

A (Very) Brief History of Artificial Intelligence

Bruce G. Buchanan

■ In this brief history, the beginnings of artificial intelligence are traced to philosophy, fiction, and imagination. Early inventions in electronics, engineering, and many other disciplines have influenced AI. Some early milestones include work in problems solving which included basic work in learning, knowledge representation, and inference as well as demonstration programs in language understanding, translation, theorem proving, associative memory, and knowledge-based systems. The article ends with a brief examination of influential organizations and current issues facing the field.

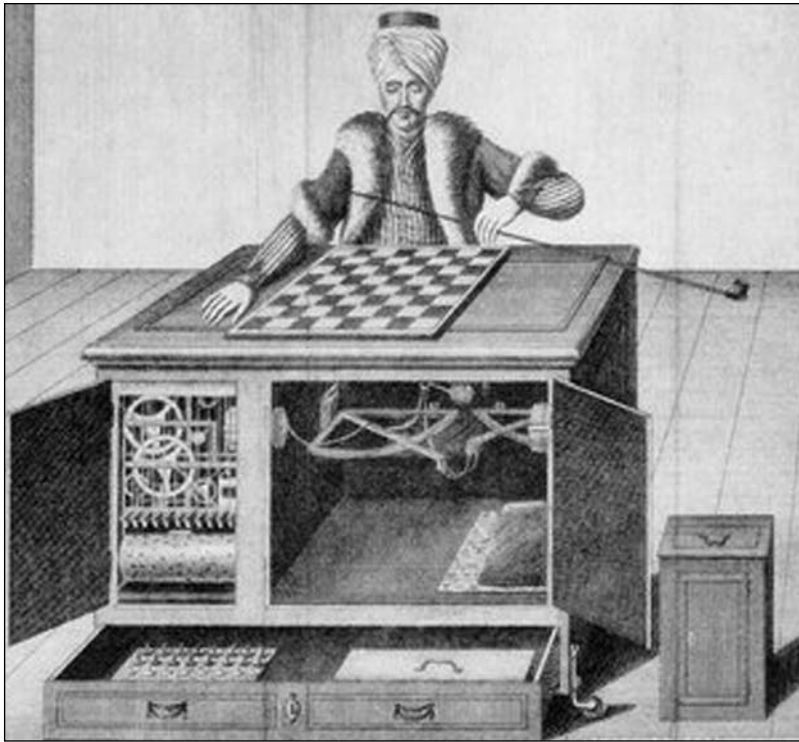
The history of AI is a history of fantasies, possibilities, demonstrations, and promise. Ever since Homer wrote of mechanical “tripods” waiting on the gods at dinner, imagined mechanical assistants have been a part of our culture. However, only in the last half century have we, the AI community, been able to build experimental machines that test hypotheses about the mechanisms of thought and intelligent behavior and thereby demonstrate mechanisms that formerly existed only as theoretical possibilities. Although achieving full-blown artificial intelligence remains in the future, we must maintain the ongoing dialogue about the implications of realizing the promise.¹

Philosophers have floated the possibility of intelligent machines as a literary device to help us define what it means to be human. René Descartes, for example, seems to have been more interested in “mechanical man” as a metaphor than as a possibility. Gottfried Wilhelm Leibniz, on the other hand, seemed to see the possibility of mechanical reasoning devices using rules of logic to settle disputes. Both Leib-

niz and Blaise Pascal designed calculating machines that mechanized arithmetic, which had hitherto been the province of learned men called “calculators,” but they never made the claim that the devices could think. Etienne Bonnot, Abbé de Condillac used the metaphor of a statue into whose head we poured nuggets of knowledge, asking at what point it would know enough to appear to be intelligent.

Science fiction writers have used the possibility of intelligent machines to advance the fantasy of intelligent nonhumans, as well as to make us think about our own human characteristics. Jules Verne in the nineteenth century and Isaac Asimov in the twentieth are the best known, but there have been many others including L. Frank Baum, who gave us the *Wizard of Oz*. Baum wrote of several robots and described the mechanical man Tiktok in 1907, for example, as an “Extra-Responsive, Thought-Creating, Perfect-Talking Mechanical Man ... Thinks, Speaks, Acts, and Does Everything but Live.” These writers have inspired many AI researchers.

Robots, and artificially created beings such as the Golem in Jewish tradition and Mary Shelly’s *Frankenstein*, have always captured the public’s imagination, in part by playing on our fears. Mechanical animals and dolls—including a mechanical trumpeter for which Ludwig van Beethoven wrote a fanfare—were actually built from clockwork mechanisms in the seventeenth century. Although they were obviously limited in their performance and were intended more as curiosities than as demonstrations of thinking, they provided some initial credibility to mechanistic views of behavior and to the idea that such behavior need not be feared. As the industrial world became more mechanized, machinery became more sophisticated



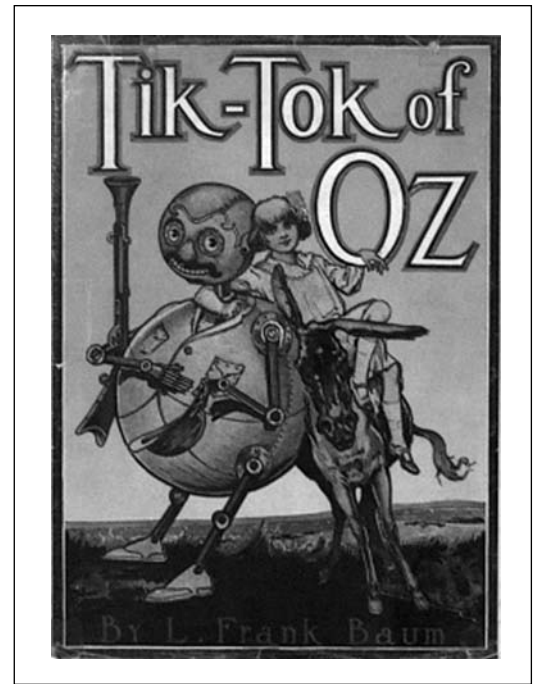
The Turk, from a 1789 Engraving by Freiherr Joseph Friedrich zu Racknitz.

and more commonplace. But it was still essentially clockwork.

Chess is quite obviously an enterprise that requires thought. It is not too surprising, then, that chess-playing machines of the eighteenth and nineteenth centuries, most notably “the Turk,” were exhibited as intelligent machines and even fooled some people into believing the machines were playing autonomously. Samuel L. Clemens (“Mark Twain”) wrote in a newspaper column, for instance, that the Turk must be a machine because it played so well! Chess was widely used as a vehicle for studying inference and representation mechanisms in the early decades of AI work. (A major milestone was reached when the Deep Blue program defeated the world chess champion, Gary Kasparov, in 1997 [McCorduck 2004].)

With early twentieth century inventions in electronics and the post-World War II rise of modern computers in Alan Turing’s laboratory in Manchester, the Moore School at Penn, Howard Aiken’s laboratory at Harvard, the IBM and Bell Laboratories, and others, possibilities have given over to demonstrations. As a result of their awesome calculating power, computers in the 1940s were frequently referred to as “giant brains.”

Although robots have always been part of the public’s perception of intelligent computers, early robotics efforts had more to do with



Baum Described Tik-Tok as an “Extra-Responsive, Thought-Creating, Perfect-Talking Mechanical Man ... Thinks, Speaks, Acts, and Does Everything but Live.”

mechanical engineering than with intelligent control. Recently, though, robots have become powerful vehicles for testing our ideas about intelligent behavior. Moreover, giving robots enough common knowledge about everyday objects to function in a human environment has become a daunting task. It is painfully obvious, for example, when a moving robot cannot distinguish a stairwell from a shadow. Nevertheless, some of the most resounding successes of AI planning and perception methods are in NASA’s autonomous vehicles in space. DARPA’s grand challenge for autonomous vehicles was recently won by a Stanford team, with 5 of 23 vehicles completing the 131.2-mile course.²

But AI is not just about robots. It is also about understanding the nature of intelligent thought and action using computers as experimental devices. By 1944, for example, Herb Simon had laid the basis for the information-processing, symbol-manipulation theory of psychology:

“Any rational decision may be viewed as a conclusion reached from certain premises.... The behavior of a rational person can be controlled, therefore, if the value and factual premises upon which he bases his decisions are specified for him.” (Quoted in the Appendix to Newell & Simon [1972]).

AI in its formative years was influenced by



On October 8, 2005, the Stanford Racing Team's Autonomous Robotic Car, Stanley, Won the Defense Advanced Research Projects Agency's (DARPA) Grand Challenge. The car traversed the off-road desert course southwest of Las Vegas in a little less than seven hours.

Photo courtesy, DARPA.



Mars Rover.

Photo Courtesy, NASA



Herb Simon.



John McCarthy.

ideas from many disciplines. These came from people working in engineering (such as Norbert Wiener's work on cybernetics, which includes feedback and control), biology (for example, W. Ross Ashby and Warren McCulloch and Walter Pitts's work on neural networks in simple organisms), experimental psychology (see Newell and Simon [1972]), communication theory (for example, Claude Shannon's theoretical work), game theory (notably by John Von Neumann and Oskar Morgenstern), mathematics and statistics (for example, Irving J. Good), logic and philosophy (for example, Alan Turing, Alonzo Church, and Carl Hempel), and linguistics (such as Noam Chomsky's work on grammar). These lines of work made their mark and continue to be felt, and our collective debt to them is considerable. But having assimilated much, AI has grown beyond them and has, in turn, occasionally influenced them.

Only in the last half century have we had computational devices and programming languages powerful enough to build experimental tests of ideas about what intelligence is. Turing's 1950 seminal paper in the philosophy journal *Mind* is a major turning point in the history of AI. The paper crystallizes ideas about the possibility of programming an electronic computer to behave intelligently, including a description of the landmark imitation game that we know as Turing's Test. Vannevar Bush's 1945 paper in the *Atlantic Monthly* lays out a prescient vision of possibilities, but Turing was actually writing programs for a computer—for example, to play chess, as laid out in Claude Elwood Shannon's 1950 proposal.

Early programs were necessarily limited in scope by the size and speed of memory and processors and by the relative clumsiness of the early operating systems and languages. (Memory management, for example, was the programmer's problem until the invention of garbage collection.) Symbol manipulation languages such as Lisp, IPL, and POP and time sharing systems—on top of hardware advances in both processors and memory—gave programmers new power in the 1950s and 1960s. Nevertheless, there were numerous impressive demonstrations of programs actually solving problems that only intelligent people had previously been able to solve.

While early conference proceedings contain descriptions of many of these programs, the first book collecting descriptions of working AI programs was Edward Feigenbaum and Julian Feldman's 1963 book, *Computers and Thought*.

Arthur Samuel's checker-playing program, described in that collection but written in the 1950s, was a tour-de-force given both the limi-

tations of the IBM 704 hardware for which the program was written as a checkout test and the limitations of the assembly language in which it was written. Checker playing requires modest intelligence to understand and considerable intelligence to master. Samuel's program (since outperformed by the Chinook program) is all the more impressive because the program learned through experience to improve its own checker-playing ability—from playing human opponents and playing against other computers. Whenever we try to identify what lies at the core of intelligence, learning is sure to be mentioned (see, for example, Marvin Minsky's 1961 paper "Steps Toward Artificial Intelligence.")

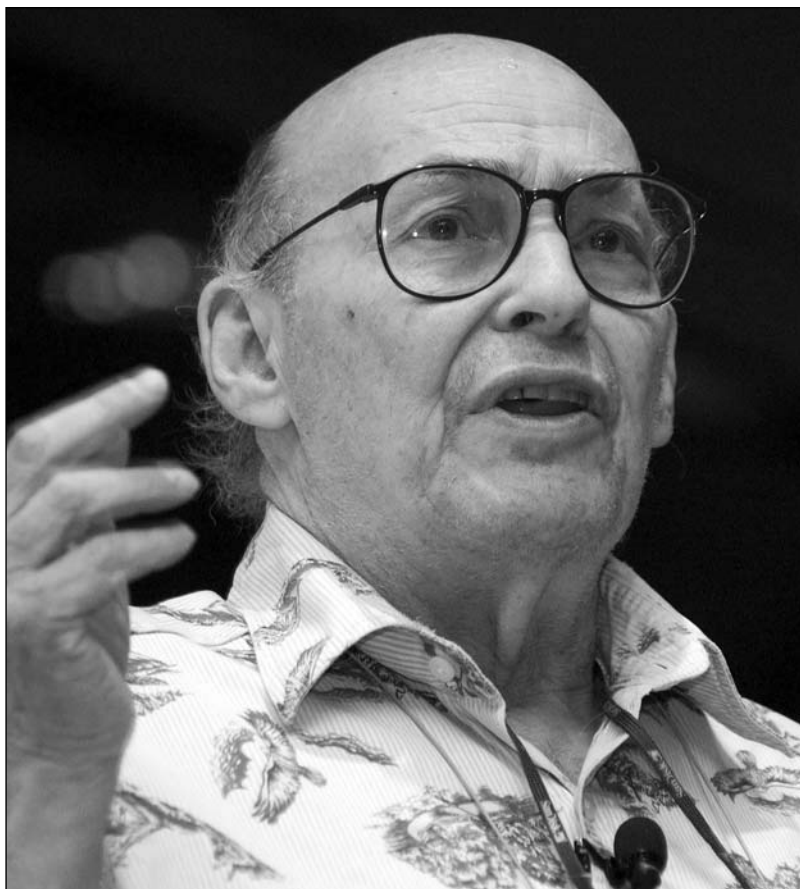
Allen Newell, J. Clifford Shaw, and Herb Simon were also writing programs in the 1950s that were ahead of their time in vision but limited by the tools. Their LT program was another early tour-de-force, startling the world with a computer that could invent proofs of logic theorems—which unquestionably requires creativity as well as intelligence. It was demonstrated at the 1956 Dartmouth conference on artificial intelligence, the meeting that gave AI its name.

Newell and Simon (1972) acknowledge the convincingness of Oliver Selfridge's early demonstration of a symbol-manipulation program for pattern recognition (see Feigenbaum and Feldman [1963]). Selfridge's work on learning and a multiagent approach to problem solving (later known as blackboards), plus the work of others in the early 1950s, were also impressive demonstrations of the power of heuristics. The early demonstrations established a fundamental principle of AI to which Simon gave the name "satisficing":

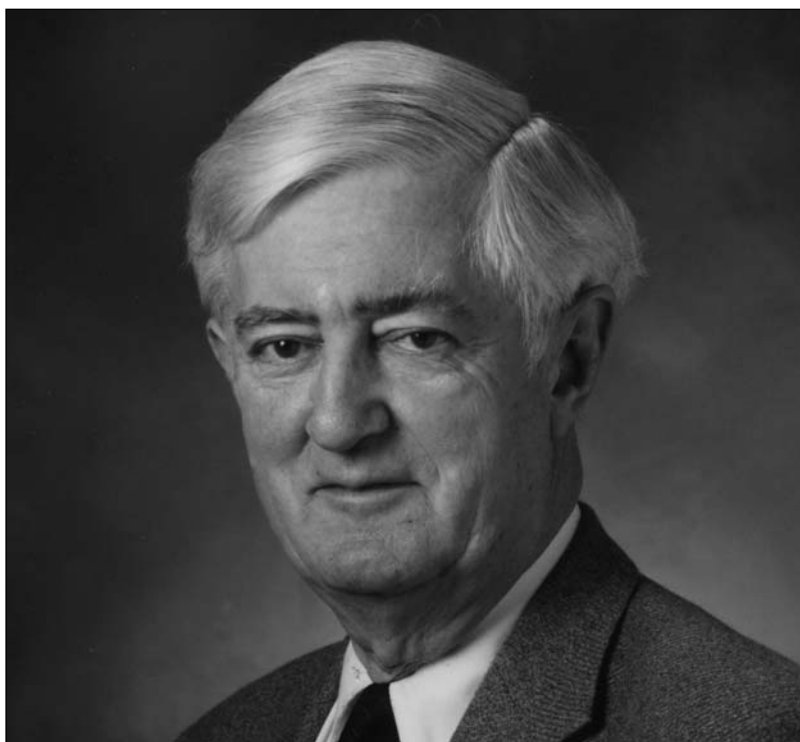
In the absence of an effective method guaranteeing the solution to a problem in a reasonable time, heuristics may guide a decision maker to a very satisfactory, if not necessarily optimal, solution. (See also Polya [1945].)

Minsky (1968) summarized much of the work in the first decade or so after 1950:

"The most central idea of the pre-1962 period was that of finding heuristic devices to control the breadth of a *trial-and-error search*. A close second preoccupation was with finding effective techniques for *learning*. In the post-1962 era the concern became less with "learning" and more with the problem of representation of knowledge (however acquired) and with the related problem of breaking through the formality and narrowness of the older systems. The problem of heuristic search efficiency remains as an underlying constraint, but it is no longer the problem one thinks about, for we are now immersed in more sophisticated subproblems, e.g., the representation and modification of plans" (Minsky 1968, p. 9).



Marvin Minsky.



Oliver Selfridge.



Donald Michie.



*Photograph Courtesy, National Library of Medicine
The Original Dendral Team, Twenty-Five Years Later.*

Minsky's own work on network representations of knowledge in frames and what he calls the "society of minds" has directed much research since then. Knowledge representation—both the formal and informal aspects—has become a cornerstone of every AI program. John McCarthy's important 1958 paper, "Programs with Common Sense" (reprinted in Minsky [1968]), makes the case for a declarative knowledge representation that can be manipulated easily. McCarthy has been an advocate for using formal representations, in particular extensions to predicate logic, ever since. Research by McCarthy and many others on nonmonotonic reasoning and default reasoning, as in planning under changing conditions, gives us important insights into what is required for intelligent action and defines much of the formal theory of AI.

GPS (by Newell, Shaw, and Simon) and much of the other early work was motivated by psychologists' questions and experimental methods (Newell and Simon 1972). Feigenbaum's EPAM, completed in 1959, for example, explored associative memory and forgetting in a program that replicated the behavior of subjects in psychology experiments (Feigenbaum and Feldman 1963). Other early programs at Carnegie Mellon University (then Carnegie Tech) deliberately attempted to replicate the reasoning steps, including the mistakes, taken by human problem solvers in puzzles such as cryptarithmic and selecting stocks for investment portfolios. Production systems, and subsequent rule-based systems, were originally conceived as simulations of human manipulations of symbols in long-term and short-term memory. Donald Waterman's 1970 dissertation at Stanford used a production system to play draw poker, and another program to learn how to play better.

Thomas Evans's 1963 thesis on solving analogy problems of the sort given on standardized IQ tests was the first to explore analogical reasoning with a running program. James Slagle's dissertation program used collections of heuristics to solve symbolic integration problems from freshman calculus. Other impressive demonstrations coming out of dissertation work at MIT in the early 1960s by Danny Bobrow, Bert Raphael, Ross Quillian, and Fischer Black are described in Minsky's collection, *Semantic Information Processing* (Minsky 1968).

Language understanding and translation were at first thought to be straightforward, given the power of computers to store and retrieve words and phrases in massive dictionaries. Some comical examples of failures of the table lookup approach to translation provided critics

with enough ammunition to stop funding on machine translation for many years. Danny Bobrow's work showed that computers could use the limited context of algebra word problems to understand them well enough to solve problems that would challenge many adults. Additional work by Robert F. Simmons, Robert Lindsay, Roger Schank, and others similarly showed that understanding—even some translation—was achievable in limited domains. Although the simple look-up methods originally proposed for translation did not scale up, recent advances in language understanding and generation have moved us considerably closer to having conversant nonhuman assistants. Commercial systems for translation, text understanding, and speech understanding now draw on considerable understanding of semantics and context as well as syntax.

Another turning point came with the development of knowledge-based systems in the 1960s and early 1970s. Ira Goldstein and Seymour Papert (1977) described the demonstrations of the Dendral program (Lindsay et al. 1980) in the mid-1960s as a “paradigm shift” in AI toward knowledge-based systems. Prior to that, logical inference, and resolution theorem proving in particular, had been more prominent. Mycin (Buchanan and Shortliffe 1984) and the thousands of expert systems following it became visible demonstrations of the power of small amounts of knowledge to enable intelligent decision-making programs in numerous areas of importance. Although limited in scope, in part because of the effort to accumulate the requisite knowledge, their success in providing expert-level assistance reinforces the old adage that knowledge is power.

The 1960s were also a formative time for organizations supporting the enterprise of AI. The initial two major academic laboratories were at the Massachusetts Institute of Technology (MIT), and CMU (then Carnegie Tech, working with the Rand Corporation) with AI laboratories at Stanford and Edinburgh established soon after. Donald Michie, who had worked with Turing, organized one of the first, if not the first, annual conference series devoted to AI, the Machine Intelligence workshops first held in Edinburgh in 1965. About the same time, in the mid-1960s, the Association for Computing Machinery's Special Interest Group on Artificial Intelligence (ACM SIGART) began an early forum for people in disparate disciplines to share ideas about AI. The international conference organization, IJCAI, started its biannual series in 1969. AAAI grew out of these efforts and was formed in 1980 to provide annual conferences for the North American AI

AAAI Today

Founded in 1980, the American Association for Artificial Intelligence has expanded its service to the AI community far beyond the National Conference. Today, AAAI offers members and AI scientists a host of services and benefits:

- **The National Conference on Artificial Intelligence** promotes research in AI and scientific interchange among AI researchers, practitioners, and scientists and engineers in related disciplines. (www.aaai.org/Conferences/National/)
- **The Conference on Innovative Applications of Artificial Intelligence** highlights successful applications of AI technology; explores issues, methods, and lessons learned in the development and deployment of AI applications; and promotes an interchange of ideas between basic and applied AI. (www.aaai.org/Conferences/IAAI/)
- **The Artificial Intelligence and Interactive Digital Entertainment Conference** is intended to be the definitive point of interaction between entertainment software developers interested in AI and academic and industrial researchers. (www.aaai.org/Conferences/AIIDE/)
- **AAAI's Spring and Fall Symposia** (www.aaai.org/Symposia/) and **Workshops** (www.aaai.org/Workshops/) programs affords participants a smaller, more intimate setting where they can share ideas and learn from each other's AI research
- **AAAI's Digital Library** (www.aaai.org/Library/), (www.aaai.org/Resources/) and **Online Services** include a host of resources for the AI professional (including more than 12,000 papers), individuals with only a general interest in the field (www.aaai.org/AITopics/), as well as the professional press (www.aaai.org/Pressroom/).
- **AAAI Press**, in conjunction with The MIT Press, publishes selected books on all aspects of AI (www.aaai.org/Press/).
- **The AI Topics** web site gives students and professionals alike links to many online resources on AI (www.aaai.org/AITopics/).
- **AAAI Scholarships** benefit students and foster new programs, meetings, and other AI programs. AAAI also recognizes those who have made significant contributions to the science of AI and AAAI through an extensive awards program (www.aaai.org/Awards/).
- **AI Magazine**, called the “journal of record for artificial intelligence,” has been published internationally for 25 years (www.aaai.org/Magazine/).
- **AAAI's Sponsored Journals** program (www.aaai.org/Publications/Journals/) gives AAAI members discounts on many of the top AI journals.

www.aaai.org

With our successes in AI, however, come increased responsibility to consider the societal implications of technological success and educate decision makers and the general public so they can plan for them.

community. Many other countries have subsequently established similar organizations.

In the decades after the 1960s the demonstrations have become more impressive, and our ability to understand their mechanisms has grown. Considerable progress has been achieved in understanding common modes of reasoning that are not strictly deductive, such as case-based reasoning, analogy, induction, reasoning under uncertainty, and default reasoning. Contemporary research on intelligent agents and autonomous vehicles, among others, shows that many methods need to be integrated in successful systems.

There is still much to be learned. Knowledge representation and inference remain the two major categories of issues that need to be addressed, as they were in the early demonstrations. Ongoing research on learning, reasoning with diagrams, and integration of diverse methods and systems will likely drive the next generation of demonstrations.

With our successes in AI, however, come increased responsibility to consider the societal implications of technological success and educate decision makers and the general public so they can plan for them. The issues our critics raise must be taken seriously. These include job displacement, failures of autonomous machines, loss of privacy, and the issue we started with: the place of humans in the universe. On the other hand we do not want to give up the benefits that AI can bring, including less drudgery in the workplace, safer manufacturing and travel, increased security, and smarter decisions to preserve a habitable planet.

The fantasy of intelligent machines still lives even as we accumulate evidence of the complexity of intelligence. It lives in part because we are dreamers. The evidence from working programs and limited successes points not only to what we don't know but

also to some of the methods and mechanisms we can use to create artificial intelligence for real. However, we, like our counterparts in biology creating artificial life in the laboratory, must remain reverent of the phenomena we are trying to understand and replicate.

Acknowledgments

My thanks to Haym Hirsch, David Leake, Ed Feigenbaum, Jon Glick, and others who commented on early drafts. They, of course, bear no responsibility for errors.

Notes

1. An abbreviated history necessarily leaves out many key players and major milestones. My apologies to the many whose work is not mentioned here. The AAAI website and the books cited contain other accounts, filling in many of the gaps left here.
2. DARPA's support for AI research on fundamental questions as well as robotics has sustained much AI research in the U.S. for many decades.

References and Some Places to Start

- AAAI 2005. AI Topics Website. (www.aaai.org/aitopics/history). Menlo Park, CA: American Association for Artificial Intelligence.
- Blake, D. V., and Uttley, A. M., eds. 1959. *Mechanisation of Thought Processes: Proceedings of a Symposium Held at the National Physical Laboratory on 24th, 25th, 26th, and 27th November, 1958*. London: Her Majesty's Stationery Office.
- Bowden, B. V., ed. 1953. *Faster Than Thought: A Symposium on Digital Computing Machines*. New York: Pitman.
- Buchanan, B. G., and Shortliffe, E. H. 1984. *Rule-Based Expert Systems: The MYCIN Experiments of the Stanford Heuristic Programming Project*. Reading, MA: Addison-Wesley.
- Bush, V. 1945. As We May Think. *Atlantic Monthly* 176(7): 101.
- Cohen, J. 1966. *Human Robots in Myth and Science*. London: George Allen & Unwin.
- Feigenbaum, E.A., and Feldman, J. 1963.

Computers and Thought. New York: McGraw-Hill (reprinted by AAAI Press).

Goldstein, I., and Papert, S., 1977. Artificial Intelligence, Language and the Study of Knowledge. *Cognitive Science* 1(1).

Lindsay, R. K.; Buchanan, B. G.; Feigenbaum, E. A.; and Lederberg, J. 1980. *Applications of Artificial Intelligence for Chemical Inference: The DENDRAL Project*. New York: McGraw-Hill.

McCorduck, P. 2004. *Machines Who Think: Twenty-Fifth Anniversary Edition*. Natick, MA: A. K. Peters, Ltd.

Minsky, M. 1968. *Semantic Information Processing*. Cambridge, MA: MIT Press.

Minsky, M. 1961. Steps Toward Artificial Intelligence. In *Proceedings of the Institute of Radio Engineers* 49:8–30. New York: Institute of Radio Engineers. Reprinted in Feigenbaum and Feldman (1963).

Newell, A., and Simon, H. 1972. *Human Problem Solving*. Englewood Cliffs, NJ: Prentice-Hall.

Polya, G. 1945. *How To Solve It*. Princeton, NJ: Princeton University Press.

Samuel, A. L. 1959. Some Studies in Machine Learning Using the Game of Checkers. *IBM Journal of Research and Development* 3: 210–229. Reprinted in Feigenbaum and Feldman (1963).

Shannon, C. 1950. Programming a Digital Computer for Playing Chess. *Philosophy Magazine* 41: 356–375.

Turing, A. M. 1950. Computing Machinery and Intelligence. *Mind* 59: 433–460. Reprinted in Feigenbaum and Feldman (1963).

Winston, P. 1988. *Artificial Intelligence: An MIT Perspective*. Cambridge, MA: MIT Press.

Bruce G. Buchanan was a founding member of AAAI, secretary-treasurer from 1986–1992, and president from 1999–2001. He received a B.A. in mathematics from Ohio Wesleyan University (1961) and M.S. and Ph.D. degrees in philosophy from Michigan State University (1966). He is University Professor emeritus at the University of Pittsburgh, where he has joint appointments with the Departments of Computer Science, Philosophy, and Medicine and the Intelligent Systems Program. He is a fellow of the American Association for Artificial Intelligence (AAAI), a fellow of the American College of Medical Informatics, and a member of the National Academy of Science Institute of Medicine. His e-mail address is buchanan@cs.pitt.edu.