of this point, but surely none of them would have bothered wondering about the *a priori* limits on a million tiny gears spinning millions of times a second. That was simply not a thinkable thought for them. It was unthinkable then, not in the familiar philosophical sense of appearing self-contradictory ("repugnant to reason") or entirely outside their conceptual scheme (like the concept of a neutrino), but in the more workaday, yet equally limiting, sense of being an idea they would have had no way to take seriously. When philosophers set out to scout large conceptual domains, they are as inhibited in the paths they take by their sense of silliness as by their insight into logical necessity. And there is something about AI that many philosophers find off-putting—if not repugnant to reason, then repugnant to their aesthetic sense.

This clash of vision was memorably displayed in a historic debate at Tufts University in March of 1978, staged, appropriately, by the Society for Philosophy and Psychology. Nominally a panel discussion on the foundations and prospects of artificial intelligence, it turned into a tag-team rhetorical wrestling match between four heavyweight ideologues: Noam Chomsky and Jerry Fodor attacking AI, and Roger Schank and Terry Winograd defending it. Schank was working at the time on programs for natural-language comprehension, and the critics focused on his scheme for representing (in a computer) the higgledy-piggledy collection of trivia we all know and somehow rely on when deciphering ordinary speech acts, allusive and truncated as they are. Chomsky and Fodor heaped scorn on this enterprise, but the grounds of their attack gradually shifted in the course of the match. It began as a straightforward, “first principles” condemnation of conceptual error—Schank was on one fool’s errand or another—but it ended with a striking concession from Chomsky: it just might turn out, as Schank thought, that the human capacity to comprehend conversation (and more generally, to think) was to be explained in terms of the interaction of hundreds or thousands of jerry-built gizmos—pseudorepresentations, one might call them—but that would be a shame, for then psychology would prove in the end not to be “interesting.” There were only two interesting possibilities, in Chomsky’s mind: psychology could turn out to be “like physics”—its regularities explainable as the consequences of a few deep, elegant, inexorable laws—or psychology could turn out to be utterly lacking in laws—in which case the only way to study or expound psychology would be the novelist’s way (and he much preferred Jane Austen to Roger Schank, if that were the enterprise).

A vigorous debate ensued among the panelists and audience, capped by an observation from Chomsky’s colleague at the Massachusetts Institute of Technology, Marvin Minsky, one of the founding fathers of AI and founder of MIT’s Artificial Intelligence Laboratory: “I think only a humanities professor at MIT could be so oblivious to the third interesting possibility: psychology could turn out to be like engineering.”

Minsky had put his finger on it. There is something about the prospect of an engineering approach to the mind that is deeply repugnant to a certain sort of humanist, and it has little or nothing to do with a distaste for materialism or science. Witness Chomsky’s

physics worship, an attitude he shares with many philosophers. The days of Berkeleyan idealism and Cartesian dualism are over (if one can judge from the current materialistic consensus among philosophers and scientists), but in their place there is a widespread acceptance of what we might call Chomsky's fork: there are only two appealing ("interesting") alternatives.

On the one hand, there is the dignity and purity of the Crystalline Mind. Recall Aristotle's prejudice against extending earthly physics to the heavens, which ought, he thought, to be bound by a higher and purer order. This was his one pernicious legacy, but now that the heavens have been stormed, we appreciate the beauty of universal physics and can hope that the mind will be among its chosen "natural kinds," not a mere gerrymandering of bits and pieces.

On the other hand, there is the dignity of ultimate mystery, the Inexplicable Mind. If our minds can't be fundamental, then let them be anomalous. A very influential view among philosophers in recent years has been Donald Davidson's "anomalous monism," the view that while the mind is the brain, there are no lawlike regularities aligning mental facts with physical facts.³ John Searle, Davidson's colleague at Berkeley, has made a different sort of mystery of the mind: the brain, thanks to some unspecified feature of its biochemistry, has some terribly important—but unspecified—"bottom-up causal powers" that are entirely distinct from the mere "control powers" studied in AI.

One feature shared by these otherwise drastically different forms of mind-body materialism is a resistance to Minsky's *tertium quid*: in between the mind as crystal and the mind as chaos lies the mind as gadget, an object that one should not expect to be governed by "deep" mathematical laws, but nevertheless a *designed* object, analyzable in functional terms: ends and means, costs and benefits, elegant solutions on the one hand, and on the other, shortcuts, jury rigs, and cheap *ad hoc* fixes.

This vision of the mind is resisted by many philosophers despite its being a straightforward implication of the current view among scientists and science-minded humanists of our place in nature: we are biological entities designed by natural selection, which is a tinker, not an ideal engineer. Computer programmers call an *ad hoc* fix a "kludge" (it rhymes with *Scrooge*), and the mixture of disdain and begrudged admiration reserved for kludges parallels the biologists'

bemusement with “the panda’s thumb” and other fascinating examples of *bricolage*, to use François Jacob’s term.⁴ The finest inadvertent spoonerism I ever heard was uttered by the linguist Barbara Partee in heated criticism of an acknowledged kludge in an AI natural-language parser: “That’s so odd hack!” Nature is full of odd hacks, many of them perversely brilliant. Although this fact is widely appreciated, its implications for the study of the mind are often repugnant to philosophers, since their traditional aprioristic methods of investigating the mind give them little power to explore phenomena that might be contrived of odd hacks. There is really only one way to study such possibilities: with the more empirical mind-set of “reverse engineering.”

The resistance is clearly manifested in Hilary Putnam’s essay in this issue of *Dædalus*, which can serve as a convenient (if not particularly florid) case of the syndrome I wish to discuss. Chomsky’s fork, the mind as crystal or as chaos, is transformed by Putnam into a pendulum swing he thinks he observes within AI itself. He claims that AI has “wobbled” over the years between looking for the Master Program and accepting the notion that “artificial intelligence is one damned thing after another.” I have not myself observed any such wobble in the field over the years, but I think I know what he is getting at. Here, then, is a different perspective on the same issue.

Among the many divisions of opinion within AI there is a faction (sometimes called the logicists) whose aspirations suggest to me that they are Putnam’s searchers for the Master Program. They were more aptly caricatured recently by a researcher in AI as searchers for “Maxwell’s equations of thought.” Several somewhat incompatible enterprises within the field can be lumped together under this rubric. Roughly, what they have in common is the idea not that there must be a Master Program but that there must be something more like a master programming language, a single, logically sound system of explicit representation for all the knowledge residing in an agent (natural or artificial). Attached to this library of represented facts (which can be treated as axioms, in effect) and operating upon it computationally will be one sort or another of “inference engine,” capable of deducing the relevant implications of the relevant axioms and eventually spewing up by this inference process the imperatives or decisions that will forthwith be implemented.

For instance, suppose perception yields the urgent new premise (couched in the master programming language) that the edge of a precipice is fast approaching; this should provoke the inference engine to call up from memory the appropriate stored facts about cliffs, gravity, acceleration, impact, damage, the paramount undesirability of such damage, and the likely effects of putting on the brakes or continuing apace. Forthwith, one hopes, the engine will deduce a theorem to the effect that halting is called for, and straightaway it will halt.

The hard part is designing a system of this sort that will actually work well in real time, even allowing for millions of operations per second in the inference engine. Everyone recognizes this problem of real-time adroitness; what sets the logicians apart is their conviction that the way to solve it is to find a truly perspicuous vocabulary and logical form for the master language. Modern logic has proven to be a powerful means of exploring and representing the stately universe of mathematics; the not unreasonable hope of the logicians is that the same systems of logic can be harnessed to capture the hectic universe of agents making their way in the protean macroscopic world. If you get the axioms and the inference system just right, they believe, the rest should be easy. The problems they encounter have to do with keeping the number of axioms down for the sake of generality (which is a must), while not requiring the system to waste time rededucing crucial intermediate-level facts every time it sees a cliff.

This idea of axiomatizing everyday reality is surely a philosophical one. Spinoza would have loved it, and many contemporary philosophers working in philosophical logic and the semantics of natural language share at least the goal of devising a rigorous logical system in which every statement, every thought, every hunch and wonder can be unequivocally expressed. The idea wasn't reinvented by AI; it was a gift from the philosophers who created modern mathematical logic: George Boole, Gottlob Frege, Alfred North Whitehead, Bertrand Russell, Alfred Tarski, and Alonzo Church. Douglas Hofstadter calls this theme in AI the Boolean dream.⁵ It has always had its adherents and critics, with many variations.

Putnam's rendering of this theme as the search for the Master Program is clear enough, but when he describes the opposite pole, he elides our two remaining prospects: the mind as gadget and the mind as chaos. As he puts it, "If AI is 'one damned thing after another,' the

number of 'damned things' the tinker may have thought of could be astronomical. The upshot is pessimistic indeed: if there is no Master Program, then we may never get very far in terms of simulating human intelligence." Here Putnam elevates a worst-case possibility (the gadget will be totally, "astronomically" *ad hoc*) as the only likely alternative to the Master Program. Why does he do this? What does he have against exploring the vast space of engineering possibilities between Crystal and Chaos? Biological wisdom, far from favoring his pessimism, holds out hope that the mix of elegance and Rube Goldberg found elsewhere in nature (in the biochemistry of reproduction, for instance) will be discernible in the mind as well.

There is, in fact, a variety of very different approaches being pursued in AI by those who hope the mind will turn out to be some sort of gadget or collection of partially integrated gadgets. All of these favor austerity, logic, and order in some aspects of their systems and yet exploit the peculiar utility of profligacy, inconsistency, and disorder in other aspects. It is not that Putnam's two themes don't exist in AI, but that by describing them as exclusive alternatives, he imposes a procrustean taxonomy on the field that makes it hard to discern the interesting issues that actually drive the field.

Most AI projects are explorations of *ways things might be done* and as such are more like thought experiments than empirical experiments. They differ from philosophical thought experiments not primarily in their content but in their methodology: they replace some—not all—of the "intuitive," "plausible," hand-waving background assumptions of philosophical thought experiments by constraints dictated by the demand that the model be made to run on the computer. These constraints of time and space and the exigencies of specification can be traded off against each other in practically limitless ways, so that new "virtual machines" or "virtual architectures" are imposed on the underlying serial architecture of the digital computer. Some choices of trade-off are better motivated, more realistic, or more plausible than others, of course, but in every case the constraints imposed serve to discipline the imagination—and hence the claims—of the thought experimenter. There is very little chance that a philosopher will be surprised (or more exactly, disappointed) by the results of his own thought experiment, but this happens all the time in AI.

A philosopher looking closely at these projects will find abundant grounds for skepticism. Many seem to be based on forlorn hopes or misbegotten enthusiasm for one architectural or information-handling feature or another, and if we extrapolate from the brief history of the field, we can be sure that most of the skepticism will be vindicated sooner or later. What makes AI an improvement on earlier philosophers' efforts at model sketching, however, is the manner in which skepticism is vindicated: by the actual failure of the system in question. Like philosophers, researchers in AI greet each new proposal with intuitive judgments about its prospects, backed up by more or less a priori arguments about why a certain feature has to be there or can't be made to work. But unlike philosophers, these researchers are not content with their arguments and intuitions; they leave themselves some room to be surprised by the results, a surprise that could only be provoked by the demonstrated, unexpected power of the actually contrived system in action.

Putnam surveys a panoply of problems facing AI: the problems of induction, of discerning relevant similarity, of learning, of modeling background knowledge. These are all widely recognized problems in AI, and the points he makes about them have all been made before by people in AI, who have then gone on to try to address the problems with various relatively concrete proposals. The devilish difficulties he sees facing traditional accounts of the process of induction, for example, are even more trenchantly catalogued by John Holland, Keith Holyoak, Richard Nisbett, and Paul Thagard in their recent book *Induction*,⁶ but their diagnosis of these ills is the preamble for sketches of AI models designed to overcome them. Models addressed to the problems of discerning similarity and mechanisms for learning can be found in abundance. The SOAR project of John Laird, Allen Newell, and Paul Rosenbloom⁷ is an estimable example. And the theme of the importance—and difficulty—of modeling background knowledge has been ubiquitous in recent years, with many suggestions for solutions under investigation. Now perhaps they are all hopeless, as Putnam is inclined to believe, but one simply cannot tell without actually building the models and testing them.

This last statement is not strictly true, of course. When an a priori refutation of an idea is sound, the doubting empirical model builder who persists despite the refutation will sooner or later have to face a chorus shouting "We told you so!" That is one of the occupational

hazards of AI. The rub is how to tell the genuine a priori proofs of impossibility from mere failures of imagination. The philosophers' traditional answer is, More a priori analysis and argument. The AI researchers' answer is, Build it and see.

Putnam offers us a striking instance of this difference in his survey of possibilities for tackling the problem of background knowledge. Like Descartes, he manages to imagine a thought-experiment fiction that is now becoming real, and like Descartes, he is prepared to dismiss it in advance. One could, Putnam says,

simply try to program into a machine all the information a sophisticated human inductive judge has (including implicit information). At the least, this would require generations of researchers to formalize the information (probably it could not be done at all, because of the sheer quantity of information involved); and it is not clear that the result would be more than a gigantic expert system. No one would find this very exciting; and such an "intelligence" would in all likelihood be dreadfully unimaginative. . . .

This almost perfectly describes Douglas Lenat's enormous CYC project.⁸ One might say that Lenat is attempting to create the proverbial walking encyclopedia: a mind-ful of commonsense knowledge in the form of a single data base containing all the facts expressed—or tacitly presupposed—in an encyclopedia! This involves handcrafting millions of representations in a single language (which must eventually be unified—no small task), from which the inference engine is expected to be able to deduce whatever it needs as it encounters novelty in its world: for instance, the fact that people in general prefer not to have their feet cut off or the fact that sunbathers are rare on Cape Cod in February.

Most of the opinion setters in AI share Putnam's jaundiced view of this project: it is not clear, as Putnam says, that the project will do anything that teaches us anything about the mind; in all likelihood, as he says, it will be dreadfully unimaginative. And many would go further and insist that its prospects are so forlorn and its cost so great that it should be abandoned in favor of more promising avenues. (The current estimate is measured in person-*centuries* of work, a figure that Putnam may not have bothered to imagine in detail.) But the project is funded, and we shall see.

What we have here is a clash of quite fundamental methodological assumptions. Philosophers are inclined to view AI projects with the

patronizing disdain one reserves for those persistent fools who keep trying to square the circle or trisect the angle with compass and straightedge: we have *proved* that it cannot be done, so drop it! But the proofs are not geometric; they are ringed with assumptions about “plausible” boundary conditions and replete with idealizations that may prove as irrelevant here as in the notorious aerodynamicists’ proofs that bumblebees cannot fly.

But still one may well inquire, echoing Putnam’s challenge, whether AI has taught philosophers anything of importance about the mind *yet*. Putnam thinks it has not and supports his view with a rhetorically curious indictment: AI has utterly failed, over a quarter century, to solve problems that philosophy has utterly failed to solve over two millennia. He is right, I guess, but I am not impressed.⁹ It is as if a philosopher were to conclude a dismissal of contemporary biology by saying that the biologists have not so much as asked the question, What is Life? Indeed, they have not; they have asked better questions that ought to dissolve or redirect the philosopher’s curiosity.

Moreover, philosophers (of all people) should appreciate that solutions to problems are not the only good gift; tough new problems are just as good! Matching Putnam’s rhetorical curiosity, I offer as AI’s best contribution to philosophy a deep, new, unsolved epistemological problem ignored by generations of philosophers: the frame problem. Plato almost saw it. In the *Theaetetus*, he briefly explored the implications of a wonderful analogy:

Socrates: Now consider whether knowledge is a thing you can possess in that way without having it about you, like a man who has caught some wild birds—pigeons or what not—and keeps them in an aviary he has made for them at home. In a sense, of course, we might say he “has” them all the time inasmuch as he possesses them, mightn’t we?

Theaetetus: Yes.

Socrates: But in another sense he “has” none of them, though he has got control of them, now that he has made them captive in an enclosure of his own; he can take and have hold of them whenever he likes by catching any bird he chooses, and let them go again; and it is open to him to do that as often as he pleases.¹⁰

Plato saw that merely possessing knowledge (like birds in an aviary) is not enough; one must be able to command what one possesses. To

perform well, one must be able to get the right bit of knowledge to fly to the edge at the right time (in real time, as the engineers say). But he underestimated the difficulty of this trick and hence underestimated the sort of theory one would have to give of the organization of knowledge in order to explain our bird-charming talents. Neither Plato nor any subsequent philosopher, so far as I can see, saw this as in itself a deep problem of epistemology, since the demands of efficiency and robustness paled into invisibility when compared with the philosophical demand for certainty, but so it has emerged in the hands of AI.¹¹

Just as important to philosophy as new problems and new solutions, however, is new raw material, and this AI has provided in abundance. It has provided a bounty of objects to think about—individual systems in all their particularity that are much more vivid and quirky than the systems I (for one) could dream up in a thought experiment. This is not a trivial harvest. Compare philosophy of mind (the analytic study of the limits, opportunities, and implications of possible theories of the mind) with the literary theory of the novel (the analytic study of the limits, opportunities, and implications of possible novels). One could in principle write excellent literary theory in the absence of novels as exemplars. Aristotle, for instance, could in principle have written a treatise on the anticipated strengths and weaknesses, powers and problems, of the various possible types of novels. Today's literary theorist is not required to examine the existing exemplars, but they are, to say the least, a useful crutch. They extend the imaginative range and the surefootedness of even the most brilliant theoretician and provide bracing checks on enthusiastic generalizations and conclusions. The minitheories, sketches, and models of AI may not be great novels, but they are the best we have to date, and just as mediocre novels are often a boon to literary theorists—they wear their deficiencies on their sleeves—so bad theories, failed models, and hopelessly confused hunches in AI are a boon to philosophers of mind. But you have to read them to get the benefit.

Perhaps the best current example of this benefit is the wave of enthusiasm for connectionist models. For years philosophers of mind have been vaguely and hopefully waving their hands in the direction of these models—utterly unable to conceive them in detail but sure in their bones that some such thing had to be possible. (My own first

book, *Content and Consciousness*, is a good example of such vague theorizing.¹²) Other philosophers have been just as sure that all such approaches were doomed (Jerry Fodor is a good example). Now, at last, we will be able to examine a host of objects in this anticipated class and find out whose hunches were correct. In principle, no doubt, it could be worked out without the crutches, but in practice, such disagreements between philosophers tend to degenerate into hardened positions defended by increasingly strained arguments, redefinitions of terms, and tendentious morals drawn from other quarters.

Putnam suggests that since AI is first and foremost a subbranch of engineering, it cannot be philosophy. He is especially insistent that we should dismiss its claim of being epistemology. I find this suggestion curious. Surely Hobbes and Leibniz and Descartes were doing philosophy, even epistemology, when they waved their hands and spoke very abstractly about the limits of mechanism. So was Kant, when he claimed to be investigating the conditions under which experience was possible. Philosophers have traditionally tried to figure out the combinatorial powers and inherent limitations of “impressions and ideas,” of “petites perceptions,” “intuitions,” and “schemata.” Researchers in AI have asked similar questions about various sorts of “data structures” and “procedural representations” and “frames” and “links” and yes, “schemata,” now rather more rigorously defined. So far as I can see, these are fundamentally the same investigations, but in AI they are conducted under additional (and generally well-motivated) constraints and with the aid of a host of more specific concepts.

Putnam sees engineering and epistemology as incompatible. I see at most a trade-off: to the extent that a speculative exploration in AI is more abstract, more idealized, less mechanistically constrained, it is “more philosophical”—but that does not mean it is thereby necessarily of more interest or value to a philosopher! On the contrary, it is probably because philosophers have been too philosophical—too abstract, idealized, and unconstrained by empirically plausible mechanistic assumptions—that they have failed for so long to make much sense of the mind. AI has not yet solved any of our ancient riddles about the mind, but it has provided us with new ways of disciplining and extending philosophical imagination that we have only begun to exploit.

ENDNOTES

- ¹Alan Turing, "Computing Machinery and Intelligence," *Mind* 59 (236) (1950): 433, reprinted in Douglas Hofstadter and Daniel Dennett, eds., *The Mind's I* (New York: Basic Books, 1981), 54–67.
- ²René Descartes, *Discourse on Method* (1637), trans. Laurence J. LaFleur, 3d ed. (New York: Bobbs-Merrill, 1960), 41–42.
- ³Donald Davidson, "Mental Events," in L. Foster and J. W. Swanson, eds., *Experience and Theory* (Amherst: University of Massachusetts Press, 1970), 79–101.
- ⁴François Jacob, "Evolution and Tinkering," *Science* 196 (1977):1161–66.
- ⁵Douglas Hofstadter, "Waking Up from the Boolean Dream, or Subcognition as Computation," chap. 26, in *Metamagical Themas* (New York: Basic Books, 1985), 631–65.
- ⁶John H. Holland, Keith J. Holyoak, Richard E. Nisbett, and Paul R. Thagard, *Induction: Processes of Inference, Learning and Discovery* (Cambridge: MIT Press, 1986).
- ⁷John E. Laird, Allen Newell and Paul S. Rosenbloom, "SOAR: An Architecture for General Intelligence," *Artificial Intelligence* 33 (September 1987):1–64.
- ⁸Douglas Lenat, Mayank Prakash, and May Shepherd, "CYC: Using Commonsense Knowledge to Overcome Brittleness and Knowledge Acquisition Bottlenecks," *AI Magazine* 6 (4) (1986):65–85.
- ⁹In "Artificial Intelligence as Philosophy and as Psychology," in *Philosophical Perspectives in Artificial Intelligence*, ed. Martin Ringle (Atlantic Highlands, N.J.: Humanities Press International, 1979), and in *Brainstorms: Philosophical Essays on Mind and Psychology* (Cambridge: MIT Press, 1978), I have argued that AI has solved what I have called Hume's Problem—the problem of breaking the threatened infinite regress of homunculi consulting (and understanding) internal representations such as Hume's impressions and ideas. I suspect Putnam would claim, with some justice, that it was computer science in general, not AI in particular, that showed philosophy the way to break this regress.
- ¹⁰*Plato's Theaetetus*, trans. Francis M. Cornford (New York: Macmillan, 1957), 197 C–D.
- ¹¹Daniel C. Dennett, "Cognitive Wheels: The Frame Problem of AI" in *Minds, Machines and Evolution: Philosophical Studies*, ed. Christopher Hookway (Cambridge: Cambridge University Press, 1985), reprinted in the new anthology *The Robot's Dilemma: The Frame Problem in Artificial Intelligence*, ed. Zenon W. Pylyshyn (Norwood, N.J.: Ablex, 1987). In this introduction to the frame problem, I explain why it is an epistemological problem and why philosophers didn't notice it.
- ¹²Daniel C. Dennett, *Content and Consciousness* (Atlantic Highlands, N.J.: Humanities Press International, 1969).