

In 1969, John McCarthy, the mathematician who coined the term "Artificial Intelligence," joined forces with another AI researcher, Patrick Hayes, and coined another term, "the frame problem." It seemed to be a devil of a problem, perhaps even a mortal blow to the hopes of AI. It looked to me like a new philosophical or epistemological problem, certainly not an electronic or computational problem. What is the frame problem? In this essay I thought I was providing interested bystanders with a useful introduction to the problem, as well as an account of its philosophical interest, but some people in AI, including McCarthy and Hayes (who ought to know), thought I was misleading the bystanders. The real frame problem was not, they said, what I was talking about. Others weren't so sure. I myself no longer have any firm convictions about which problems are which, or even about which problems are really hard after all—but something is still obstreperously resisting solution, that's for sure. Happily, there are now three follow-up volumes that pursue these disagreements through fascinating thickets of controversy: The Robot's Dilemma (Pylyshyn, 1987), Reasoning Agents in a Dynamic World (Ford and Hayes, 1991), and The Robot's Dilemma Revisited (Ford and Pylyshyn, 1996). If you read and understand those volumes you will understand just about everything anybody understands about the frame problem. Good luck.

Once upon a time there was robot, named R1 by its creators. Its only task was to fend for itself. One day its designers arranged for it to learn that its spare battery, its precious energy supply, was locked in a room with a time bomb set to go off soon. R1 located the room and the key to the door, and formulated a plan to rescue its battery. There was

a wagon in the room, and the battery was on the wagon, and R1 hypothesized that a certain action which it called PULLOUT (WAGON, ROOM) would result in the battery being removed from the room. Straightway it acted, and did succeed in getting the battery out of the room before the bomb went off. Unfortunately, however, the bomb was also on the wagon. R1 *knew* that the bomb was on the wagon in the room, but didn't realize that pulling the wagon would bring the bomb out along with the battery. Poor R1 had missed that obvious implication of its planned act.

Back to the drawing board. "The solution is obvious," said the designers. "Our next robot must be made to recognize not just the intended implications of its acts, but also the implications about their side effects, by deducing these implications from the descriptions it uses in formulating its plans." They called their next model, the robot-deducer, R1D1. They placed R1D1 in much the same predicament that R1 had succumbed to, and as it too hit upon the idea of PULLOUT (WAGON, ROOM) it began, as designed, to consider the implications of such a course of action. It had just finished deducing that pulling the wagon out of the room would not change the color of the room's walls, and was embarking on a proof of the further implication that pulling the wagon out would cause its wheels to turn more revolutions than there were wheels on the wagon—when the bomb exploded.

Back to the drawing board. "We must teach it the difference between relevant implications and irrelevant implications," said the designers, "and teach it to ignore the irrelevant ones." So they developed a method of tagging implications as either relevant or irrelevant to the project at hand, and installed the method in their next model, the robot-relevant-deducer, or R2D1 for short. When they subjected R2D1 to the test that had so unequivocally selected its ancestors for extinction, they were surprised to see it sitting, Hamlet-like, outside the room containing the ticking bomb, the native hue of its resolution sicklied o'er with the pale cast of thought, as Shakespeare (and more recently Fodor) has aptly put it. "Do something!" they yelled at it. "I am," it retorted. "I'm busy ignoring some thousands of implications I have determined to be irrelevant. Just as soon as I find an irrelevant implication, I put it on the list of those I must ignore, and . . ." the bomb went off.

All these robot suffer from the *frame problem*.¹ If there is ever to be

1. The problem is introduced by John McCarthy and Patrick Hayes in their 1969 paper. The task in which the problem arises was first formulated in McCarthy 1960. I am grateful to Bo Dahlbohm, Pat Hayes, John Haugeland, John McCarthy, Bob Moore, and Zenon

a robot with the fabled perspicacity and real-time adroitness of R2D2, robot designers must solve the frame problem. It appears at first to be at best an annoying technical embarrassment in robotics, or merely a curious puzzle for the bemusement of people working in Artificial Intelligence (AI). I think, on the contrary, that it is a new, deep epistemological problem—accessible in principle but unnoticed by generations of philosophers—brought to light by the novel methods of AI, and still far from being solved. Many people in AI have come to have a similarly high regard for the seriousness of the frame problem. As one researcher has quipped, “We have given up the goal of designing an intelligent robot, and turned to the task of designing a gun that will destroy any intelligent robot that anyone else designs!”

I will try here to present an elementary, nontechnical, philosophical introduction to the frame problem, and show why it is so interesting. I have no solution to offer, or even any original suggestions for where a solution might lie. It is hard enough, I have discovered, just to say clearly what the frame problem is—and is not. In fact, there is less than perfect agreement in usage within the AI research community. McCarthy and Hayes, who coined the term, use it to refer to a particular, narrowly conceived problem about representation that arises only for certain strategies for dealing with a broader problem about real-time planning systems. Others call this broader problem the frame problem—“the whole pudding,” as Hayes has called it (personal correspondence)—and this may not be mere terminological sloppiness. If “solutions” to the narrowly conceived problem have the effect of driving a (deeper) difficulty into some other quarter of the broad problem, we might better reserve the title for this hard-to-corner difficulty. With apologies to McCarthy and Hayes for joining those who would appropriate their term, I am going to attempt an introduction to the whole pudding, calling *it* the frame problem. I will try in due course to describe the narrower version of the problem, “the frame problem proper” if you like, and show something of its relation to the broader problem.

Pylyshyn for the many hours they have spent trying to make me understand the frame problem. It is not their fault that so much of their instruction has still not taken.

I have also benefited greatly from reading an unpublished paper, “Modeling Change—the Frame Problem,” by Lars-Erik Janlert, Institute of Information Processing, University of Umea, Sweden. It is to be hoped that a subsequent version of that paper will soon find its way into print, since it is an invaluable *vademecum* for any neophyte, in addition to advancing several novel themes. [This hope has been fulfilled: Janlert, 1987.—DCD, 1997]

Since the frame problem, whatever it is, is certainly not solved yet (and may be, in its current guises, insoluble), the ideological foes of AI such as Hubert Dreyfus and John Searle are tempted to compose obituaries for the field, citing the frame problem as the cause of death. In *What Computers Can't Do* (Dreyfus 1972), Dreyfus sought to show that AI was a fundamentally mistaken method for studying the mind, and in fact many of his somewhat impressionistic complaints about AI models and many of his declared insights into their intrinsic limitations can be seen to hover quite systematically in the neighborhood of the frame problem. Dreyfus never explicitly mentions the frame problem,² but is it perhaps the smoking pistol he was looking for but didn't quite know how to describe? Yes, I think AI can be seen to be holding a smoking pistol, but at least in its "whole pudding" guise it is everybody's problem, not just a problem for AI, which, like the good guy in many a mystery story, should be credited with a discovery, not accused of a crime.

One does not have to hope for a robot-filled future to be worried by the frame problem. It apparently arises from some very widely held and innocuous-seeming assumptions about the nature of intelligence, the truth of the most undoctinaire brand of physicalism, and the conviction that it must be possible to explain how we think. (The dualist evades the frame problem—but only because dualism draws the veil of mystery and obfuscation over all the tough how-questions; as we shall see, the problem arises when one takes seriously the task of answering certain how-questions. Dualists inexcusably excuse themselves from the frame problem.) One utterly central—if not defining—feature of an intelligent being is that it can "look before it leaps." Better, it can *think* before it leaps. Intelligence is (at least partly) a matter of using well what you know—but for what? For improving the fidelity

2. Dreyfus mentions McCarthy (1960, pp. 213–214), but the theme of his discussion there is that McCarthy ignores the difference between a *physical state* description and a *situation* description, a theme that might be succinctly summarized: a house is not a home. Similarly, he mentions *ceteris paribus* assumptions (in the Introduction to the Revised Edition, p. 56ff), but only in announcing his allegiance to Wittgenstein's idea that "when-ever human behavior is analyzed in terms of rules, these rules must always contain a *ceteris paribus* condition." But this, even if true, misses the deeper point: the need for something like *ceteris paribus* assumptions confronts Robinson Crusoe just as ineluctably as it confronts any protagonist who finds himself in a situation involving human culture. The point is not, it seems, restricted to *Geisteswissenschaft* (as it is usually conceived); the "intelligent" robot on an (otherwise?) uninhabited but hostile planet faces the frame problem as soon as it commences to plan its days.

of your expectations about what is going to happen next, for planning, for considering courses of action, for framing further hypotheses with the aim of increasing the knowledge you will use in the future, so that you can preserve yourself, by letting your hypotheses die in your stead (as Sir Karl Popper once put it). The stupid—as opposed to ignorant—being is the one who lights the match to peer into the fuel tank,³ who saws off the limb he is sitting on, who locks his keys in his car and then spends the next hour wondering how on earth to get his family out of the car.

But when we think before we leap, *how do we do it?* The answer seems obvious: an intelligent being learns from experience, and then uses what it has learned to guide expectations in the future. Hume explained this in terms of habits of expectation, in effect. *But how do the habits work?* Hume had a hand-waving answer—associationism—to the effect that certain transition paths between ideas grew more likely-to-be-followed as they became well worn, but since it was not Hume's job, surely, to explain in more detail the mechanics of these links, problems about how such paths could be put to good use—and not just turned into an impenetrable maze of untraversable alternatives—were not discovered.

Hume, like virtually all other philosophers and “mentalistic” psychologists, was unable to see the frame problem because he operated at what I call a purely semantic level, or a *phenomenological* level. At the phenomenological level, all the items in view are *individuated by their meanings*. Their meanings are, if you like, “given”—but this just means that the theorist helps himself to all the meanings he wants. In this way the semantic relation between one item and the next is typically plain to see, and one just assumes that the items behave as items with those meanings *ought* to behave. We can bring this out by concocting a Humean account of a bit of learning.

Suppose there are two children, both of whom initially tend to grab cookies from the jar without asking. One child is allowed to do this unmolested but the other is spanked each time she tries. What is the result? The second child learns not to go for the cookies. Why? Because she has had experience of cookie-reaching followed swiftly by spanking. What good does that do? Well, the *idea* of cookie-reaching becomes connected by a habit path to the idea of spanking, which in turn is

3. The example is from an important discussion of rationality by Christopher Cherniak, in “Rationality and the Structure of Memory” (1983).

connected to the idea of pain . . . so *of course* the child refrains. Why? Well, that's just the effect of that idea on that sort of circumstance. But why? Well, what else ought the idea of pain to do on such an occasion? Well, it might cause the child to pirouette on her left foot, or recite poetry or blink or recall her fifth birthday. But given what the idea of pain *means*, any of those effects would be absurd. True; now *how* can ideas be designed so that their effects are what they ought to be, given what they mean? Designing some internal thing—an idea, let's call it—so that it behaves vis-à-vis its brethren as if it meant *cookie* or *pain* is the only way of endowing that thing with that meaning; it couldn't mean a thing if it didn't have those internal behavioral dispositions.

That is the mechanical question the philosophers left to some dimly imagined future researcher. Such a division of labor might have been all right, but it is turning out that most of the truly difficult and deep puzzles of learning and intelligence get kicked downstairs by this move. It is rather as if philosophers were to proclaim themselves expert explainers of the methods of a stage magician, and then, when we ask them to explain how the magician does the sawing-the-lady-in-half trick, they explain that it is really quite obvious: the magician doesn't really saw her in half; he simply makes it appear that he does. "But how does he do *that*?" we ask. "Not our department," say the philosophers—and some of them add, sonorously: "Explanation has to stop somewhere."⁴

When one operates at the purely phenomenological or semantic level, where does one get one's data, and how does theorizing proceed? The term "phenomenology" has traditionally been associated with an introspective method—an *examination* of what is presented or given to consciousness. A person's phenomenology just was by definition the contents of his or her consciousness. Although this has been the ideology all along, it has never been the practice. Locke, for instance, may have thought his "historical, plain method" was a method of unbiased self-observation, but in fact it was largely a matter of disguised aprioristic reasoning about what ideas and impressions *had to be* to do the jobs they "obviously" did.⁵ The myth that each of us can observe our

4. Note that on this unflattering portrayal, the philosophers might still be doing *some* valuable work; the wild goose chase one might avert for some investigator who had rashly concluded that the magician really did saw the lady in half and then miraculously reunite her. People have jumped to such silly conclusions, after all; many philosophers have done so, for instance.

5. See my 1982d, a commentary on Goodman (1982).

mental activities has prolonged the illusion that major progress could be made on the theory of thinking by simply reflecting carefully on our own cases. For some time now we have known better: we have conscious access to only the upper surface, as it were, of the multilevel system of information-processing that occurs in us. Nevertheless, the myth still claims its victims.

So the analogy of the stage magician is particularly apt. One is not likely to make much progress in figuring out *how* the tricks are done by simply sitting attentively in the audience and watching like a hawk. Too much is going on out of sight. Better to face the fact that one must either rummage around backstage or in the wings, hoping to disrupt the performance in telling ways; or, from one's armchair, think aprioristically about how the tricks *must* be done, given whatever is manifest about the constraints. The frame problem is then rather like the unsettling but familiar "discovery" that so far as armchair thought can determine, a certain trick we have just observed is flat impossible.

Here is an example of the trick. Making a midnight snack. How is it that I can get myself a midnight snack? What could be simpler? I suspect there is some sliced leftover turkey and mayonnaise in the fridge, and bread in the bread box—and a bottle of beer in the fridge as well. I realize I can put these elements together, so I concoct a childish simple plan: I'll just go and check out the fridge, get out the requisite materials, and make myself a sandwich, to be washed down with a beer. I'll need a knife, a plate, and a glass for the beer. I forthwith put the plan into action and it works! Big deal.

Now of course I couldn't do this without knowing a good deal—about bread, spreading mayonnaise, opening the fridge, the friction and inertia that will keep the turkey between the bread slices and the bread on the plate as I carry the plate over to the table beside my easy chair. I also need to know about how to get the beer out of the bottle and into the glass.⁶ Thanks to my previous accumulation of experience in the world, fortunately, I am equipped with all this worldly knowledge. Of course some of the knowledge I need *might* be innate. For instance, one trivial thing I have to know is that when the beer gets into the glass it is no longer in the bottle, and that if I'm holding the mayonnaise jar in my left hand I cannot also be spreading the mayonnaise with the knife in my left hand. Perhaps these are straightforward

6. This knowledge of physics is not what one learns in school, but in one's crib. See Hayes 1978, 1980.

implications—instantiations—of some more fundamental things that I was in effect *born knowing* such as, perhaps, the fact that if something is in one location it isn't also in another, different location; or the fact that two things can't be in the same place at the same time; or the fact that situations change as the result of actions. It is hard to imagine just how one could learn these facts from experience.

Such utterly banal facts escape our notice as we act and plan, and it is not surprising that philosophers, thinking phenomenologically *but introspectively*, should have overlooked them. But if one turns one's back on introspection, and just thinks "hetero-phenomenologically" about the purely informational demands of the task—what *must* be known by any entity that can perform this task—these banal bits of knowledge rise to our attention. We can easily satisfy ourselves that no agent that did not *in some sense* have the benefit of information (that beer in the bottle is not in the glass, etc.) could perform such a simple task. It is one of the chief methodological beauties of AI that it makes one be a phenomenologist, one reasons about what the agent must "know" or figure out *unconsciously or consciously* in order to perform in various ways.

The reason AI forces the banal information to the surface is that the tasks set by AI start at zero: the computer to be programmed to simulate the agent (or the brain of the robot, if we are actually going to operate in the real, nonsimulated world), initially knows nothing at all "about the world." The computer is the fabled tabula rasa on which every required item must somehow be impressed, either by the programmer at the outset or via subsequent "learning" by the system.

We can all agree, today, that there could be no learning at all by an entity that faced the world at birth as a tabula rasa, but the dividing line between what is innate and what develops maturationally and what is actually learned is of less theoretical importance than one might have thought. **While some information has to be innate, there is hardly any particular item that must be: an appreciation of *modus ponens*, perhaps, and the law of the excluded middle, and some sense of causality.** And while some things we know must be learned—for example, that Thanksgiving falls on a Thursday, or that refrigerators keep food fresh—many other "very empirical" things could in principle be innately known—for example, that smiles mean happiness, or that un-

7. For elaborations of hetero-phenomenology, see Dennett 1978a, chapter 10, "Two Approaches to Mental Images," and Dennett 1982b. See also Dennett 1982a and 1991.

suspended, unsupported things fall. (There is some evidence, in fact, that there is an innate bias in favor of perceiving things to fall with gravitational acceleration).⁸

Taking advantage of this advance in theoretical understanding (if that is what it is), people in AI can frankly ignore the problem of learning (it seems) and take the shortcut of *installing* all that an agent has to "know" to solve a problem. After all, if God made Adam as an adult who could presumably solve the midnight snack problem *ab initio*, AI agent-creators can *in principle* make an "adult" agent who is equipped with worldly knowledge *as if* it had laboriously learned all the things it needs to know. This may of course be a dangerous shortcut.

The installation problem is then the problem of installing in one way or another all the information needed by an agent to plan in a changing world. It is a difficult problem because the information must be installed in a usable format. The problem can be broken down initially into the semantic problem and the syntactic problem. The semantic problem—called by Allan Newell the problem at the "knowledge level" (Newell, 1982)—is the problem of just what information (on what topics, to what effect) must be installed. The syntactic problem is what system, format, structure, or mechanism to use to put that information in.⁹

The division is clearly seen in the example of the midnight snack problem. I *listed* a few of the very many humdrum facts one needs to know to solve the snack problem, but I didn't mean to suggest that those facts are stored in me—or in any agent—piecemeal, in the form

8. Gunnar Johannsen has shown that animated films of "falling" objects in which the moving spots drop with the normal acceleration of gravity are unmistakably distinguished by the casual observer from "artificial" motions. I do not know whether infants have been tested to see if they respond selectively to such displays.

9. McCarthy and Hayes (1969) draw a different distinction between the "epistemological" and the "heuristic." The difference is that they include the question "In what kind of internal notation is the system's knowledge to be expressed?" in the epistemological problem (see p. 466), dividing off *that* syntactic (and hence somewhat mechanical) question from the procedural questions of the design of "the mechanism that on the basis of the information solves the problem and decides what to do."

One of the prime grounds for controversy about just which problem the frame problem is springs from this attempted division of the issue. For the answer to the syntactical aspects of the epistemological question makes a large difference to the nature of the heuristic problem. After all, if the syntax of the expression of the system's knowledge is sufficiently perverse, then in spite of the *accuracy* of the representation of that knowledge, the heuristic problem will be impossible. And some have suggested that the heuristic problem would virtually disappear if the world knowledge were felicitously couched in the first place.

of a long list of sentences explicitly declaring each of these facts for the benefit of the agent. That is of course one possibility, officially: it is a preposterously extreme version of the “language of thought” theory of mental representation, with each distinguishable “proposition” separately inscribed in the system. No one subscribes to such a view; even an encyclopedia achieves important economies of explicit expression via its organization, and a walking encyclopedia—not a bad caricature of the envisaged AI agent—must use different systemic principles to achieve efficient representation and access. We know trillions of things; we know that mayonnaise doesn’t dissolve knives on contact, that a slice of bread is smaller than Mount Everest, that opening the refrigerator doesn’t cause a nuclear holocaust in the kitchen.

There must be in us—and in any intelligent agent—some highly efficient, partly generative or productive system of representing—storing for use—all the information needed. Somehow, then, we must store many “facts” at once—where facts are presumed to line up more or less one-to-one with nonsynonymous declarative sentences. Moreover, we cannot realistically hope for what one might call a Spinozistic solution—a *small* set of axioms and definitions from which all the rest of our knowledge is deducible on demand—since it is clear that there simply are no entailment relations between vast numbers of these facts. (When we rely, as we must, on experience to tell us how the world is, experience tells us things that do not at all follow from what we have heretofore known.)

The demand for an efficient system of information storage is in part a space limitation, since our brains are not all that large, but more importantly it is a time limitation, for stored information that is not reliably accessible for use in the short real-time spans typically available to agents in the world is of no use at all. A creature that can solve any problem given enough time—say a million years—is not in fact intelligent at all. We live in a time-pressured world and must be able to think quickly before we leap. (One doesn’t have to view this as an *a priori* condition on intelligence. One can simply note that we do in fact think quickly, so there is an empirical question about how we manage to do it.)

The task facing the AI researcher appears to be designing a system that can plan by using well-selected elements from its store of knowledge about the world it operates in. “Introspection” on how *we* plan yields the following description of a process: one envisages a certain situation (often very sketchily); one then imagines performing a certain

act in that situation; one then "sees" what the likely outcome of that envisaged act in that situation would be, and evaluates it. What happens backstage, as it were, to permit this "seeing" (and render it as reliable as it is) is utterly inaccessible to introspection.

On relatively rare occasions we all experience such bouts of thought, unfolding in consciousness at the deliberate speed of pondering. These are the occasions in which we are faced with some novel and relatively difficult problem, such as: How can I get the piano upstairs? or Is there any way to electrify the chandelier without cutting through the plaster ceiling? It would be quite odd to find that one had to think *that* way (consciously and slowly) in order to solve the midnight snack problem. But the suggestion is that even the trivial problems of planning and bodily guidance that are beneath our notice (though in some sense we "face" them) are solved by similar processes. Why? I don't *observe* myself planning in such situations. This fact suffices to convince the traditional, introspective phenomenologist that no such planning is going on.¹⁰ The hetero-phenomenologist, on the other hand, reasons that *one way or another* information about the objects in the situation, and about the intended effects and side effects of the candidate actions, *must* be used (considered, attended to, applied, appreciated). Why? Because otherwise the "smart" behavior would be sheer luck or magic. (Do we have any model for how such unconscious information-appreciation might be accomplished? The only model we have *so far* is conscious, deliberate information-appreciation. Perhaps, AI suggests, this is a good model. If it isn't, we are all utterly in the dark for the time being.)

We assure ourselves of the intelligence of an agent by considering counterfactuals: if I had been told that the turkey was poisoned, or the beer explosive, or the plate dirty, or the knife too fragile to spread mayonnaise; would I have acted as I did? If I were a stupid "automaton"—or like the *Sphex* wasp you "mindlessly" repeats her stereotyped burrow-checking routine till she drops¹¹—I might infelicitously "go

10. Such observations also convinced Gilbert Ryle, who was, in an important sense, an introspective phenomenologist (and not a "behaviorist"). See Ryle 1949.

One can readily imagine Ryle's attack on AI: "And *how many* inferences do I perform in the course of preparing my sandwich? What syllogisms convince me that the beer will stay in the glass?" For a further discussion of Ryle's skeptical arguments and their relation to cognitive science, see my (1983b) "Styles of Representation."

11. "When the time comes for egg laying the wasp *Sphex* builds a burrow for the purpose and seeks out a cricket which she stings in such a way as to paralyze but not kill it. She drags the cricket into her burrow, lays her eggs alongside, closes the burrow, then flies away, never to return. In due course, the eggs hatch and the wasp grubs feed off the paralyzed cricket, which has not decayed, having been kept in the wasp equivalent of

through the motions" of making a midnight snack oblivious to the recalcitrant features of the environment.¹² But in fact, my midnight-snack-making behavior is multifariously sensitive to current and background information about the situation. The only way it could be so sensitive—runs the tacit hetero-phenomenological reasoning—is for it to examine, or test for, the information in question. This information manipulation may be unconscious and swift, and it need not (it *better* not) consist of hundred or thousands of *seriatim* testing procedures, but it must occur somehow, and its benefits must appear in time to help me as I commit myself to action.

I may of course have a midnight snack routine, developed over the years, in which case I can partly rely on it to pilot my actions. Such a complicated "habit" would have to be under the control of a mechanism of some complexity, since even a rigid sequence of steps would involve periodic testing to ensure that subgoals had been satisfied. And even if I am an infrequent snacker, I no doubt have routines for mayonnaise-spreading, sandwich-making, and getting something out of the fridge, from which I could compose my somewhat novel activity. Would such ensembles of routines, nicely integrated, suffice to solve the frame problem for me, at least in my more "mindless" endeavors? That is an open question to which I will return below.

It is important in any case to acknowledge at the outset, and remind oneself frequently, that even very intelligent people do make mistakes; we are not only not infallible planners, we are quite prone to overlooking large and retrospectively obvious flaws in our plans. This foible manifests itself in the familiar case of "force of habit" errors (in which our stereotypical routines reveal themselves to be surprisingly insensi-

a deep freeze. To the human mind, such an elaborately organized and seemingly purposeful routine conveys a convincing flavor of logic and thoughtfulness—until more details are examined. For example, the wasp's routine is to bring the paralyzed cricket to the burrow, leave it on the threshold, go inside to see that all is well, emerge, and then drag the cricket in. If, while the wasp is inside making her preliminary inspection the cricket is moved a few inches away, the wasp, on emerging from the burrow, will bring the cricket back to the threshold, but not inside, and will then repeat the preparatory procedure of entering the burrow to see that everything is all right. If again the cricket is removed a few inches while the wasp is inside, once again the wasp will move the cricket up to the threshold and re-enter the burrow for a final check. The wasp never thinks of pulling the cricket straight in. On one occasion, this procedure was repeated forty times, always with the same result" (Wooldridge 1963).

This vivid example of a familiar phenomenon among insects is discussed by me in *Brainstorms*, and in Douglas R. Hofstadter 1982.

12. See my 1982a: 58–59, on "Robot Theater."

tive to some portentous environmental changes while surprisingly sensitive to others). The same weakness also appears on occasion in cases where we have consciously deliberated with some care. How often have you embarked on a project of the piano-moving variety—in which you've thought through or even "walked through" the whole operation in advance—only to discover that you must backtrack or abandon the project when some perfectly foreseeable but unforeseen obstacle or unintended side effect loomed? If we smart folk seldom actually paint ourselves into corners, it may not be because we plan ahead so well as that we supplement our sloppy planning powers with a combination of recollected lore (about fools who paint themselves into corners, for instance) and frequent progress checks as we proceed. Even so, we must know enough to call up the right lore at the right time, and to recognize impending problems as such.

To summarize: we have been led by fairly obvious and compelling considerations to the conclusion that an intelligent agent must engage in swift information-sensitive "planning" which has the effect of producing reliable but not foolproof expectations of the effects of its actions. That these expectations are normally in force in intelligent creatures is testified to by the startled reaction they exhibit when their expectations are thwarted. This suggests a graphic way of characterizing the minimal goal that can spawn the frame problem: we want a midnight-snack-making robot to be "surprised" by the trick plate, the unspreadable concrete mayonnaise, the fact that we've glued the beer glass to the shelf. To be surprised you have to have expected something else, and in order to have expected the right something else, you have to have *and use* a lot of information about the things in the world.¹³

The central role of expectation has led some to conclude that the

13. Hubert Dreyfus has pointed out that *not expecting x* does not imply *expecting y* (where $x \neq y$), so one can be startled by something one didn't expect without its having to be the case that one (unconsciously) expected something else. But this sense of *not expecting* will not suffice to explain startle. What are the odds against your seeing an Alfa Romeo, a Buick, a Chevrolet, and a Dodge parked in alphabetical order some time or other within the next five hours? Very high, no doubt, all things considered, so I would not expect you to expect this; I also would not expect you to be startled by seeing this unexpected sight—except in the sort of special case where you had reason to expect something else at that time and place.

Startle reactions are powerful indicators of cognitive state—a fact long known by the police (and writers of detective novels). *Only* someone who expected the refrigerator to contain Smith's corpse (say) would be startled (as opposed to mildly interested) to find it to contain the rather unlikely trio: a bottle of vintage Chablis, a can of cat food, and a dishrag.

frame problem is not a new problem at all, and has nothing particularly to do with planning actions. It is, they think, simply the problem of having good expectations about any future events, whether they are one's own actions, the actions of another agent, or mere happenings of nature. That is the problem of induction—noted by Hume and intensified by Goodman (Goodman 1965), but still not solved to anyone's satisfaction. We know today that the problem is a nasty one indeed. Theories of subjective probability and belief fixation have not been stabilized in reflective equilibrium, so it is fair to say that no one has a good, principled answer to the general question: given that I believe all *this* (have all this evidence), what *ought* I to believe as well (about the future, or about unexamined parts of the world)?

The reduction of one unsolved problem to another is some sort of progress, unsatisfying though it may be, but it is not an option in this case. The frame problem is not the problem of induction in disguise. For suppose the problem of induction were solved. Suppose—perhaps miraculously—that our agent has solved all its induction problems or had them solved by fiat; it believes, then, all the right generalizations from its evidence, and associates with all of them the appropriate probabilities and conditional probabilities. This agent, *ex hypothesi*, believes just what it ought to believe about all empirical matters in its ken, including the probabilities of future events. It might still have a bad case of the frame problem, for that problem concerns how to represent (so it can be *used*) all that hard-won empirical information—a problem that arises independently of the truth value, probability, warranted assertability, or subjective certainty of any of it. Even if you have excellent *knowledge* (and not mere belief) about the changing world, how can this knowledge be represented so that it can be efficaciously brought to bear?

Recall poor R1D1, and suppose for the sake of argument that it had perfect empirical knowledge of the probabilities of all the effects of all its actions that would be detectable by it. Thus it believes that with probability 0.7864, executing PULLOUT (WAGON, ROOM) will cause the wagon wheels to make an audible noise; and with probability 0.5, the door to the room will open in rather than out; and with probability 0.999996, there will be no live elephants in the room, and with probability 0.997 the bomb will remain on the wagon when it is moved. How is R1D1 to find this last, relevant needle in its haystack of empirical knowledge? A walking encyclopedia will walk over a cliff, for all its knowledge of cliffs and the effects of gravity, unless it is designed in

such a fashion that it can find the right bits of knowledge at the right times, so it can plan its engagements with the real world.

The earliest work on planning systems in AI took a deductive approach. Inspired by the development of Robinson's methods of resolution theorem proving, designers hoped to represent all the system's "world knowledge" explicitly as axioms, and use ordinary logic—the predicate calculus—to deduce the effects of actions. Envisaging a certain situation *S* was modeled by having the system entertain a set of axioms describing the situation. Added to this were background axioms (the so-called frame axioms that give the frame problem its name) which describe general conditions and the general effects of every action type defined for the system. To this set of axioms the system would apply an action—by postulating the occurrence of some action *A* in situation *S*—and then deduce the effect of *A* in *S*, producing a description of the outcome situation *S'*. While all this logical deduction looks like nothing at all in our conscious experience, research on deductive approach could proceed on either or both of two enabling assumptions: the methodological assumption that psychological realism was a gratuitous bonus, not a goal, of "pure" AI, or the substantive (if still vague) assumption that the deductive processes described would somehow model the backstage processes beyond conscious access. In other words, either we don't do our thinking deductively in the predicate calculus but a robot might; or we do (unconsciously) think deductively in the predicate calculus. Quite aside from doubts about its psychological realism, however, the deductive approach has not been made to work—the proof of the pudding for any robot—except for deliberately trivialized cases.

Consider some typical frame axioms associated with the action type: *move x onto y*.

1. If $z \neq x$ and I move *x* onto *y*, then if *z* was on *w* before, then *z* is on *w* after
2. If *x* is blue before, and I move *x* onto *y*, then *x* is blue after

Note that (2), about being blue, is just one example of the many boring "no-change" axioms we have to associate with this action type. Worse still, note that a cousin of (2), also about being blue, would have to be associated with every other action-type—with *pick up x* and with *give x to y*, for instance. One cannot save this mindless repetition by postulating once and for all something like

3. If anything is blue, it stays blue

for that is false, and in particular we will want to leave room for the introduction of such action types as *paint x red*. Since virtually any aspect of a situation can change under some circumstance, this method requires introducing for each aspect (each predication in the description of S) an axiom to handle whether that aspect changes for each action type.

This representational profligacy quickly gets out of hand, but for some "toy" problems in AI, the frame problem can be overpowered to some extent by a mixture of the toyness of the environment and brute force. The early version of SHAKEY, the robot at S.R.I., operated in such a simplified and sterile world, with so few aspects it could worry about that it could get away with an exhaustive consideration of frame axioms.¹⁴

Attempts to circumvent this explosion of axioms began with the proposal that the system operate on the tacit assumption that nothing changes in a situation but what is explicitly asserted to change in the definition of the applied action (Fikes and Nilsson, 1971). The problem here is that, as Garrett Hardin once noted, you can't do just one thing. This was RI's problem, when it failed to notice that it would pull the bomb out with the wagon. In the explicit representation (a few pages back) of my midnight snack solution, I mentioned carrying the plate over to the table. On this proposal, my model of S' would leave the turkey back in the kitchen, for I didn't explicitly say the turkey would come along with the plate. One can of course patch up the definition of "bring" or "plate" to handle just this problem, but only at the cost of creating others. (Will a few more patches tame the problem? At what point should one abandon patches and seek an altogether new approach? Such are the methodological uncertainties regularly encountered in this field, and of course no one can responsibly claim in advance to have a good rule for dealing with them. Premature counsels of despair or calls for revolution are as clearly to be shunned as the dogged pursuit of hopeless avenues; small wonder the field is contentious.)

While one cannot get away with the tactic of supposing that one can do just one thing, it remains true that very little of what could (logically) happen in any situation does happen. Is there some way of falli-

14. This early feature of SHAKEY was drawn to my attention by Pat Hayes. See also Dreyfus (1972, p. 26). SHAKEY is put to quite a different use in Dennett (1982b).

bly marking the likely area of important side effects, and assuming the rest of the situation to stay unchanged? Here is where relevance tests seem like a good idea, and they may well be, but not within the deductive approach. As Minsky notes:

Even if we formulate relevancy restrictions, logistic systems have a problem using them. In any logistic system, all the axioms are necessarily "permissive"—they all help to permit new inferences to be drawn. Each added axiom means more theorems; none can disappear. There simply is no direct way to add information to tell such a system about kinds of conclusions that should *not* be drawn! . . . If we try to change this by adding axioms about relevancy, we still produce all the unwanted theorems, plus annoying statements about their irrelevancy. (Minsky, 1981, p. 125)

What is needed is a system that genuinely *ignores* most of what it knows, and operates with a well-chosen portion of its knowledge at any moment. Well-chosen, but not chosen by exhaustive consideration. How, though, can you give a system *rules* for ignoring—or better, since explicit rule following is not the problem, how can you design a system that reliably ignores what it ought to ignore under a wide variety of different circumstances in a complex action environment?

John McCarthy calls this the *qualification problem*, and vividly illustrates it via the famous puzzle of the missionaries and the cannibals.

Three missionaries and three cannibals come to a river. A rowboat that seats two is available. If the cannibals ever outnumber the missionaries on either bank of the river, the missionaries will be eaten. How shall they cross the river?

Obviously the puzzler is expected to devise a strategy of rowing the boat back and forth that gets them all across and avoids disaster. . . .

Imagine giving someone the problem, and after he puzzles for awhile, he suggests going upstream half a mile and crossing on a bridge. "What bridge?" you say. "No bridge is mentioned in the statement of the problem." And this dunce replies, "Well, they don't say there isn't a bridge." You look at the English and even at the translation of the English into first order logic, and you must admit that "they don't say" there is no bridge. So you modify the problem to exclude bridges and pose it again, and the dunce proposes a helicopter, and after you exclude that, he proposes a winged horse or that the others hang onto the outside of the boat while two row.

You now see that while a dunce, he is an inventive dunce. Despairing of getting him to accept the problem in the proper puzzler's spirit, you tell him the solution. To your further annoyance, he attacks your solution on the grounds that the boat might have a leak or lack oars. After you rectify that omission from the statement of problem, he suggests that a sea monster may swim up the river and may swallow the boat. Again you are frustrated, and you look for a mode of reasoning that will settle his hash once and for all. (McCarthy, 1980, pp. 29–30)

What a normal, intelligent human being does in such a situation is to engage in some form of *nonmonotonic inference*. In a classical, monotonic logical system, *adding* premises never *diminishes* what can be proved from the premises. As Minsky noted, the axioms are essentially permissive, and once a theorem is permitted, adding more axioms will never invalidate the proofs of earlier theorems. But when we think about a puzzle or a real life problem, we can achieve a solution (and even prove that it is a solution, or even the only solution to *that* problem), and then discover our solution invalidated by the addition of a new element to the posing of the problem; For example, "I forgot to tell you—there are no oars" or "By the way, there's a perfectly good bridge upstream."

What such late additions show us is that, contrary to our assumption, other things weren't equal. We had been reasoning with the aid of a *ceteris paribus* assumption, and now our reasoning has just been jeopardized by the discovery that something "abnormal" is the case. (Note, by the way, that the abnormality in question is a much subtler notion than anything anyone has yet squeezed out of probability theory. As McCarthy notes, "The whole situation involving cannibals with the postulated properties cannot be regarded as having a probability, so it is hard to take seriously the conditional probability of a bridge given the hypothesis" [*ibid.*].)

The beauty of a *ceteris paribus* clause in a bit of reasoning is that one does not have to say exactly what it means. "What do you mean, 'other things being equal'? Exactly which arrangements of which other things count as being equal?" If one had to answer such a question, invoking the *ceteris paribus* clause would be pointless, for it is precisely in order to evade that task that one uses it. If one could answer that question, one wouldn't need to invoke the clause in the first place. One way of viewing the frame problem, then, is as the attempt to get a computer to avail itself of this distinctively human style of mental operation. There are several quite different approaches to nonmonotonic inference being pursued in AI today. They have in common only the goal of capturing the human talent for *ignoring* what should be ignored, while staying alert to relevant recalcitrance when it occurs.

One family of approaches, typified by the work of Marvin Minsky and Roger Schank (Minsky 1981; Schank and Abelson 1977), gets its ignoring-power from the attention-focussing power of stereotypes. The inspiring insight here is the idea that all of life's experiences, for all

their variety, boil down to variations on a manageable number of stereotypic themes, paradigmatic scenarios—"frames" in Minsky's terms, "scripts" in Schank's.

An artificial agent with a well-stocked compendium of frames or scripts, appropriately linked to each other and to the impingements of the world via its perceptual organs, would face the world with an elaborate system of what might be called habits of attention and benign tendencies to leap to particular sorts of conclusions in particular sorts of circumstances. It would "automatically" pay attention to certain features in certain environments and assume that certain unexamined normal features of those environments were present. Concomitantly, it would be differentially alert to relevant divergences from the stereotypes it would always be "expecting."

Simulations of fragments of such an agent's encounters with its world reveal that in many situations it behaves quite felicitously and apparently naturally, and it is hard to say, of course, what the limits of this approach are. But there are strong grounds for skepticism. Most obviously, while such systems perform creditably when the world cooperates with their stereotypes, and even with *anticipated* variations on them, when their worlds turn perverse, such systems typically cannot recover gracefully from the misanalyses they are led into. In fact, their behavior in extremis looks for all the world like the preposterously counterproductive activities of insects betrayed by their rigid tropisms and other genetically hard-wired behavioral routines.

When these embarrassing misadventures occur, the system designer can improve the design by adding provisions to deal with the particular cases. It is important to note that in these cases, the system does not redesign itself (or learn) but rather must wait for an external designer to select an improved design. This process of redesign recapitulates the process of natural selection in some regards; it favors minimal, piecemeal, ad hoc redesign which is tantamount to a wager on the likelihood of patterns in future events. So in some regards it is faithful to biological themes.¹⁵

Several different sophisticated attempts to provide the representational framework for this deeper understanding have emerged from

15. In one important regard, however, it is dramatically unlike the process of natural selection, since the trial, error and selection of the process is far from blind. But a case can be made that the impatient researcher does nothing more than telescope time by such foresighted interventions in the redesign process.

the deductive tradition in recent years. Drew McDermott and Jon Doyle have developed a "nonmonotonic logic" (1980), Ray Reiter has a "logic for default reasoning" (1980), and John McCarthy has developed a system of "circumscription," a formalized "rule of conjecture that can be used by a person or program for 'jumping to conclusions'" (1980). None of these is, or is claimed to be, a complete solution to the problem of *ceteris paribus* reasoning, but they might be components of such a solution. More recently, McDermott (1982) has offered a "temporal logic for reasoning about processes and plans." I will not attempt to assay the formal strengths and weaknesses of these approaches. Instead I will concentrate on another worry. From one point of view, nonmonotonic or default logic, circumscription, and temporal logic all appear to be radical improvements to the mindless and clanking deductive approach, but from a slightly different perspective they appear to be more of the same, and at least as unrealistic as frameworks for psychological models.

They appear in the former guise to be a step toward greater psychological realism, for they take seriously, and attempt to represent, the phenomenologically salient phenomenon of common sense *ceteris paribus* "jumping to conclusions" reasoning. But do they really succeed in offering any plausible suggestions about how the backstage implementation of that conscious thinking is accomplished *in people*? Even if on some glorious future day a robot with debugged circumscription methods maneuvered well in a non-toy environment, would there be much likelihood that its constituent processes, *described at levels below the phenomenological*, would bear informative relations to the unknown lower-level backstage processes in human beings? To bring out better what my worry is, I want to introduce the concept of a *cognitive wheel*.

We can understand what a cognitive wheel might be by reminding ourselves first about ordinary wheels. Wheels are wonderful, elegant triumphs of technology. The traditional veneration of the mythic inventor of the wheel is entirely justified. But if wheels are so wonderful, why are there no animals with wheels? Why are no wheels to be found (functioning as wheels) in nature? First, the presumption of these questions must be qualified. A few years ago the astonishing discovery was made of several microscopic beasties (some bacteria and some unicellular eukaryotes) that have wheels of sorts. Their propulsive tails, long thought to be flexible flagella, turn out to be more or less rigid corkscrews, which rotate continuously, propelled by microscopic motors

of sorts, complete with main bearings.¹⁶ Better known, if less interesting for obvious reasons, are tumbleweeds. So it is not quite true that there are no wheels (or wheeliform designs) in nature.

Still, macroscopic wheels—reptilian or mammalian or avian wheels—are not to be found. Why not? They would seem to be wonderful retractable landing gear for some birds, for instance. Once the question is posed, plausible reasons rush in to explain their absence. Most important, probably, are the considerations about the topological properties of the axle / bearing boundary that make the transmission of material or energy across it particularly difficult. How could the life-support traffic arteries of a living system maintain integrity across this boundary? But once that problem is posed, solutions suggest themselves; suppose the living wheel grows to mature form in a nonrotating, nonfunctional form, and is then hardened and sloughed off, like antlers or an outgrown shell, but not completely off: it then rotates freely on a lubricated fixed axle. Possible? It's hard to say. Useful? Also hard to say, especially since such a wheel would have to be freewheeling. This is an interesting speculative exercise, but certainly not one that should inspire us to draw categorical, *a priori* conclusions. It would be foolhardy to declare wheels biologically impossible, but at the same time we can appreciate that they are at least very distant and unlikely solutions to *natural* problems of design.

Now a cognitive wheel is simply any design proposal in cognitive theory (at any level from the purest semantic level to the most concrete level of "wiring diagrams" of the neurons) that is profoundly unbiological, however wizardly and elegant it is as a bit of technology.

Clearly this is a vaguely defined concept, useful only as a rhetorical abbreviation, as a gesture in the direction of real difficulties to be spelled out carefully. "Beware of postulating cognitive wheels" masquerades as good advice to the cognitive scientist, while courting vacuity as a maxim to follow.¹⁷ It occupies the same rhetorical position as

16. For more details, and further reflections on the issues discussed here, see Diamond (1983).

17. I was interested to discover that at least one researcher in AI mistook the rhetorical intent of my new term on first hearing; he took "cognitive wheels" to be an accolade. If one thinks of AI as he does, not as a research method in psychology but as a branch of engineering attempting to extend human cognitive powers, then of course cognitive wheels are breakthroughs. The vast and virtually infallible memories of computers would be prime examples; others would be computers' arithmetical virtuosity and invulnerability to boredom and distraction. See Hofstadter (1982) for an insightful discussion of the relation of boredom to the structure of memory and the conditions for creativity.

the stockbroker's maxim: buy low and sell high. Still, the term is a good theme-fixer for discussion.

Many critics of AI have the conviction that *any* AI system is and must be nothing but a gearbox of cognitive wheels. This could of course turn out to be true, but the usual reason for believing it is based on a misunderstanding of the methodological assumptions of the field. When an AI model of some cognitive phenomenon is proposed, the model is describable at many different levels, from the most global, phenomenological level at which the behavior is described (with some presumptuousness) in ordinary mentalistic terms, down through various levels of implementation all the way to the level of program code—and even further down, to the level of fundamental hardware operations if anyone cares. No one supposes that the model maps onto the processes of psychology and biology *all the way down*. The claim is only that for some high level or levels of description below the phenomenological level (which merely *sets* the problem) there is a mapping of model features onto what is being modeled: the cognitive processes in living creatures, human or otherwise. It is understood that all the implementation details below the level of intended modeling will consist, no doubt, of cognitive wheels—bits of unbiological computer activity mimicking the gross effects of cognitive subcomponents by using methods utterly unlike the methods still to be discovered in the brain. Someone who failed to appreciate that a model composed microscopically of cognitive wheels could still achieve a fruitful isomorphism with biological or psychological processes at a higher level of aggregation would suppose there were good *a priori* reasons for generalized skepticism about AI.

But allowing for the possibility of valuable intermediate levels of modeling is not ensuring their existence. In a particular instance a model might descend directly from a phenomenologically recognizable level of psychological description to a cognitive wheels implementation without shedding any light at all on how we human beings manage to enjoy that phenomenology. I *suspect* that all current proposals in the field for dealing with the frame problem have that shortcoming. Perhaps one should dismiss the previous sentence as mere autobiography. I find it hard to imagine (for what that is worth) that any of the *procedural details* of the mechanization of McCarthy's circumscriptions, for instance, would have suitable counterparts in the backstage story yet to be told about how human common-sense reasoning is accom-

plished. If these procedural details lack "psychological reality" then there is nothing left in the proposal that might model psychological processes except the phenomenological-level description in terms of jumping to conclusions, ignoring and the like—and we already know we do that.

There is an alternative defense of such theoretical explorations, however, and I think it is to be taken seriously. One can claim (and I take McCarthy to claim) that while formalizing common-sense reasoning in his fashion would not tell us anything *directly* about psychological processes of reasoning, it would clarify, sharpen, systematize the purely semantic-level characterization of the demands on any such implementation, biological or not. Once one has taken the giant step forward of taking information-processing seriously as a real process in space and time, one can then take a small step back and explore the implications of that advance at a very abstract level. Even at this very formal level, the power of circumscription and the other versions of nonmonotonic reasoning remains an open but eminently explorable question.¹⁸

Some have thought that the key to a more realistic solution to the frame problem (and indeed, in all likelihood, to any solution at all) must require a complete rethinking of the semantic-level setting, prior to concern with syntactic-level implementation. The more or less standard array of predicates and relations chosen to fill out the predicate-calculus format when representing the "propositions believed" may embody a fundamentally inappropriate parsing of nature for this task. Typically, the interpretation of the formulae in these systems breaks the world down along the familiar lines of objects with properties at times and places. Knowledge of situations and events in the world is represented by what might be called sequences of verbal snapshots. State S , constitutively described by a list of sentences true at time t asserting various n -adic predicates true of various particulars, gives way to state S' , a similar list of sentences true at t' . Would it perhaps be better to reconceive of the world of planning in terms of histories and processes?¹⁹ Instead of trying to model the capacity to *keep track of*

18. McDermott (1982, "A Temporal Logic for Reasoning about Processes and Plans," Section 6, "A Sketch of an Implementation, ") shows strikingly how many *new* issues are raised once one turns to the question of implementation, and how indirect (but still useful) the purely formal considerations are.

19. Patrick Hayes has been exploring this theme, and a preliminary account can be found in "Naive Physics 1: The Ontology of Liquids" (1978).

things in terms of principles for passing through temporal cross-sections of knowledge expressed in terms of terms (*names* for *things*, in essence) and predicates, perhaps we could model keeping track of things more directly, and let all the cross-sectional information about what is deemed true moment by moment be merely implicit (and hard to extract—as it is for us) from the format. These are tempting suggestions, but so far as I know they are still in the realm of handwaving.²⁰

Another, perhaps related, handwaving theme is that the current difficulties with the frame problem stem from the conceptual scheme engendered by the serial-processing von Neumann architecture of the computers used to date in AI. As large, fast parallel processors are developed, they will bring in their wake huge conceptual innovations which are now of course only dimly imaginable. Since brains are surely massive parallel processors, it is tempting to suppose that the concepts engendered by such new hardware will be more readily adaptable for realistic psychological modeling. But who can say? For the time being, most of the optimistic claims about the powers of parallel processing belong in the same camp with the facile observations often encountered in the work of neuroscientists, who postulate marvelous cognitive powers for various portions of the nervous system without a clue of how they are realized.²¹

Filling in the details of the gap between the phenomenological magic show and the well-understood powers of small tracts of brain tissue is the immense research task that lies in the future for theorists of every persuasion. But before the problems can be solved they must be encountered, and to encounter the problems one must step resolutely into the gap and ask how-questions. What philosophers (and everyone else) have always known is that people—and no doubt all intelligent agents—can engage in swift, sensitive, risky-but-valuable *ceteris pari-*

20. Oliver Selfridge's unpublished monograph, *Tracking and Trailing*, promises to push back this frontier, I think, but I have not yet been able to assimilate its messages. [Nor, after many years of cajoling, have I succeeded in persuading Selfridge to publish it. It is still, in 1997, unpublished.] There are also suggestive passages on this topic in Ruth Garrett Millikan's *Language, Thought, and Other Biological Categories*, Cambridge, MA: Bradford Books / The MIT Press, 1984.

21. To balance the "top-down" theorists' foible of postulating cognitive wheels, there is the "bottom-up" theorists' penchant for discovering *wonder tissue*. Wonder tissue appears in many locales. J.J. Gibson's (1979) theory of perception, for instance, seems to treat the whole visual system as a hunk of wonder tissue, resonating with marvelous sensitivity to a host of sophisticated "affordances."

bus reasoning. How do we do it? AI may not yet have a good answer, but at least it has encountered the question.²²

22. One of the few philosophical articles I have uncovered that seem to contribute to the thinking about the frame problem—though not in those terms—is Ronald de Sousa's "The Rationality of Emotions" (de Sousa, 1979). In the section entitled "What are Emotions For?" de Sousa suggests, with compelling considerations, that:

the function of emotion is to fill gaps left by [mere wanting plus] "pure reason" in the determination of action and belief. Consider how Iago proceeds to make Othello jealous. His task is essentially to direct Othello's attention, to suggest questions to ask . . . Once attention is thus directed, inferences which, before on the same evidence, would not even have been thought of, are experienced as compelling.

In de Sousa's understanding, "emotions are determinate patterns of salience among objects of attention, lines of inquiry, and inferential strategies" (p. 50) and they are not "reducible" in any way to "articulated propositions. Suggestive as this is, it does not, of course, offer any concrete proposals for how to endow an inner (emotional) state with these interesting powers. Another suggestive—and overlooked—paper is Howard Darmstadter's "Consistency of Belief" (Darmstadter, 1971, pp. 301–310). Darmstadter's exploration of *ceteris paribus* clauses and the relations that might exist between beliefs as psychological states and sentences believers may utter (or have uttered about them) contains a number of claims that deserve further scrutiny.