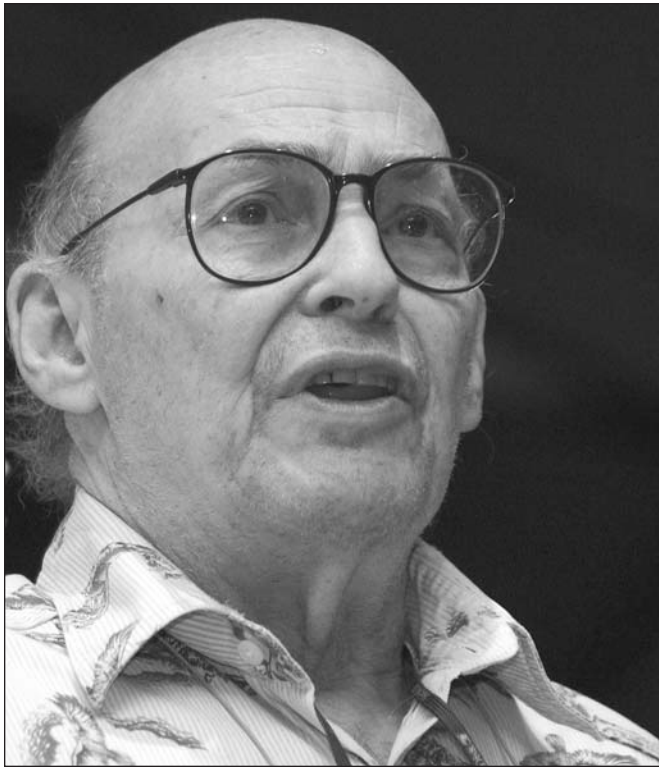




In Honor of Marvin Minsky's Contributions on his 80th Birthday

*Danny Hillis,
John McCarthy,
Tom M. Mitchell,
Erik T. Mueller,
Doug Riecken,
Aaron Sloman,
and Patrick Henry Winston*

■ Marvin Lee Minsky, a founder of the field of artificial intelligence and professor at MIT, celebrated his 80th birthday on August 9, 2007. This article seizes an opportune time to honor Marvin and his contributions and influence in artificial intelligence, science, and beyond. The article provides readers with some personal insights of Minsky from Danny Hillis, John McCarthy, Tom Mitchell, Erik Mueller, Doug Riecken, Aaron Sloman, and Patrick Henry Winston—all members of the AI community that Minsky helped to found. The article continues with a brief resume of Minsky's research, which spans an enormous range of fields. It concludes with a short biographical account of Minsky's personal history.

Personal Insights

Marvin Minsky has been so influential in so many fields that a brief summary cannot do him justice. He is one of the founders of artificial intelligence and robotics, and he has also made significant contributions in psychology and the theory of computing. He has received many awards and honors, including the A. M. Turing Award (1970), the Japan Prize (1990), the IJCAI Research Excellence Award (1991), and the Benjamin Franklin Medal (2001). He is a fellow of the American Academy of Arts and Sciences

and the Institute of Electrical and Electronic Engineers and a member of the U.S. National Academy of Sciences. He served as president of AAAI from 1981 to 1982.

Our goal in this article is to provide readers with personal insights of Marvin from members of our AI community along with some brief discussion of his contributions.

Danny Hillis

When I first arrived as a freshman at MIT, I was determined to meet the legendary Marvin Minsky. I hung out around the AI lab hoping to run into him, but he never seemed to be around. Finally, I heard a rumor that he was down in the basement, working every night on some new kind of computer. I went down one evening to take a look, and sure enough, there he was, surrounded by red-lined diagrams, wire-wrap guns, half-finished computer boards, and a few very busy assistants. I was too shy to introduce myself, so I just stayed out of the way and watched for a while, quietly examining the circuit drawings that were strewn around the room.

When I noticed an error in one of the drawings I bravely went up to Marvin and pointed it out. "Well, fix it," he said. So, I looked around and found the right tools and rewired the emerging computer to fix the problem. Then Marvin asked me to fix something else,

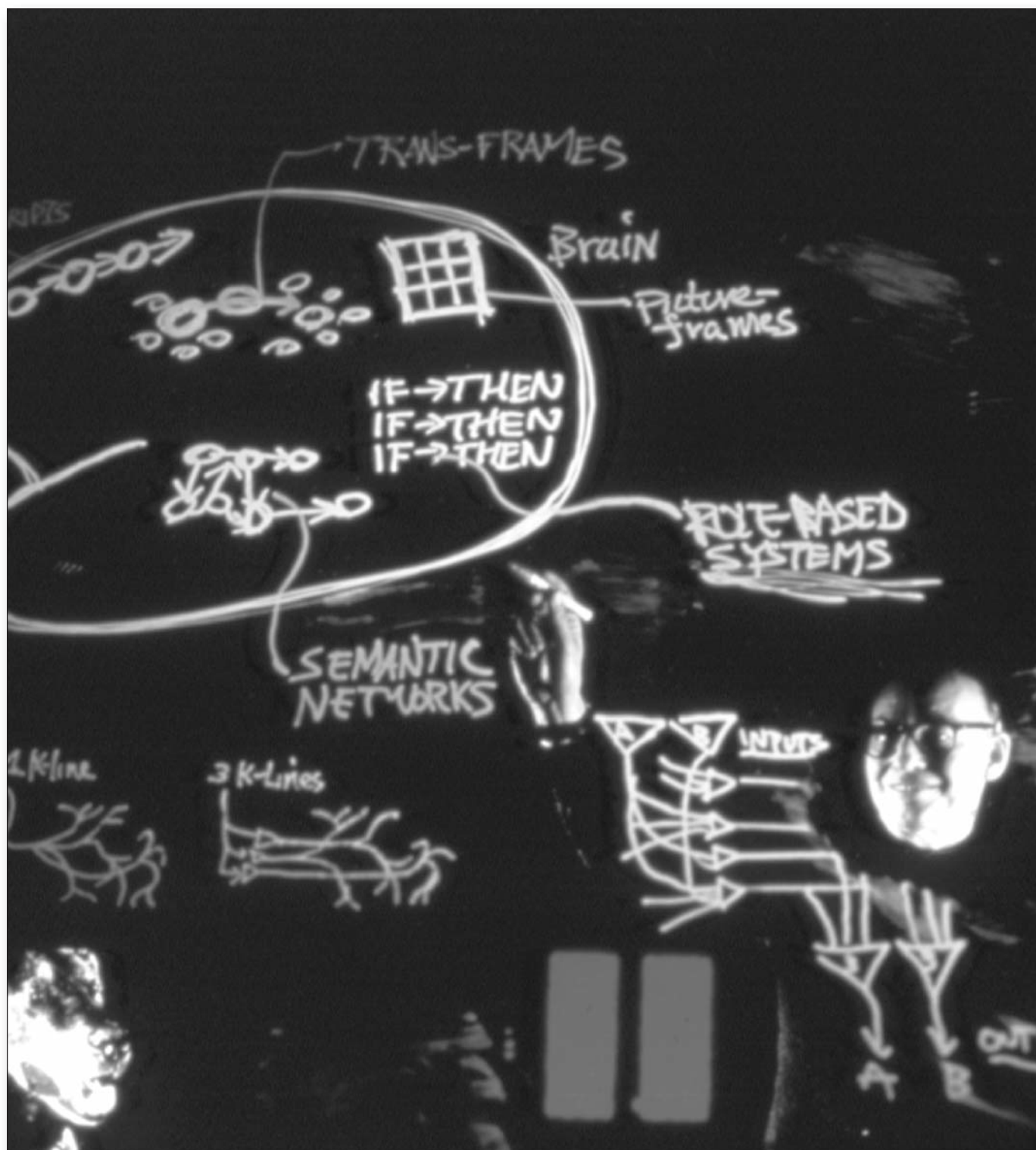


Figure 1. Marvin Minsky. Photo from S.P.A. Magazine, courtesy MIT Media Laboratory.

explaining to me how it was supposed to work. After a while, I guess he just assumed that I worked for him.

So, that's how I started working for Marvin

Minsky. Later, as I came to know Marvin as friend and mentor, I began to understand that this was a pretty normal interview process for him. Lots of people would hang around and

say clever things, but he always paid attention to the ones who actually did something useful.

John McCarthy

Marvin came to Princeton as a graduate student in mathematics in 1950, a year after me. We quickly realized that we both were interested in artificial intelligence, as we later came to call it. Unlike me, I think Marvin already had some definite ideas about how to achieve (human-level) AI. These then resulted in his design for the SNARC neural net learner and later led to his 1954 Ph.D. thesis. I had had ideas different from his but didn't consider them good enough and didn't come to a definite approach (through logic) until 1957. Neither approach has yet reached human level—nor have any of the others.

Tom M. Mitchell

Marvin Minsky has had a significant impact on my own AI research over the years, despite the fact that we have had relatively few opportunities for face-to-face discussions. How can a person with whom I've had little personal contact have such a strong influence? It's easy—I have been inspired by Marvin's style of out-of-the-box thinking and by his vision and aspirations for the field of AI. His work with his students on machine learning, integrated robot/language/perception/planning systems, and frame representations has helped shape the field and no doubt my own thinking about it. But for me personally, Marvin's strongest influence has been serving as an existence proof that we mortals are capable of great ideas and great aspirations. He has strengthened my own courage to tackle problems that I might otherwise not and to avoid spending too much time on incremental extensions of well-worn ideas. Marvin is a creative genius, full of ideas and willing to pursue them to lengths that others might not.

I recall one episode, when I served with Marvin in the 1980s on a technical advisory panel to NASA regarding a potential space station project. Marvin suggested the space station should be self-assembling so that it wouldn't require costly and dangerous involvement of astronauts. This was heresy to the NASA establishment, which was then justifying the space shuttle in part as the key to assembling a space station. But Marvin stuck to his guns and proceeded to lay out how he would design general fittings that snapped themselves together automatically, self-propelling parts, and pointing out that teleoperation of such smart devices would work fine despite multisecond communication delays with earthbound human controllers.

I believe Marvin was right that day, and NASA was wrong. But the point of this episode is that it exemplifies what I admire in Marvin's style. He avoided getting trapped in other people's framings of the question, he thought the problem through in his own way, and he came to a different solution. And he had fun doing it. To me, AI has always been the most creative, most ambitious, and most fun part of computer science. Marvin helped set that tone for AI, and in this way he has touched all of us.

Erik T. Mueller

The grand vision of human-level artificial intelligence that Marvin presents has been enormously important to the field and to my own research. At regular intervals I pick up and read a copy of *The Emotion Machine* to keep myself on track and prevent myself from losing sight of the goal—understanding how the mind works. Marvin has taught us many things about artificial intelligence research. Four that stand out are the following:

(1) We should think big and try to solve hard artificial intelligence problems like common-sense reasoning and story understanding. (2) We shouldn't attach too much importance to particular words, especially ones like *consciousness* and *emotion* that refer to complex sets of phenomena. We shouldn't rush to define them. Minsky says "words are designed to keep you from thinking." (3) If we can't find an elegant solution to an artificial intelligence problem, we shouldn't give up. An elegant solution is unlikely given what we know about the brain and the fact that it is a product of many years of evolution. (4) We should build large-scale systems consisting of many different methods that work together.

I deeply appreciate Marvin's support over the years for research on the unsolved problems of the field. Many members of our community working on building computer programs with common sense, emotions, and imagination have received the benefit of his constant advocacy for research on these aspects of human intelligence. His work is an inspiration to us all.

Doug Riecken

"What is your theory?" This was the first question from Marvin following our being introduced in 1986. I replied that I was working on a machine-learning music composing system called *Wolfgang*. Marvin asked if Wolfgang composed interesting music, to which I replied that it appears to, but there is an issue. You see, I programmed how Wolfgang was going to learn to learn (based on current machine-learning

ing techniques at that time). I then explained that Wolfgang should be put aside for a newer version system based on a different approach to bias learning and reasoning in a composing system; the newer Wolfgang system should develop biases for things it wants to learn based on system characteristics representing “its emotions and instincts.” That was the flash point for a friendship and scientific journey with Marvin that continues to be the opportunity of a lifetime.

My graduate studies with Marvin, and the years since, have enlightened me to Marvin’s insight, creativity, and perpetual curiosity; his unique thinking and comments are always many steps ahead, and as many colleagues have indicated, it is typical to obtain a eureka effect from a Minsky suggestion after several days of thought. While we recognize Marvin for his numerous contributions, it is most compelling to consider how he continues to point us as a community towards many of the most demanding questions. He has always commented “that you can not learn something until you learn it many different ways.” Thus it is essential we develop significant theories of architecture that will demonstrate this multi-learning-reasoning-representation ability. Marvin’s thinking and publications (such as his latest book, *The Emotion Machine: Commonsense Thinking, Artificial Intelligence, and the Future of the Human Mind* [Minsky 2006]) provide important insights by which to advance our field of research.

Perhaps one of Marvin’s most compelling contributions is his legacy as a teacher. Please visit Marvin’s website¹ and locate his list of students (located from the Minsky homepage through a link labeled “people”). It is an impressive, large list of many of the top AI researchers in our community. I am most grateful for Marvin’s friendship and his unique thinking and framing of the hard research problems that matter.

Aaron Sloman

Over 35 years ago, Marvin’s writings, such as the amazing “Steps” paper (Minsky 1961b), which contained many important ideas that illuminated a major breadth of later work in the field of AI, convinced me that the “design stance” should be used to address philosophical problems, that is, trying to understand how various mental processes could actually work: the basis of my 1978 book. Later, I learned from his frequent contributions to comp.ai.philosophy (which should be collected somewhere). In 1995 he agreed to join the Philosophical Encounter at IJCAI (with McCarthy), after

which our interactions became more frequent, including visits he paid to Birmingham and meetings at conferences and workshops.

Although I keep learning from Marvin, we do not agree about everything. His emphasis on “common sense” focuses on the uniquely human aspects of what minds do, which I think underplays some of our most important evolved abilities shared with some other animals, especially perception and manipulation of complex three-dimensional structures with complex causal interactions. Explaining common sense strikes me as probably the hardest unsolved problem in AI, with implications for many other aspects of intelligence. We have both worked on requirements for “layered” architectures for humanlike systems. Marvin’s work showed that I was not making enough distinctions between levels of processing, partly because I was more concerned with how humans fit into a space of possible designs, which includes knowing how we overlap with other species, and also how the architecture evolved and how it develops from infancy onwards. This difference leads to different ways of structuring discussion of architectures, but without substantial disagreement. His latest book is a major achievement, despite stretching the word *emotion* to refer to too many different control functions! But the important thing is not the label he uses but his myriad insights into what we are and, above all, how we work. We have reached some similar conclusions by different routes that bring out insights that I had missed, for instance in his paper “Interior Grounding, Reflection, and Self-Consciousness” (Minsky 2005) arguing that apparently simple sensations must be “complex reflective activities.” Finally, his website is full of treasures, for example the extraordinary BBC Three Music Interview in 2004.

Patrick Henry Winston

In my first term or two of graduate school, all I really knew was that I wanted to understand thinking, so I stumbled around studying all sorts of subjects, including neuroanatomy, psychology, control, information theory, and operations research. Then, one day I wandered into one of Marvin’s lectures, and I felt, for the first time, that I was listening to a genius thinking out loud. By the time the lecture was over, the die was cast, and my career objectives were settled, once and for all. He said things then that I’m still thinking about now.

A few months later, I scribbled out a thesis proposal vaguely aimed at dealing with analyzing scene descriptions and building on Adolfo Guzman’s pioneering work on understanding



Photo courtesy Jim Hendler.

Minsky During a Break at AAI-05 in Pittsburgh, Pennsylvania.

line drawings. In a footnote, I mumbled something about learning from description differences. When Marvin thumbed through the proposal, the footnote caught his eye, and turning to Seymour Papert, he said, "Well, this learning hack looks interesting." In that few seconds, my thesis proposal was reduced to the 20 words in the footnote.

Many years later, Danny Hillis came into my office, and as we talked, we noted that we both had many experiences in which Marvin fashioned good ideas from our not-so-promising and confused thinking. The most extreme, but common scenario goes like this: You think you have a great idea, so you try to tell Marvin about it. He has a very short attention span, so he leaps ahead and guesses what the idea is. Invariably, his guess is much better than the idea you were trying to explain. Danny pointed out that something similar might go on when we talk to ourselves. The words and phrases provide access to ideas that, when expressed in words and phrases, provide access to still more ideas, sort of like Marvin's K-line activation cycle.

Eventually, I have come to know a lot of geniuses, but Marvin is the only genius I know who is so smart it's scary. I worry that it is fool-

hardy to disagree with him, as I do sometimes, but then I remind myself that if we switched sides in any argument I would still get crushed with some refulgent insight, expressed concisely and with astonishing clarity. Anyway, when a computational account of intelligence is finally fully worked out, my bet is that Marvin's ideas and those he champions will be at center stage and everything else will seem like epicycles.

Minsky's Research and Contributions

Minsky's research spans an enormous range of fields, including mathematics, computer science, the theory of computation, neural networks, artificial intelligence, robotics, commonsense reasoning, natural language processing, and psychology.

For his bachelor's thesis in mathematics at Harvard University, Marvin used the topology of knots to prove a generalization of Kakutani's fixed-point theorem (Minsky 1950). At Princeton University, he focused on neural networks. In 1951, he and fellow graduate student Dean Edmonds built the stochastic neural analog reinforcement computer (SNARC), the first

hardware implementation of an artificial neural network (Bernstein 1981). His Ph.D. thesis in mathematics presented theories and proved theorems about learning in neural networks (Minsky 1954).

In the following years, he published a number of results in the theory of computation. He proved theorems about small universal sets of digital logic elements (Minsky 1956), proved that tag systems can be universal (Minsky 1961a, 1967), and discovered a four-symbol seven-state universal Turing machine (Minsky 1962). He later published a textbook on the theory of computation (Minsky 1967).

With John McCarthy, Nathaniel Rochester, and Claude Shannon, Marvin organized the 1956 Dartmouth Summer Research Project on Artificial Intelligence that is considered to be the birthplace of the field of artificial intelligence. Over the next few years, he wrote a number of papers on artificial intelligence. They included the seminal “Steps Toward Artificial Intelligence” (Minsky 1961b), which presented the major problems for future research in symbolic artificial intelligence—search, pattern recognition, learning, planning, and induction—and “Matter, Mind, and Models” (Minsky 1965), which investigated introspection and free will.

In 1968, Minsky published the influential collection *Semantic Information Processing*, which presented the state of the art in symbolic artificial intelligence at the time (Minsky 1968). M. Ross Quillian’s chapter on semantic networks spawned much further research on the topic in psychology and artificial intelligence. The book included a chapter by Thomas G. Evans on a geometric analogy program, a chapter by Daniel G. Bobrow on a program that solves algebra problems, and chapters on question-answering systems by Bertram Raphael and Fischer Black. The collection also contained reprints of several papers by Minsky as well as John McCarthy.

In 1969 with Seymour Papert, Minsky published the controversial book *Perceptrons* (Minsky and Papert 1969), which is summarized by Rumelhart and Zipser (1986) as follows:

The central theme of this work is that parallel recognizing elements, such as perceptrons, are beset by the same problems of scale as serial pattern recognizers. Combinatorial explosion catches you sooner or later, although sometimes in different ways in parallel than in serial. (p. 158)

In 1974, Minsky published the famous AI Memo No. 306 titled “A Framework for Representing Knowledge” (Minsky 1974). This seminal memo introduced the notion of frames, which was highly influential in artificial intel-

ligence and psychology. It also discussed frame systems and default assignments and how all these concepts could be applied in such areas as control of thought, imagery, language and story understanding, learning and memory, and vision. The memo is also well known for its appendix criticizing the logical approach to artificial intelligence and its computational inefficiency, insistence on consistency, and monotonicity. His criticism of logic’s monotonicity triggered work on nonmonotonic logics and led to the field of formal nonmonotonic reasoning (Brewka, Dix, and Konolige 1997; Ginsberg 1987).

Starting in the early 1970s with Seymour Papert, Marvin began developing his most famous theory of how the mind works, the society of mind (Minsky and Papert 1972, pp. 92–99). Minsky continued to evolve the theory in several papers (Minsky 1974, 1977, 1980) and published his acclaimed book *The Society of Mind* in 1986 (Minsky 1986). Singh (2004) summarizes the society of mind theory and work inspired by it, including combined symbolic-connectionist methods, the method of derivational analogy, and society-of-mind-like programming languages. In 2006, Minsky published a sequel to *The Society of Mind* titled *The Emotion Machine* (Minsky 2006).

Both *The Society of Mind* and *The Emotion Machine* are giant catalogs of what Minsky calls “resources” and “ways to think” that allow the human mind to be resourceful. Here is a sampling:

A censor is a resource that prevents an idea from coming to mind.

A k-line is a resource that records what resources were used to solve a problem so that they can be used to solve future similar problems.

A paranome is a resource that connects several representations so that the mind can easily switch between those representations.

A selector is a resource that engages a set of resources likely to help with a problem recognized by a critic.

A suppressor is a resource that prevents a dangerous action from being performed.

A trouble detector is a resource that engages higher-level resources when the usual ones for achieving a goal do not succeed.

A value is a resource that is used by resources that deal with self-conscious reflection.

The emotion of fear is a way of thinking whose evolutionary purpose is to allow us to protect ourselves.

Aspects of the society of mind have been

implemented by Minsky's students (Riecken 1994, Cassimatis 2002, Singh 2005). Minsky's books will continue to serve as a treasure trove of ideas to be mined for years to come.

Minsky's Personal History

Minsky's wide-ranging scientific and mathematical curiosity started in his childhood in New York City where he amused himself by taking apart his father's ophthalmological instruments. He thoroughly enjoyed the opportunity to focus on his interests at the Bronx High School of Science in the challenging company of several classmates who went on to become Nobel Prize-winning physicists. After spending a year at Phillips Academy, Andover, he served in the United States Navy in the final year of World War II. His undergraduate experience at Harvard University was academically dazzling; he studied classical mechanics with Herbert Goldstein, mathematics with Andrew Gleason, neuroanatomy with Marcus Singer, neurophysiology with John Welsh, and psychology with George Miller. At Harvard, he first envisioned a machine that could learn, an idea he has pursued for his entire career. He went on to graduate work in mathematics at Princeton University, where the department chair Solomon Lefschetz gave him great freedom to develop his interdisciplinary ideas on learning and intelligence. After receiving his Ph.D. in 1954, he continued this work as a junior fellow at Harvard and at the MIT Lincoln Laboratory where he wrote one of the first papers on artificial intelligence, titled "Heuristic Aspects of the Artificial Intelligence Problem." He joined the MIT faculty in 1958, where along with John McCarthy he founded the preeminent world laboratory for artificial intelligence research. He is currently a professor of electrical engineering and computer science and the Toshiba Professor of Media Arts and Sciences.

Marvin has boundless energy and creativity. He studied musical composition with the composer Irving Fine and has theorized about the appeal of music (Minsky 1981). In his spare time, he often improvises Bachlike fugues on the piano. He invented a machine that composes and plays its own music, the Triadex Muse, with Edward Fredkin (U.S. Patent 3,610,801). His other inventions include the confocal scanning microscope (1955, U.S. Patent 3,013,467), the first head-mounted graphical display (1963), the concept of a binary-tree robotic manipulator (1963), and the serpentine hydraulic robot arm (1967). He was a founder of Logo Computer Systems (publisher

of educational software for children) and Thinking Machines Corporation (maker of the Connection Machine supercomputer).

He has written about possibilities for a digital physics with particles and fields existing within cellular automata (Minsky 1982). His interest in science also extends to science fiction. He served as a technical consultant for the movie *2001: A Space Odyssey*. With Harry Harrison, he wrote a science fiction novel, *The Turing Option*.

Marvin and his wife, Gloria Rudisch, a pediatrician, have three children, Julie Minsky, Henry Minsky, and Margaret Minsky, and four grandchildren. His living room, immortalized in the *Society of Mind* CD, is packed with interesting things, including two pianos, a harp, sculptures, paintings, old Macs, a SNARC neuron, a small rocket, mementos from many distinguished personalities/friends (such as Bono from U2, Larry Bird of the Boston Celtics, Gene Roddenberry and the cast from *Star Trek*, and so on), and plastic storage bins full of fun components, gadgets, and toys.

Notes

1. www.media.mit.edu/~minsky.

References

- Bernstein, J. 1981. Profiles (Marvin Minsky). *The New Yorker* 50–126, December 14.
- Brewka, G.; Dix, J.; and Konolige, K. 1997. *Nonmonotonic Reasoning: An Overview*. Stanford, CA: Center for the Study of Language and Information (CSLI).
- Cassimatis, N. L. 2002. Polyscheme: A Cognitive Architecture for Integrating Multiple Representation and Inference Schemes. Doctoral dissertation, Program in Media Arts and Sciences, School of Architecture and Planning, Massachusetts Institute of Technology, Cambridge, MA.
- Ginsberg, M. L., ed. 1987. *Readings in Nonmonotonic Reasoning*. Los Altos, CA: Morgan Kaufmann.
- Minsky, M. 1950. A Generalization of Kakutani's Fixed-Point Theorem. Bachelor's Thesis, Department of Mathematics, Harvard University, Cambridge, MA.
- Minsky, M. 1954. Neural Nets and the Brain Model Problem. Doctoral dissertation, Department of Mathematics, Princeton University, Princeton, NJ.
- Minsky, M. 1956. Some Universal Elements for Finite Automata. In *Automata Studies*, Volume 34, ed. C. E. Shannon and J. McCarthy, 117–128. Princeton, NJ: Princeton University Press.
- Minsky, M. 1961a. Recursive Unsolvability of Post's Problem of "Tag" and Other Topics in the Theory of Turing Machines. *Annals of Mathematics* 74(3): 437–455.
- Minsky, M. 1961b. Steps Toward Artificial Intelligence. *Proceedings of the IRE* 49(1): 8–30.
- Minsky, M. 1962. Size and Structure of Universal Turing Machines Using Tag Systems. In *Recursive Function Theory, Proceedings of Symposia in Pure Mathematics*



Illustration © Julie Ridge. All rights reserved.

ics, Volume 5, ed. J. C. E. Dekker, 229–238. Providence, RI: American Mathematical Society.

Minsky, M. 1965. Matter, Mind, and Models. In *Proceedings of the International Federation for Information Processing (IFIP) Congress*, 45–49. Washington, DC: Spartan Books.

Minsky, M. 1967. *Computation: Finite and Infinite Machines*. Englewood Cliffs, NJ: Prentice-Hall.

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

Minsky, M. 1974. A Framework for Representing Knowledge. A. I. Memo 306, Cambridge, MA: Artificial Intelligence Laboratory, Massachusetts Institute of Technology.

Minsky, M. 1977. Plain Talk about Neurodevelopmental Epistemology. In *Proceedings of the Fifth International Joint Conference on Artificial Intelligence*, ed. R. Reddy, 1083–1092. Los Altos, CA: William Kaufmann.

Minsky, M. 1980. K-lines: A Theory of Memory. *Cognitive Science* 4: 117–133.

Minsky, M. 1981. Music, Mind, and Meaning. *Computer Music Journal* 5(3): 8–44.

Minsky, M. 1982. Cellular Vacuum. *International Journal of Theoretical Physics* 21(6/7): 537–551.

Minsky, M. 1986. *The Society of Mind*. New York: Simon and Schuster.

Minsky, M. 2005. Interior Grounding, Reflection, and Self-Consciousness, Proceedings of an International Conference on Brain, Mind, and Society, Tohoku University, Graduate School of Information Sciences, Sendai, Japan.

Minsky, M. 2006. *The Emotion Machine: Commonsense Thinking, Artificial Intelligence, and the Future of the Human Mind*. New York: Simon and Schuster.

Minsky, M., and Papert, S. 1969. *Perceptrons: An Introduction to Computational Geometry*. Cambridge, MA: The MIT Press.

Minsky, M., and Papert, S. 1972. Artificial Intelligence Progress Report. A.I. Memo 252, Cambridge, MA: Artificial Intelligence Laboratory, Massachusetts Institute of Technology.

Riecken, D. 1994. M: An Architecture of Integrated

Agents. *Communications of the ACM* 37(7): 107–116, 146.

Rumelhart, D. E., and Zipser, D. 1986. Feature Discovery by Competitive Learning. In *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Volume 1: Foundations*, ed. D. E. Rumelhart, J. L. McClelland, and PDP Research Group, 151–193. Cambridge, MA: The MIT Press.

Singh, P. 2004. Examining the Society of Mind. *Computing and Informatics* 22(6): 521–543.

Singh, P. 2005. EM-ONE: An Architecture for Reflective Commonsense Thinking. Doctoral dissertation, Department of Electrical Engineering and Computer Science, Massachusetts Institute of Technology, Cambridge, MA.

Danny Hillis, Ph.D., is a scientist and currently cochairman and chief technical officer of Applied Minds, Inc. Hillis also developed the Connection Machine, along with many other significant inventions, and is a cofounder of Thinking Machines Corp.

John McCarthy, Ph.D., is a scientist, a Turing Award recipient, a former president of AAAI, and a cofounder of the field of artificial intelligence. McCarthy is a professor emeritus of computer science at Stanford University. Long ago he originated the Lisp programming language and the initial research on general-purpose timesharing computer systems.

Tom M. Mitchell, Ph.D., is a scientist and a preeminent member of the machine learning community. A former president of AAAI, Mitchell is the Fredkin Professor of AI and Machine Learning, and is chair of the Machine Learning Department at Carnegie Mellon University.

Erik T. Mueller is a scientist at the IBM Thomas J. Watson Research Center. He is the author of two books and creator of the ThoughtTreasure commonsense knowledge base. His research focuses on common sense reasoning and its application to story understanding and intelligent user interfaces.

Doug Riecken, Ph.D., is a scientist and department head of the Commonsense Computing Research Department at IBM Research. Riecken developed the Wolfgang and M systems. His research is in common sense reasoning, agent-based systems, and intelligent user interfaces.

Aaron Sloman, Ph.D., is a scientist and honorary professor of artificial intelligence and cognitive science at the University of Birmingham. Sloman developed Poplog. His research focus is in the study of mind, along with agent architectures, motivation, emotion, vision, causation, and consciousness.

Patrick H. Winston, Ph.D., is a scientist and Ford Professor of Artificial Intelligence and Computer Science at MIT. A former president of AAAI, Winston is also chairman and cofounder of Ascent Technology, Inc. Winston's research is focused on developing an account of human intelligence, with particular emphasis on the role of language, vision, and tightly coupled loops connecting the two.