

Book Reviews

Actors: A Model of Concurrent Computation in Distributed Systems

Varol Akman

Gul A. Agha's *Actors: A Model of Concurrent Computation in Distributed Systems* (The MIT Press, Cambridge, Mass., 1987, 144 pages, \$25.00, ISBN 0-262-01092-5) is part of the MIT Press Series in Artificial Intelligence. This volume is edited by Patrick Winston, Michael Brady, and Daniel Bobrow.

In the actor formalism, pioneered by Carl Hewitt (1977), one perceives abstract computational agents, called *actors*, that are distributed in space. Each actor has a mailbox (and a mail address) and associated with each actor is a behavior. One actor can influence the actions of another actor only by sending it a communication. An actor can send another actor a communication only when it knows the address of the second actor. An actor can (should be able to) receive more than one message simultaneously (in some sense). In case of synchronous communication in nonactor paradigms, this simultaneous communication requires the sender to wait until the recipient of the communication is ready to accept it. Once the recipient starts processing a given communication, it is blocked from accepting any other messages until after it is done with the message—all for the sake of maintaining the integrity of communications, no doubt.

The actor model achieves the same effect through buffering. A recipient's mailbox is always ready to receive messages. Furthermore, in the actor framework, there are no restrictions on freely communicating mail addresses; thus, the interconnections of the actors can be dynamically modified, and resource management is more flexibly done.

Should the reader begin pondering the difference between actors and objects (as found in object-oriented programming), I would like to make a few clarifying observations. Objects

in Simula (a language owed to O. J. Dahl, B. Myhrhaug, and K. Nygaard and developed in the Norwegian Computing Center in 1970) can be conceived of as ancestors of actors because they represent an individual data item and all the operations and procedures on it. However, actors are more powerful than sequential processes and data flow systems. One can represent a purely functional system with an actor system. Similarly, it is possible to specify arbitrary sequential processes with a suitable actor system. However, the converse is not true. Although actor creation is an integral part of the computational model of Hewitt and Agha, sequential processes such as communicating sequential processes (CSP), owed to Tony Hoare, do not create other sequential processes (but can activate other sequential processes). The creation of actors permits one to increase the distributivity of computation as it evolves.

Is programming a collection of actors easy? It is commonly accepted that parallel programming is, in general, hard, and the need for yet another programming language such as actors might seem difficult to argue. One crucial observation here is that because message passing is fundamental to computation in actors, the time complexity of communication might become the dominant factor in program execution. Thus, architectural considerations might play a crucial role in the realistic implementations of actor languages. Note, however, that the most attractive feature of actor languages is that programmers are not required to code details such as when and where to use parallelism. Also, although some functional programming languages might have problems with history-sensitive shared objects, actors can easily model such objects.

Recently, there has been interesting research into what is somewhat exotically termed the ecology of computation (Hubermann 1988). In this philosophy, distributed compu-

tational systems are seen as models of social, economical, biological, and so on, organizations. The goal is to understand and build such computational ecologies. Three issues of concern in this regard are (1) the general issues underlying open systems (Hewitt 1985), (2) the actual implementations of distributed computation in computational ecologies, and (3) the design of appropriate languages for open systems. Agha's book is a must for gaining an appreciation of the problems to be solved to obtain these actual implementations and also deals with problems of designing languages for open systems. If you want to understand the more practical ideas of ecological computation, theory of games, market economies and underlying mechanisms, evolution and biological systems, and so on, there is a lot to gain in terms of conceptual clarity and computational fundamentals by starting with this elegant book. Agha also commented on this sociological aspect of the actor model elsewhere:

The term actor brings with it the image of an active entity, acting out its role in concert with others. Each actor can act only according to its script, which represents its view of the world. At this level, the actor paradigm stands in sharp contrast to the logicist approach, which not only postulates the existence of a unique reality, but commits us to representing our knowledge in terms of a consistent collection of information.

The book has several additional strengths that make it a useful source for people working on distributed computation. Problems of distributed computing such as divergence, deadlock, and mutual exclusion are treated. There is also an excellent chapter on abstraction and compositionality. An excellent glossary of actor terms is included as well as a list of references and a useful index. Although the book is a revised version of the author's

Readings from AI Magazine

The First Five Years: 1980-1985

Edited with a Preface by Robert Englemore

AAAI is pleased to announce publication of Readings from AI Magazine, the complete collection of all the articles that appeared during AI Magazine's first five years. Within this 650-page indexed volume, you will find articles on AI written by the foremost practitioners in the field—articles that earned AI Magazine the title “journal of record for the artificial intelligence community.” This collection of classics from the premier publication devoted to the entire field of artificial intelligence is available in one large, paperbound desktop reference.

Subjects Include:

- | | | |
|-----------------------------------|---------------------------------------|----------------------------------|
| ■ Automatic Programming | ■ Distributed Artificial Intelligence | ■ Games |
| ■ Infrastructure | ■ Learning | ■ Natural Language Understanding |
| ■ Problem Solving | ■ Robotics | ■ Education |
| ■ General Artificial Intelligence | ■ Knowledge Acquisition | ■ Legal Issues |
| ■ Object Oriented Programming | ■ Programming Language | ■ Simulation |
| ■ Technology Transfer | ■ Discovery | ■ Expert Systems |
| ■ Historical Perspectives | ■ Knowledge Representation | ■ Logic |
| ■ Partial Evaluation | ■ Reasoning with Uncertainty | ■ Computer Architectures |

\$74.95 plus \$2 postage and handling. 650 pages, illus., appendix, index ISBN 0-929280-01-6.

Send prepaid orders to The MIT Press, 55 Hayward Street, Cambridge, MA 02142.

doctoral dissertation, it doesn't carry the usual blemish associated with theses in general; that is, it is not written in the stream-of-consciousness mode.

There is a well-known saying: Any theory that can be put in a nutshell belongs there. Well, running this risk, I'll say that in a nutshell, the actor approach is about the future of computing.

Postscript: A new (and probably the definitive) source has joined the ranks of actor literature: *Concurrent Systems for Knowledge Processing: An Actor Perspective*, edited by C. Hewitt and G. Agha (Cambridge, Mass., The MIT Press, 1989). This book brings together over 20 important contributions on the actor concept and its applications. Although I haven't seen this book, given my positive views about the book under review, I believe that this new text will also be a pleasure to read.

Varol Akman is an assistant professor in the Department of Computer Engineering and Information Science, Bilkent University, Ankara, Turkey. His current research concentrates on various theoretical aspects of AI, especially from the angle of intelligent computer-aided design systems, for example, commonsense reasoning, naive

physics, knowledge representation, and mathematical logic.

References

Agha, G. Foundational Issues in Concurrent Computing, Dept. of Computer Science, Yale Univ. Unpublished. Also in Foundational Issues in Concurrent Computing. In Proceedings of the ACM/SIGPLAN Workshop on Object-Based Concurrent Programming, SIGPLAN Notices, eds. G. Agha, P. Wagner, and A. Yonezawa. New York: Association of Computing Machinery.

Hewitt, C. 1977. Viewing Control Structures as Patterns of Passing Messages. *Artificial Intelligence* 8(3): 323-364.

Hewitt, C. 1985. The Challenge of Open Systems. *Byte* 10(4): 223-242.

Hubermann, B., ed. 1988. *The Ecology of Computation: Studies in Computer Science and Artificial Intelligence*, volume 2. The Hague, The Netherlands: Elsevier.

Simple Minds

Lee A. Gladwin

Of what are minds made? Internal mental representations? Matter? In this provocative and engaging work (*Simple Minds*, Cambridge, Mass.: The

MIT Press, 1989, 266 pages, \$25.00, ISBN 0-262-12140-9), Dan Lloyd seeks to provide answers that will bridge the gap between computational and connectionist models of the mind. He notes that naturalism assumes that the human brain and the brains of some animals are organs of thought or representation, but what is it about these brains that makes them thoughtful? Additionally, just what sorts of physical systems are sufficient to embody the representations typical of thinking? Finally, could a nonbiological system—a computer perhaps—think? With these fundamental questions, Lloyd introduces the central question to be answered by *Simple Minds*: How does the brain comprise the mind?

Lloyd seeks to outline a theory of representing systems and lay the foundations of a natural philosophy of mind. Such a reductionist theory, he observes, must explain how representations are made out of nonrepresentational parts. Only by this means can we solve the problem of mind and brain and the open question of the convergence of psychology and neuroscience.

Simple Minds is divided into three parts. Part 1, Representation from the

Bottom Up, describes what a theory of material representation must do and looks at two current approaches (internalism and externalism). It then proceeds to develop a dialectic theory of representation based on thought experiments involving a simple device called Squint that integrates input from multiple channels and can come to be a representative device. Part 2, Interpreting Neural Networks, applies the theory to biological and simulated neural networks. Part 3, From Simple Minds to Human Minds, extends the theory to language use, consciousness, and reasoning.

Representation, according to Lloyd, depends on relations involving the world outside the representing system. He establishes three criteria for basic representation: multiple channels, convergence, and uptake. The first requires the presence of multiple information channels leading to the representation; the second requires the installation of an 'and' gate at the confluence of the channels to act as an indicator when conjunctions of antecedent conditions occur. To work effectively, the convergence condition further requires that the multiple-channel device be uniquely situated so that sometimes one event (mediated by various channels) is sufficient for its activation. In this way, channels converge on a single event. The uptake criterion requires that the event representation have the capacity to cause either another representation or a salient behavioral event.

The introduction of these concepts does treble violence to Occam's Razor, adding unnecessary levels of complexity to an already difficult topic. Most of the same ground is better covered by Gary Lynch in *Synapses, Circuits, and the Beginnings of Memory* (Cambridge, Mass.: The MIT Press, 1986), with far greater attention to existing neurobiological concepts and studies. There are also numerous fine illustrations in Lynch to aid the reader in understanding the neurological foundations of representation.

Having set forth his dialectic theory of representation in part 1, Lloyd notes that his theory will illuminate the human mind only if it can encompass the complexity of its subject. Bridging the gap between the brain, neural networks, and computational models is the task of the remaining two parts of the book.

Three forms of representation are

noted with regard to neural networks: local; distributed; and, what Lloyd calls, featural. In local representation, each unit is dedicated to one content—understood as a proposition or hypothesis—and the activation of this unit indicates how committed the system is to the truth of the representation. Distributed and featural representation are nonlocal, with individual units participating in a variety of representations. In distributed representation, however, no single unit is exclusively associated with either a specific input or a specific output, and all the active units must be interpreted together. Finally, in featural representation, every unit is dedicated to the representation of a feature, and this interpretation remains constant over time. Higher cognitive processes can thus be implemented by lower-level units. Again, there are clearer accounts of feature detection, perception, and representation in Lynch and others.

With this foundation, the reader is prepared to cross over the metarepresentational bridge between the connectionist and computational models of the mind, which occurs in chapter 6, "The Language of Thought." First, Lloyd summarizes Jerry Fodor's hypothesis that the inner representational system uses a language-like medium: Thought is the manipulation of internal symbols with cognitively relevant constituent structure. In particular, Lloyd examines the concept of *mentalese*, the unconscious, innate language of thought.

He then introduces his own concept of *metarepresentation*, a representation inheriting content from another representation. Lloyd offers the familiar example of newspaper photos that are reconstructed from a digitized representation of a photograph and concludes that what makes representation is not the resemblance of features shared by representation and metarepresentation but rather the representation of original representational features by the metarepresentation.

Lloyd's bridge struck me as promising in design but difficult to cross. Rather large leaps are required to get from one part of the bridge to the next. The reader jumps from a well-described, easily visualized theory in part 1 to artificial-biological networks in part 2 to metarepresentation in part 3. There is scant but suggestive evidence provided that electrochemi-

cal activity in the immense pharmacy of the brain is transformed by some alchemy into the language of thought, but the reader is left far from seeing how the hardware (or wetware) supports the psychological software.

A more orderly and better researched bridging theory can be found in Jean-Pierre Changeux's *Neuronal Man: The Biology of Mind* (New York: Pantheon Books, 1985). Changeux, a neurobiologist, provides a thorough and readable description of the brain's physiology and then presents support for his own thesis that the brain contains representations of the outside world in the anatomical organization of its cortex and is capable of building representations of its own and using them in its computations. Much of Changeux's description of how multiple representations are coordinated within the visual region of the cortex, for example, would have proved useful to Lloyd in his discussions of featural representation and metarepresentation.

Similar comments could be made with reference to Donald O. Hebb's cell assembly theory. First presented in *The Organization of Behavior* (New York: John Wiley and Sons, 1949), it was expanded and further documented in *Essay on Mind* (Hillsdale, N.J.: Lawrence Erlbaum Associates, 1980). In the latter work, Hebb described the development of visual perception based on cell assemblies and sub-assemblies in the striate and peristriate cortex. His description of first- and second-order assemblies accords well with Lloyd's metarepresentation:

When a group of assemblies are repeatedly activated, simultaneously or in sequence, those cortical cells that are regularly active following this primary activity can themselves become organized as a superordinate (second-order) assembly. The activity of this assembly then is representative of the combination of events that individually organized the first-order assemblies.

It is these higher-order cell assemblies that Hebb believed were the foundation of abstract thought.

Although Lloyd's metarepresentational bridge is seemingly sound in design, it would have gained considerable support from a broader reading of the neurobiological literature and the inclusion of references to the discoveries, concepts, models, and

theories already available. A new bridge can be built without inventing new materials if existing ones are already available.

In some future golden age of science, Lloyd notes, it might be possible to test theories of the mind against the details of the brain by directly observing inner representations in their physical form. The results might vindicate a naturalism about the mind and show that the representations posited by cognitive psychologists or the beliefs and desires alluded to in folk psychology will turn out to be real, physical entities in the world. Alternatively, on this last day when all mysteries will be revealed, we might all stop and stammer, "So that's how it works!"

Lee A. Gladwin received his degree from Carnegie-Mellon University. His dissertation was on historical problem solving. He is currently interested in combining hypermedia with neural networks.

Machine Intelligence: A Critique of Arguments against the Possibility of Artificial Intelligence

Achim Hoffman

Stuart Goldkind's book *Machines and Intelligence: A Critique of Arguments against the Possibility of Artificial Intelligence* (Westport, Conn.: Greenwood Press, 1987, 139 pages, \$29.95, ISBN 0-313-25450-8) considers and criticizes a collection of arguments against the possibility of AI. It does not claim to contain an exhaustive collection of such arguments or even to cover the most important arguments raised in the mainly philosophical literature of the past few decades. Rather, as the author points out in the preface, the book is intended to examine the arguments for and against the possibility of AI to give the reader more accurate notions of terms such as intelligence and cognition as well as make the reader more aware of the implicit suppositions involved in the debate around the possibility of AI.

To pursue this issue, the book starts by examining the well-known Turing test as well as discussing some arguments (for example, the argument of consciousness) against the possibility of AI that Turing himself addressed

in his paper "Computing Machinery and Intelligence" (*Mind* 59, 1950). The author discusses the validity of the Turing test and concludes that there are no means for showing that in contrast to a human being, a machine that passes the test is not able to think or understand. Thereby prepared, in the chapters that follow, the reader finds a more or less compact presentation of arguments against the possibility of AI that have been established by philosophers in recent decades.

The book reviews and critiques the major lines of argument in Dreyfus's book *What Computers Can't Do* (New York: Harper & Row, 1972). To each line of argument, the author responds by pointing out logical gaps in Dreyfus's argument. Although Goldkind shows that Dreyfus's arguments do not prove the impossibility of AI, he nevertheless emphasizes the value of Dreyfus's reasoning in giving profound insight into the problems of constructing intelligent machines as well as an understanding of the nature of human intelligence.

Goldkind then carries on a dialogue about the possibility of machines making mistakes, as raised in Turing's paper. Here, in contrast to the other chapters, the author mainly follows his own thoughts. Goldkind's case is that for a machine to err, it must have intentions, which is the topic of the remainder of the book.

Goldkind begins the examination of this topic with a discussion of Richard Taylor's 1966 monograph *Action and Purpose* (Englewood Cliffs, N.J.: Prentice-Hall) in which Taylor argues that machines are incapable of purposeful behavior. Goldkind points out the restricted generality of Taylor's arguments to refute his line of thinking. Goldkind claims that only machines that are capable of knowledge are able to behave purposefully. Unfortunately, Goldkind does not exactly explain what he means by knowledge. As a result, the reader can neither decide about the conclusiveness of the author's claim nor survey its implications.

In the final chapter, parts of Norman Malcolm's 1968 article "The Conceivability of Mechanism" (*The Philosophical Review*, Vol. 77, pp. 45-72) are reviewed. Malcolm argues for an incompatibility of the causal explanation of behavior on the one hand and a purposive explanation of the same behavior on the other.

Goldkind uses elegant arguments to reject Malcolm's conclusions. Nevertheless, in examining an argument of Malcolm's, Goldkind fails to mention its lack of logical conclusiveness, instead attacking its premises, which is a bit confusing. The author closes his last chapter with a repetition of the claim that a machine must be capable of knowledge to act purposefully, that is, to act at all.

I found the critiques of Taylor's and Malcolm's positions overly long and boring. It would have been better to shorten the discussion in these last two chapters and more deeply explore key terms in the critique.

In summary, this book discusses a couple of philosophical issues concerning the possibility of AI. The first two chapters, covering about one half of the book, provide an especially good insight into common lines of argument against the possibility of AI and their implicit suppositions. In each chapter, the author tries to weaken the arguments against the possibility of AI by pointing out certain unproven claims within each one. In making implicit assumptions explicit, the book certainly contributes to the understanding of AI.

Achim G. Hoffman is a member of the academic staff of the Computer Science Department at Technische Universität Berlin. His current research interests are machine learning and the application of AI to very large-scale integrated design automation.

**This publication
is available
in microform
from University
Microfilms
International.**



Call toll-free 800-521-3044. In Michigan,
Alaska and Hawaii call collect 313-761-4700. Or
mail inquiry to: University Microfilms International,
300 North Zeeb Road, Ann Arbor, MI 48106

For free information, circle no. 174