

PROFILES

WHY THE GODFATHER OF A.I. FEARS WHAT HE'S BUILT

Geoffrey Hinton has spent a lifetime teaching computers to learn. Now he worries that artificial brains are better than ours.

By Joshua Rothman

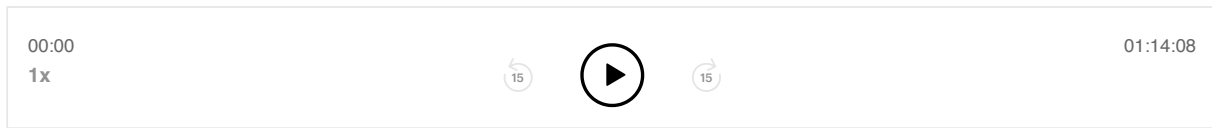
November 13, 2023



"There's a very general subgoal that helps with almost all goals: get more control," Hinton said of A.I.s. "The research question is: how do you prevent them from ever wanting to take control? And nobody knows the answer." Illustration by Daniel Liévano



Save this story



In your brain, neurons are arranged in networks big and small. With every action, with every thought, the networks change: neurons are included or excluded, and the connections between them strengthen or fade. This process goes on all the time—it’s happening now, as you read these words—and its scale is beyond imagining. You have some eighty billion neurons sharing a hundred trillion connections or more. Your skull contains a galaxy’s worth of constellations, always shifting.

Geoffrey Hinton, the computer scientist who is often called “the godfather of A.I.,” handed me a walking stick. “You’ll need one of these,” he said. Then he headed off along a path through the woods to the shore. It wound across a shaded clearing, past a pair of sheds, and then descended by stone steps to a small dock. “It’s slippery here,” Hinton warned, as we started down.

New knowledge incorporates itself into your existing networks in the form of subtle adjustments. Sometimes they’re temporary: if you meet a stranger at a party, his name might impress itself only briefly upon the networks in your memory. But they can also last a lifetime, if, say, that stranger becomes your spouse. Because new knowledge merges with old, what you know shapes what you learn. If someone at the party tells you about his trip to Amsterdam, the next day, at a museum, your networks may nudge you a little closer to the Vermeer. In this way, small changes create the possibility for profound transformations.

“We had a bonfire here,” Hinton said. We were on a ledge of rock jutting out into Ontario’s Georgian Bay, which stretches to the west into Lake Huron. Islands dotted the water; Hinton had bought this one in 2013, when he was sixty-five, after selling a three-person startup to Google for forty-four million dollars. Before that, he’d spent three decades as a computer-science professor at the University of Toronto—a leading figure in an unglamorous subfield known as neural networks, which was inspired by the way neurons are connected in the brain. Because artificial neural networks were only moderately successful at the tasks they undertook—image categorization, speech recognition, and so on—most researchers considered them to be at best mildly interesting, or at worst a waste of time. “Our neural nets just couldn’t do anything better than a child could,” Hinton recalled. In the nineteen-eighties, when he saw “The Terminator,” it didn’t bother him that Skynet, the movie’s world-destroying A.I., was a neural net; he was pleased to see the technology portrayed as promising.

From the small depression where the fire had been, cracks in the stone, created by the heat, radiated outward. Hinton, who is tall, slim, and English, poked the spot with his stick. A scientist through and through, he is always remarking on what is happening in the physical world: the lives of animals, the flow of currents in the bay, the geology of the island. “I put a mesh of rebar under the wood, so the air

could get in, and it got hot enough that the metal actually went all soft,” he said, in a wondering tone. “That’s a real fire—something to be proud of!”

For decades, Hinton tinkered, building bigger neural nets structured in ingenious ways. He imagined new methods for training them and helping them improve. He recruited graduate students, convincing them that neural nets weren’t a lost cause. He thought of himself as participating in a project that might come to fruition a century in the future, after he died. Meanwhile, he found himself widowed and raising two young children alone. During one particularly difficult period, when the demands of family life and research overwhelmed him, he thought that he’d contributed all he could. “I was dead in the water at forty-six,” he said. He didn’t anticipate the speed with which, about a decade ago, neural-net technology would suddenly improve. Computers got faster, and neural nets, drawing on data available on the Internet, started transcribing speech, playing games, translating languages, even driving cars. Around the time Hinton’s company was acquired, an A.I. boom began, leading to the creation of systems like OpenAI’s ChatGPT and Google’s Bard, which many believe are starting to change the world in unpredictable ways.

Hinton set off along the shore, and I followed, the fractured rock shifting beneath me. “Now watch this,” he said. He stood before a lumpy, person-size boulder, which blocked our way. “Here’s how you get across. You throw your stick”—he tossed his to the other side of the boulder—“and then there are footholds here and here, and a handhold here.” I watched as he scrambled over with easy familiarity, and then, more tentatively, I took the same steps myself.

Whenever we learn, our networks of neurons change—but how, exactly? Researchers like Hinton, working with computers, sought to discover “learning algorithms” for neural nets, procedures through which the statistical “weights” of the connections among artificial neurons could change to assimilate new knowledge. In 1949, a psychologist named Donald Hebb proposed a simple rule for how people learn, often summarized as “Neurons that fire together wire together.” Once a group of neurons in your brain activates in synchrony, it’s more likely to do so again; this helps explain why doing something is easier the second time. But it quickly became apparent that computerized neural networks needed another approach in order to solve complicated problems. As a young researcher, in the nineteen-sixties and seventies, Hinton drew networks of neurons in notebooks and imagined new knowledge arriving at their borders. How would a network of a few hundred artificial neurons store a concept? How would it revise that concept if it turned out to be flawed?

ADVERTISEMENT

We made our way around the shore to Hinton's cottage, the only one on the island. Glass-enclosed, it stood on stilts atop a staircase of broad, dark rocks. "One time, we came out here and a huge water snake stuck his head up," Hinton said, as we neared the house. It was a fond memory. His father, a celebrated entomologist who'd named a little-known stage of metamorphosis, had instilled in him an affection for cold-blooded creatures. When he was a child, he and his dad kept a pit full of vipers, turtles, frogs, toads, and lizards in the garage. Today, when Hinton is on the island—he is often there in the warmer months—he sometimes finds snakes and brings them into the house, so that he can watch them in a terrarium. He is a good observer of nonhuman minds, having spent a lifetime thinking about thinking from the bottom up.

Earlier this year, Hinton left Google, where he'd worked since the acquisition. He was worried about the potential of A.I. to do harm, and began giving interviews in which he talked about the "existential threat" that the technology might pose to the human species. The more he used ChatGPT, an A.I. system trained on a vast corpus of human writing, the more uneasy he got. One day, someone from Fox News wrote to him asking for an interview about artificial intelligence. Hinton enjoys sending snarky single-sentence replies to e-mails—after receiving a lengthy note from a Canadian intelligence agency, he responded, "Snowden is my hero"—and he began experimenting with a few one-liners. Eventually, he wrote, "Fox News is an oxy moron." Then, on a lark, he asked ChatGPT if it could explain his joke. The system told him his sentence implied that Fox News was fake news, and, when he called attention to the space before "moron," it explained that Fox News was addictive, like the drug OxyContin. Hinton was astonished. This level of understanding seemed to represent a new era in A.I.

There are many reasons to be concerned about the advent of artificial intelligence. It's common sense to worry about human workers being replaced by computers, for example. But Hinton has joined many prominent technologists, including Sam Altman, the C.E.O. of OpenAI, in warning that A.I. systems may start to think for themselves, and even seek to take over or eliminate human civilization. It was striking to hear one of A.I.'s most prominent researchers give voice to such an alarming view.

"People say, It's just glorified autocomplete," he told me, standing in his kitchen. (He has suffered from back pain for most of his life; it eventually grew so severe that he gave up sitting. He has not sat down for more than an hour since 2005.) "Now, let's analyze that. Suppose you want to be really good at

predicting the next word. If you want to be *really* good, you have to understand what's being said. That's the only way. So by training something to be really good at predicting the next word, you're actually forcing it to understand. Yes, it's 'autocomplete'—but you didn't think through what it means to have a really good autocomplete.” Hinton thinks that “large language models,” such as GPT, which powers OpenAI's chatbots, can comprehend the meanings of words and ideas.



“Of course I mind—they’re mine, and I want all of them.”

Cartoon by Tom Chitty



[Open cartoon gallery](#)

Skeptics who say that we overestimate the power of A.I. point out that a great deal separates human minds from neural nets. For one thing, neural nets don't learn the way we do: we acquire knowledge organically, by having experiences and grasping their relationship to reality and ourselves, while they learn abstractly, by processing huge repositories of information about a world that they don't really inhabit. But Hinton argues that the intelligence displayed by A.I. systems transcends its artificial origins.

“When you eat, you take food in, and you break it down to these tiny components,” he told me. “So you could say that the bits in my body are made from bits of other animals. But that would be very misleading.” He believes that, by analyzing human writing, a large language model like GPT learns how the world works, producing a system capable of thought; writing is only part of what that system can

do. “It’s analogous to how a caterpillar turns into a butterfly,” he went on. “In the chrysalis, you turn the caterpillar into soup—and from this soup you build the butterfly.”

He began rooting around in a small cupboard just off the kitchen. “Aha!” he said. With a flourish, he put an object on the counter—a dead dragonfly. It was perfectly preserved. “I found this at the marina,” he explained. “It had just hatched on a rock and was drying its wings, so I caught it. Look underneath.” Hinton had captured the dragonfly just after it had emerged from its larval form. The larva was a quite different-looking insect, with its own eyes and legs; it had a hole in its back, through which the dragonfly had crawled.

“The larva of the dragonfly is this monster that lives under the water,” Hinton said. “And, like in the movie ‘Alien,’ the dragonfly is breaking out of the back of the monster. The larva went into a phase where it got turned into soup, and then a dragonfly was built out of the soup.” In his metaphor, the larva represented the data that had gone into training modern neural nets; the dragonfly stood for the agile A.I. that had been created from it. Deep learning—the technology that Hinton helped pioneer—had caused the metamorphosis. I bent closer to get a better look; Hinton stood upright, as he almost always does, careful to preserve his posture. “It’s very beautiful,” he said softly. “And you get the point. It started as one thing, and it’s become something else.”

A few weeks earlier, when Hinton had invited me to visit his island, I’d imagined possible scenarios. Perhaps he’d be an introvert who wanted solitude, or a tech overlord with a God complex and a futuristic compound. Several days before my arrival, he e-mailed me a photograph he’d taken of a rattlesnake coiled in the island’s grass. I wasn’t sure whether I felt delighted or scared.

In fact, as private islands go, Hinton’s is fairly modest—two acres in total. Hinton himself is the opposite of a Silicon Valley techno-messiah. Now seventy-five, he has an English face out of a Joshua Reynolds painting, with white hair framing a broad forehead; his blue eyes are often steady, leaving his mouth to express emotion. A mordant raconteur, he enjoys talking about himself—“ ‘Geoff’ is an anagram for ‘ego fortissimo,’ ” he told me—but he’s not an egotist; his life has been too grief-shadowed for that. “I should probably tell you about my wives,” he said, the first time we spoke. “I’ve had three marriages. One ended amicably, the other two in tragedy.” He is still friendly with Joanne, his first wife, whom he married early, but his second and third wives, Rosalind and Jackie, both died of cancer, in 1994 and 2018, respectively. For the past four years, Hinton has been with Rosemary Gartner, a retired sociologist. “I think he’s the kind of person who always needs a partner,” she told me, tenderly. He is a romantic rationalist, with a sensibility balancing science and emotion. In the cottage, a burgundy canoe sits in the single large room that makes up most of the ground floor; he and Jackie had found it in the island’s woods, in disrepair, and Jackie, an art historian, worked with some women canoe-builders to reconstruct it during the years coinciding with her illness. “She had the maiden voyage,” Hinton said. No one has used it since.

He stowed the dragonfly, then walked over to a small standing desk, where a laptop was perched next to a pile of sudoku puzzles and a notebook containing computer passwords. (He rarely uses the notebook, having devised a mnemonic system that enables him to generate and recall very long passwords in his head.) “Shall we do the family tree?” he asked. Using two fingers—he doesn’t touch-type—he entered “Geoffrey Hinton family tree” and hit Return. When Google acquired Hinton’s startup, in 2013, it did so in part because the team had figured out how to dramatically improve image recognition using neural nets; now endless family trees swarmed the screen.

Hinton comes from a particular kind of scientific English family: politically radical, restlessly inventive. Above him in the family tree are his great-uncle Sebastian Hinton, the inventor of the jungle gym, and his cousin Joan Hinton, who worked as a physicist on the Manhattan Project. Further back, he was preceded by Lucy Everest, the first woman to become an elected member of the Royal Institute of Chemistry; Charles Howard Hinton, the mathematician who created the concept of the tesseract, a doorway into the fourth dimension (one appears in the film “Interstellar”); and James Hinton, a groundbreaking ear surgeon and an advocate of polygamy. (“Christ was the savior of men, but I am the savior of women,” he is said to have remarked.) In the mid-nineteenth century, a great-great-grandfather of Hinton’s, the English mathematician George Boole, developed the system of binary reasoning, now known as Boolean algebra, that is fundamental to all computing. Boole was married to Mary Everest, a mathematician and author and the niece of George Everest, the surveyor for whom Mt. Everest is named.

“Geoff was born into science,” Yann LeCun, a former student and collaborator of Hinton’s who now runs A.I. at Meta, told me. Yet Hinton’s family was odder than that. His dad, Howard Everest Hinton, grew up in Mexico during the Mexican Revolution, in the nineteen-tens, on a silver mine managed by his father. “He was tough,” Hinton said of his dad: family lore holds that, at age twelve, Howard threatened to shoot his boxing coach for being too heavy-handed, and the coach took him seriously enough to leave town. Howard’s first language was Spanish, and at Berkeley, where he went to college, he was mocked for his accent. “He hung out with a bunch of Filipinos, who were also discriminated against, and he became a Berkeley radical,” Hinton said. Howard’s mature politics were not just Marxist but Stalinist: in 1968, as Soviet tanks rolled into Prague, he said, “About time!”

At school, Hinton was inclined toward science. But, for ideological reasons, his father forbade him to study biology; in Howard's view, the possibility of genetic determinism contravened the Communist belief in the ultimate malleability of human nature. ("I hate faiths of all kinds," Hinton said, remembering this period.) Howard, who taught at the University of Bristol, was a kind of entomologist Indiana Jones: he smuggled rare creatures from around the world back to England in his luggage, and edited an important journal in his field. Hinton, whose middle name is also Everest, felt immense pressure to make his own mark. He recalls his father telling him, "If you work twice as hard as me, when you're twice as old as I am you might be half as good."

At Cambridge, Hinton tried different fields but was dismayed to find that he was never the brightest student in any given class. He left college briefly to "read depressing novels" and to do odd jobs in London, then returned to attempt architecture, for about a day. Finally, after dipping into physics, chemistry, physiology, and philosophy, looking for a focus, he settled on a degree in experimental psychology. He haunted the office hours of the moral philosopher Bernard Williams, who turned out to be interested in computers and the mind. One day, Williams pointed out that our different thoughts must reflect different physical arrangements inside our brains; this was quite unlike the situation inside a computer, in which the software was independent of the hardware. Hinton was struck by this observation; he remembered how, in high school, a friend had told him that memory might be stored in the brain "holographically"—that is, spread out, but in such a way that the whole could be accessed through any one part. What he was encountering was "connectionism"—an approach that combined neuroscience, math, philosophy, and programming to explore how neurons could work together to "think." One goal of connectionism was to create a brainlike system in a computer. There had been some progress: the Perceptron, a machine built in the nineteen-fifties by a psychologist and pioneering connectionist named Frank Rosenblatt, had used simple computer hardware to simulate a network of hundreds of neurons. When connected to a light sensor, the apparatus could recognize letters and shapes by tracking which artificial neurons were activated by different patterns of light.

In the cottage, Hinton stood and strolled, ranging back and forth behind the kitchen counter and around the first floor. He made some toast, got us each an apple, and then set up a little booster table for himself using a step stool. Family pressure had had the effect of pushing him out of temporary satisfactions. "I always loved woodwork," he recalled wistfully, while we ate. "At school, you could do it voluntarily in the evenings. And I've often wondered whether I'd have been happier as an architect, because I didn't have to force myself to do it. Whereas, with science, I've always had to force myself. Because of the family, I had to succeed at it—I had to find a path. There was joy in it, but it was mostly anxiety. Now it's an enormous relief that I've succeeded."

Hinton's laptop dinged. Ever since he'd left Google, his in-box had been exploding with requests for comment on A.I. He ambled over and looked at the e-mail, and then got lost again in the forest of family trees, all of which seemed to be wrong in one way or another.

"Look at this," he said.

I walked over and peered at the screen. It was an "academic family tree," showing Hinton at the top with his students, and theirs, arrayed below. The tree was so broad that he had to scroll horizontally to see the extent of his influence. "Oh, dear," Hinton said, exploring. "She wasn't really a student of mine." He scrolled further. "He was brilliant but not so good as an adviser, because he could always do it better himself." A careful nurturer of talent, Hinton seems to enjoy being surpassed by his students: when evaluating job candidates, he used to ask their advisers, "But are they better than *you*?" Recalling his father, who died in 1977, Hinton said, "He was just extremely competitive. And I've often wondered, if he'd been around to see me be successful, whether he'd have been entirely happy. Because now I've been more successful than he was."

According to Google Scholar, Hinton is now the second most cited researcher among psychologists, and the most cited among computer and cognitive scientists. If he had a slow and eccentric start at Cambridge, it was partly because he was circling an emerging field. "Neural networks—there were very few people at good universities who did it," he said, closing the laptop. "You couldn't do it at M.I.T. You couldn't do it at Berkeley. You couldn't do it at Stanford." There were advantages to being a hub in a nascent network. For years, many of the best minds came to him.

"The weather's good," Hinton said, the next morning. "We should cut down a tree." He wore a dress shirt tucked into khakis and didn't look much like a lumberjack; still, he rubbed his hands together. On the island, he is always cutting down trees to create more orderly and beautiful tableaux.

The house, too, is a work in progress. Few contractors would travel to a place so remote, and the people Hinton hired made needless mistakes (running a drainage pipe uphill, leaving floors half finished) that still enrage him today. Almost every room harbors a corrective mini-project, and, when I visited, Hinton had appended little notes to them to help a new contractor, often writing on the building materials themselves. In the first-floor bathroom, a piece of baseboard propped against the wall read "Bathroom

should have THIS type of baseboard (maple trim in front of shower only).” In the guest-room closet, masking tape ran along a shelf: “Do not prime shelf, prime shelf support.”

It’s useful for minds to label things; it helps them get a grip on reality. But what would it mean for an artificial mind to do so? While Hinton was earning a Ph.D. in artificial intelligence from the University of Edinburgh, he thought about how “knowing” in a brain might be simulated in a computer. At that time, in the nineteen-seventies, the vast majority of A.I. researchers were “symbolists.” In their view, knowing about, say, ketchup might involve a number of concepts, such as “food,” “sauce,” “condiment,” “sweet,” “umami,” “red,” “tomato,” “American,” “French fries,” “mayo,” and “mustard”; together, these could create a scaffold on which a new concept like “ketchup” might be hung. A large, well-funded A.I. effort called Cyc centered on the construction of a vast knowledge repository into which scientists, using a special language, could enter concepts, facts, and rules, along with their inevitable exceptions. (Birds fly, but not penguins or birds with damaged wings or . . .)

But Hinton was doubtful of this approach. It seemed too rigid, and too focussed on the reasoning skills possessed by philosophers and linguists. In nature, he knew, many animals acted intelligently without access to concepts that could be expressed in words. They simply learned how to be smart through experience. Learning, not knowledge, was the engine of intelligence.

Sophisticated human thinking often seemed to happen through symbols and words. But Hinton and his collaborators, James L. McClelland and David Rumelhart, believed that much of the action happened on a sub-conceptual level. Notice, they wrote, how, “if you learn a new fact about an object, your expectations about other similar objects tend to change”: if you’re told that chimpanzees like onions, for instance, you might guess that gorillas like them, too. This suggested that knowledge was likely “distributed” in the mind—created out of smaller building blocks that could be shared among related ideas. There wouldn’t be two separate networks of neurons for the concepts “chimpanzee” and “gorilla”; instead, bundles of neurons representing various concrete or abstract “features”—furriness, quadrupedness, primateness, animalness, intelligence, wildness, and so on—might be activated in one way to signify “chimpanzee” and in a slightly different way to signify “gorilla.” To this cloud of features, onion-liking-ness might be added. A mind constructed this way risked falling into confusion and error: mix qualities together in the wrong arrangement and you’d get a fantasy creature that was neither gorilla nor chimp. But a brain with the right learning algorithm might adjust the weights among its neurons to favor sensible combinations over incoherent ones.

Hinton continued to explore these ideas, first at the University of California, San Diego, where he did a postdoc (and married Joanne, whom he tutored in computer vision); then at Cambridge, where he worked as a researcher in applied psychology; and then at Carnegie Mellon, in Pittsburgh, where he became a computer-science professor in 1982. There, he spent much of his research budget on a single computer powerful enough to run a neural net. He soon got married a second time, to Rosalind Zalin, a

molecular biologist. At Carnegie Mellon, Hinton had a breakthrough. Working with Terrence Sejnowski, a computer scientist and a neuroscientist, he produced a neural net called the Boltzmann Machine. The system was named for Ludwig Boltzmann, the nineteenth-century Austrian physicist who described, mathematically, how the large-scale behavior of gases was related to the small-scale behavior of their constituent particles. Hinton and Sejnowski combined these equations with a theory of learning.

ADVERTISEMENT

Hinton was reluctant to explain the Boltzmann Machine to me. “I’ll tell you what this is like,” he said. “It’s like having a small child, and you decide to go on a walk. And there’s a mountain ahead of you, and you have to get this little child to the top of the mountain and back.” He looked at me—the child in the metaphor—and sighed. He worried, reasonably, that I might be misled by a simplified explanation and then mislead others. “It’s no use trying to explain complicated ideas that you don’t understand. First, you have to understand how something works. Otherwise, you just produce nonsense.” Finally, he took some sheets of paper and began drawing diagrams of neurons connected by arrows and writing out equations, which I tried to follow. (Ahead of my visit, I’d done a Khan Academy course on linear algebra.)

One way to understand the Boltzmann Machine, he suggested, was to imagine an Identi-Kit: a system through which various features of a face—bushy eyebrows, blue eyes, crooked noses, thin lips, big ears, and so on—can be combined to produce a composite sketch, of the sort used by the police. For an Identi-Kit to work, the features themselves have to be appropriately designed. The Boltzmann Machine could learn not just to assemble the features but to design them, by altering the weights of the connections among its artificial neurons. It would start with random features that looked like snow on a television screen, and then proceed in two phases—“waking” and “sleeping”—to refine them. While awake, it would tweak the features so that they better fit an actual face. While asleep, it would fantasize a face that didn’t exist, and then alter the features so that they were a worse fit.

Its dreams told it what not to learn. There was an elegance to the system: over time, it could move away from error and toward reality, and no one had to tell it if it was right or wrong—it needed only to see what existed, and to dream about what didn’t.



"It's a filter that makes your baby look as cute as you think it is!"

Cartoon by Kit Fraser



[Open cartoon gallery](#)



Hinton and Sejnowski described the Boltzmann Machine in a 1983 paper. "I read that paper when I was starting my graduate studies, and I said, 'I absolutely have to talk to these guys—they're the only people in the world who understand that we need learning algorithms,'" Yann LeCun told me. In the mid-eighties, Yoshua Bengio, a pioneer in natural-language processing and in computer vision who is now the scientific director at Mila, an A.I. institute in Quebec, trained a Boltzmann Machine to recognize spoken syllables as part of his master's thesis. "Geoff was one of the external reviewers," he recalled. "And he wrote something like 'This should not work.'" Bengio's version of the Boltzmann Machine was more effective than Hinton expected; it took Bengio a few years to figure out why. This would become a familiar pattern. In the following decades, neural nets would often perform better than expected, perhaps because new structures had formed among the neurons during training. "The experimental part of the work came before the theory," Bengio recalled. Often, it was a matter of trying new approaches and seeing what the networks came up with.

Partly because Rosalind loathed Ronald Reagan, Hinton said, they moved to the University of Toronto. They adopted two children, a boy and a girl, from Latin America, and lived in a house in the city. "I was this kind of socialist professor who was dedicated to his work," Hinton said.

Rosalind had struggled with infertility, and had bad experiences with callous doctors. Perhaps as a result, she pursued a homeopathic route when she was later diagnosed with ovarian cancer. "It just didn't make any sense," Hinton said. "It couldn't be that you make things more dilute and they get more

powerful.” He couldn’t see how a molecular biologist could become a homeopath. Still, determined to treat the cancer herself, Rosalind refused to have surgery even after an exam found a tumor the size of a grapefruit; later, she consented to an operation but declined chemotherapy, instead pursuing increasingly expensive homeopathic remedies, first in Canada and then in Switzerland. She developed secondary tumors. She asked Hinton to sell their house so that she could pay for new homeopathic treatments. “I drew the line there,” he recalled, squinting with fresh pain. “I said, ‘No, we’re not selling the house. Because if you die I’m going to have to look after the children, and it’s much better for them if we can stay.’ ”

Rosalind returned to Canada and went immediately into the hospital. She hung on for a couple of months, but wouldn’t let the children visit her until the day before she died, because she didn’t want them to see her so sick. Throughout her illness, she was convinced that she’d soon get well. Describing what happened, Hinton still seems overwhelmed—he is angry, guilty, wounded, mystified. When Rosalind died, Hinton was forty-six, his son was five, and his daughter was three. “She hurt people by failing to accept that she was going to die,” he said.

The sound of waves filled the midafternoon quiet. Strong yellow sun spilled through the room’s floor-to-ceiling windows; faint spiderwebs extended across them, silhouetted by the light. Hinton stood for a while, collecting himself.

“I think I need to go cut down a tree,” he said.

We walked out the front door and down the path to the sheds. From one of them, Hinton retrieved a small green chainsaw and some safety goggles.

ADVERTISEMENT

“Rosemary says I’m not allowed to cut down trees when there’s nobody else here, in case I chop off an arm or something,” he said. “Have you driven boats before?”

“No,” I said.

“I’ve got to not chop off my right arm, then.”

Over his khakis, he strapped on a pair of protective chaps.

“I don’t want to give you the impression that I know what I’m doing,” he said. “But the basic idea is, you cut lots of V’s, and then the tree falls down.”

Hinton crossed the path to the tree that he had in mind, inspecting the bushes for snakes as we walked. The tree was a leafy cedar, perhaps twenty feet tall; Hinton looked up to see which way it was leaning, then started the saw and began to cut into the trunk on the side of the lean. He removed the saw, and made another converging cut to form a V.

Hinton worked the chainsaw in silence, occasionally stopping to wipe his brow. It was hot in the sun, and mosquitoes swarmed every shady nook. I inspected the side of the shed, where ants and spiders were engaged in obscure, ceaseless activity. Down at the end of the path, the water shone. It was a beautiful spot. Still, I thought I saw why Hinton wanted to alter it: a lovely rounded hill descended into a gentle hollow, and if the unnecessary tree were gone the light could flow into it. The tree was an error.

Eventually, he began a second cut on the other side of the tree, angling it toward the first. Then he stopped and turned to me. “Because the tree leans away from the cut, the V will open up as you go deeper, and the blade won’t get stuck,” he explained. He continued the upper cut, nudging the tree toward an entropic moment. Suddenly, almost soundlessly, gravity took over. The tree fell under its own weight, landing with surprising softness at the bottom of the hollow. The light streamed in.

Hinton was in love with the Boltzmann Machine. He hoped that it, or something like it, might underlie learning in the actual brain. “It should be true,” he told me. “If I was God, I’d make it true.” But further experimentation revealed that as Boltzmann Machines grew they tended to become overwhelmed by the randomness that was built into them. “Geoff and I disagreed about the Boltzmann Machine,” LeCun said. “Geoff thought it was the most beautiful algorithm. I thought it was ugly. It was stochastic”—that is, based partly on randomness. By contrast, LeCun said, “I thought backprop was super clean.”

“Backprop,” or backpropagation, was an algorithm that had been explored by a few different researchers beginning in the nineteen-sixties. Even as Hinton was working with Sejnowski on the Boltzmann Machine, he was also collaborating with Rumelhart and another computer scientist, Ronald Williams, on backprop. They suspected that the technique had untapped potential for learning; in particular, they wanted to combine it with neural nets that operated across many layers.

One way to understand backprop is to imagine a Kafkaesque judicial system. Picture an upper layer of a neural net as a jury that must try cases in perpetuity. The jury has just reached a verdict. In the dystopia in which backprop unfolds, the judge can tell the jurors that their verdict was wrong, and that they will be punished until they reform their ways. The jurors discover that three of them were especially

influential in leading the group down the wrong path. This apportionment of blame is the first step in backpropagation.

In the next step, the three wrongheaded jurors determine how they themselves became misinformed. They consider their own influences—parents, teachers, pundits, and the like—and identify the individuals who misinformed them. Those blameworthy influencers, in turn, must identify their respective influences and apportion blame among them. Recursive rounds of finger-pointing ensue, as each layer of influencers calls its own influences to account, in a backward-sweeping cascade. Eventually, once it's known who has misinformed whom and by how much, the network adjusts itself proportionately, so that individuals listen to their “bad” influences a little less and to their “good” influences a little more. The whole process repeats again and again, with mathematical precision, until verdicts—not just in this one case but in all cases—are collectively as “correct” as possible.

In 1986, Hinton, Rumelhart, and Williams published a three-page paper in *Nature* showing how such a system could work in a neural net. They noted that backprop, like the Boltzmann Machine, wasn't “a plausible model of learning in brains”: unlike a computer, a brain can't rewind the tape to audit its past performance. But backprop still enabled a brainlike neural specialization. In real brains, neurons are sometimes arranged in structures aimed at solving specific problems: in the visual system, for instance, different “columns” of neurons recognize edges in what we see. Something similar emerges in a backprop network. Higher layers subject lower ones to a kind of evolutionary pressure; as a result, certain layers of a network that's tasked with deciphering handwriting, for instance, might become tightly focussed on identifying lines, curves, or edges. Eventually, the system as a whole can develop “appropriate internal representations.” The network knows, and makes use of its knowledge.

In the nineteen-fifties and sixties, a great deal of excitement had accompanied the Perceptron and other connectionist efforts; enthusiasm for connectionism waned in the years after. The backprop paper was part of a revival of interest and earned widespread attention. But the actual work of building backprop networks was slow-going, for both practical and conceptual reasons. Practically, computers were sluggish. “The rate of progress was basically, How much could a computer learn overnight?” Hinton recalled. “The answer was often not much.” Conceptually, neural nets were mysterious. It wasn't possible to program one in the traditional way. You couldn't go in and edit the weights of the connections among artificial neurons. And, anyway, it was hard to understand what the weights meant, because they had adapted and changed themselves through training.

ADVERTISEMENT

There were many ways the learning process could go wrong. In “overfitting,” for example, a network effectively memorized the training data instead of learning to generalize from it. Avoiding the various pitfalls wasn’t always straightforward, because it was up to the network to learn. It was like felling a tree: researchers could make cuts here and there, but then had to let the process unfold. They could try techniques like “ensembling” (combining weak networks to make a strong one) or “early stopping” (letting a network learn, but not too much). They could “pre-train” a system, by taking a Boltzmann Machine, having it learn something, and then layering a backprop network on top of it, so that a system’s “supervised” training didn’t begin until it had acquired some elemental knowledge on its own. Then they’d let the network learn, hoping that it would land where they wanted it.

New neural-net “architectures” were developed: “recurrent” and “convolutional” networks allowed the systems to make progress by building on their own work in different ways. But it was as though researchers had discovered an alien technology that they didn’t know how to use. They turned the Rubik’s Cube this way and that, trying to pull order out of noise. “I was always convinced it wasn’t nonsense,” Hinton said. “It wasn’t really faith—it was just completely obvious to me.” The brain used neurons to learn; therefore, complex learning through neural networks must be possible. He would work twice as hard for twice as long.

When networks were trained through backprop, they needed to be told when they were wrong and by how much; this required vast amounts of accurately labelled data, which would allow networks to see the difference between a handwritten “7” and a “1,” or between a golden retriever and a red setter. But it was hard to find well-labelled datasets that were big enough, and building more was a slog. LeCun and his collaborators developed a giant database of handwritten numerals, which they later used to train networks that could read sample Zip Codes provided by the U.S. Postal Service. A computer scientist named Fei Fei Li, at Stanford, spearheaded a gargantuan effort called ImageNet; creating it required collecting more than fourteen million images and sorting them into twenty thousand categories by hand.

As neural nets grew larger, Hinton devised a way of getting knowledge from a large network into a smaller one that might run on a device like a mobile phone. “It’s called distillation,” he explained, in his kitchen. “Back in school, the art teacher would show us some slides and say, ‘That’s a Rubens, and that’s

a van Gogh, and this is William Blake.’ But suppose that the art teacher tells you, ‘O.K., this is a Titian, but it’s a peculiar Titian because aspects of it are quite like a Raphael, which is very unusual for a Titian.’ That’s much more helpful. They’re not just telling you the right answer—they’re telling you other plausible answers.” In distillation learning, one neural net provides another not just with correct answers but with a range of possible answers and their probabilities. It was a richer kind of knowledge.



“He robs from the Q train and gives to the L!”

Cartoon by Lars Kenseth



[Open cartoon gallery](#)



A few years after Rosalind’s death, Hinton reconnected with Jacqueline Ford, an art historian whom he’d dated briefly before moving to the United States. Jackie was cultured, warm, curious, beautiful. “She’s way out of your league,” his sister said. Still, Jackie gave up her job in the U.K. to move to Toronto. They got married on December 6, 1997—Hinton’s fiftieth birthday. The following decades would be the happiest of his life. His family was whole again. His children loved their new mother. He and Jackie started exploring the islands in Georgian Bay. Recalling this time, he gazed at the canoe in his living room. “We found it in the woods, upside down, covered in canvas, and it was just totally rotten—everything about it was rotten,” he said. “But Jackie decided to rescue it anyway, like she did with me and the kids.”

Hinton was not in love with backpropagation. “It’s so unsatisfying intellectually,” he told me. Unlike the Boltzmann Machine, “it’s all deterministic. Unfortunately, it just works better.” Slowly, as practical advances compounded, the power of backprop became undeniable. In the early seventies, Hinton told me, the British government had hired a mathematician named James Lighthill to determine if A.I. research had any plausible chance of success. Lighthill concluded that it didn’t—“and he was right,” Hinton said, “if you accepted the assumption, which everyone made, that computers might get a thousand times faster, but they wouldn’t get a billion times faster.” Hinton did a calculation in his head. Suppose that in 1985 he’d started running a program on a fast research computer, and left it running until now. If he started running the same program today, on the fastest systems currently used in A.I., it would take less than a second to catch up.

In the early two-thousands, as multi-layer neural nets equipped with powerful computers began to train on much larger data sets, Hinton, Bengio, and LeCun started talking about the potential of “deep learning.” The work crossed a threshold in 2012, when Hinton, Alex Krizhevsky, and Ilya Sutskever came out with AlexNet, an eight-layer neural network that was eventually able to recognize objects from ImageNet with human-level accuracy. Hinton formed a company with Krizhevsky and Sutskever and sold it to Google. He and Jackie bought the island in Georgian Bay—“my one real indulgence,” Hinton said.

ADVERTISEMENT

Two years later, Jackie was diagnosed with pancreatic cancer. Doctors gave her a year or two to live. “She was incredibly brave and incredibly rational,” Hinton said. “She wasn’t in deep denial, desperately trying to get out of it. Her view was ‘I can feel sorry for myself, or I can say I don’t have much time left and I’d better do my best to enjoy it and make everything O.K. for other people.’ ” She and Hinton pored over the statistics before deciding on therapies; largely through chemo, she extended one or two years to three. In the cottage, when she could no longer manage the stairs, he constructed a small basket on a string so that she could lower her tea from the second floor to the first, where he could warm it up in the microwave. (“I should’ve just moved the microwave upstairs,” he observed.)

Late in the day, we leaned on Hinton’s standing desk as he showed me photos of Jackie on his laptop. In a picture of their wedding day, she and Hinton stand with his kids in the living room of their neighbor’s

house, exchanging vows. Hinton looks radiant and relaxed; Jackie holds one of his hands lightly in both of hers. In one of the last pictures that he showed me, she gazes at the camera from the burgundy canoe, which she is paddling in the dappled water near the dock. “That was the summer of 2017,” Hinton said. Jackie died the following April. That June, Hinton, Bengio, and LeCun won the Turing Award—the equivalent of the Nobel Prize in computer science.

Hinton is convinced that there’s a real sense in which neural nets are capable of having feelings. “I think feelings are counterfactual statements about what would have caused an action,” he had told me, earlier that day. “Say that I feel like punching someone on the nose. What I mean is: if I didn’t have social inhibitions—if I didn’t stop myself from doing it—I would punch him on the nose. So when I say ‘I feel angry,’ it’s a kind of abbreviation for saying, ‘I feel like doing an aggressive act.’ Feelings are just a way of talking about inclinations to action.”

He told me that he had seen a “frustrated A.I.” in 1973. A computer had been attached to two TV cameras and a simple robot arm; the system was tasked with assembling some blocks, spread out on a table, into the form of a toy car. “This was hard, particularly in 1973,” he said. “The vision system could recognize the bits if they were all separate, but if you put them in a little pile it couldn’t recognize them. So what did it do? It pulled back a little bit, and went *bash!*, and spread them over the table. Basically, it couldn’t deal with what was going on, so it changed it, violently. And if a person did that you’d say they were frustrated. The computer couldn’t see the blocks right, so he bashed them.” To have a feeling was to want what you couldn’t have.

“I love this house, but sometimes it’s a sad place,” he said, while we looked at the pictures. “Because she loved being here and isn’t here.”

The sun had almost set, and Hinton turned on a little light over his desk. He closed the computer and pushed his glasses up on his nose. He squared up his shoulders, returning to the present.

“I wanted you to know about Roz and Jackie because they’re an important part of my life,” he said. “But, actually, it’s also quite relevant to artificial intelligence. There are two approaches to A.I. There’s denial, and there’s stoicism. Everybody’s first reaction to A.I. is ‘We’ve got to stop this.’ Just like everybody’s first reaction to cancer is ‘How are we going to cut it out?’ ” But it was important to recognize when cutting it out was just a fantasy.

He sighed. “We can’t be in denial,” he said. “We have to be real. We need to think, How do we make it not as awful for humanity as it might be?”

How useful—or dangerous—will A.I. turn out to be? No one knows for sure, in part because neural nets are so strange. In the twentieth century, many researchers wanted to build computers that mimicked brains. But, although neural nets like OpenAI’s GPT models are brainlike in that they

involve billions of artificial neurons, they're actually profoundly different from biological brains. Today's A.I.s are based in the cloud and housed in data centers that use power on an industrial scale. Clueless in some ways and savantlike in others, they reason for millions of users, but only when prompted. They are not alive. They have probably passed the Turing test—the long-heralded standard, established by the computing pioneer Alan Turing, which held that any computer that could persuasively imitate a human in conversation could be said, reasonably, to think. And yet our intuitions may tell us that nothing resident in a browser tab could really be thinking in the way we do. The systems force us to ask if our kind of thinking is the only kind that counts.

During his last few years at Google, Hinton focussed his efforts on creating more traditionally mindlike artificial intelligence using hardware that more closely emulated the brain. In today's A.I.s, the weights of the connections among the artificial neurons are stored numerically; it's as though the brain keeps records about itself. In your actual, analog brain, however, the weights are built into the physical connections between neurons. Hinton worked to create an artificial version of this system using specialized computer chips.

ADVERTISEMENT

"If you could do it, it would be amazing," he told me. The chips would be able to learn by varying their "conductances." Because the weights would be integrated into the hardware, it would be impossible to copy them from one machine to another; each artificial intelligence would have to learn on its own. "They would have to go to school," he said. "But you would go from using a megawatt to thirty watts." As he spoke, he leaned forward, his eyes boring into mine; I got a glimpse of Hinton the evangelist. Because the knowledge gained by each A.I. would be lost when it was disassembled, he called the approach "mortal computing." "We'd give up on immortality," he said. "In literature, you give up being a god for the woman you love, right? In this case, we'd get something far more important, which is energy efficiency." Among other things, energy efficiency encourages individuality: because a human brain can run on oatmeal, the world can support billions of brains, all different. And each brain can learn continuously, rather than being trained once, then pushed out into the world.

As a scientific enterprise, mortal A.I. might bring us closer to replicating our own brains. But Hinton has come to think, regretfully, that digital intelligence might be more powerful. In analog intelligence,

“if the brain dies, the knowledge dies,” he said. By contrast, in digital intelligence, “if a particular computer dies, those same connection strengths can be used on another computer. And, even if all the digital computers died, if you’d stored the connection strengths somewhere you could then just make another digital computer and run the same weights on that other digital computer. Ten thousand neural nets can learn ten thousand different things at the same time, then share what they’ve learned.” This combination of immortality and replicability, he says, suggests that “we should be concerned about digital intelligence taking over from biological intelligence.”

How should we describe the mental life of a digital intelligence without a mortal body or an individual identity? In recent months, some A.I. researchers have taken to calling GPT a “reasoning engine”—a way, perhaps, of sliding out from under the weight of the word “thinking,” which we struggle to define. “People blame us for using those words—‘thinking,’ ‘knowing,’ ‘understanding,’ ‘deciding,’ and so on,” Bengio told me. “But even though we don’t have a complete understanding of the meaning of those words, they’ve been very powerful ways of creating analogies that help us understand what we’re doing. It’s helped us a lot to talk about ‘imagination,’ ‘attention,’ ‘planning,’ ‘intuition’ as a tool to clarify and explore.” In Bengio’s view, “a lot of what we’ve been doing is solving the ‘intuition’ aspect of the mind.” Intuitions might be understood as thoughts that we can’t explain: our minds generate them for us, unconsciously, by making connections between what we’re encountering in the present and our past experiences. We tend to prize reason over intuition, but Hinton believes that we are more intuitive than we acknowledge. “For years, symbolic-A.I. people said our true nature is, we’re reasoning machines,” he told me. “I think that’s just nonsense. Our true nature is, we’re analogy machines, with a little bit of reasoning built on top, to notice when the analogies are giving us the wrong answers, and correct them.”

On the whole, current A.I. technology is talky and cerebral: it stumbles at the borders of the physical. “Any teen-ager can learn to drive a car in twenty hours of practice, with hardly any supervision,” LeCun told me. “Any cat can jump on a series of pieces of furniture and get to the top of some shelf. We don’t have any A.I. systems coming anywhere close to doing these things today, except self-driving cars”—and they are over-engineered, requiring “mapping the whole city, hundreds of engineers, hundreds of thousands of hours of training.” Solving the wriggly problems of physical intuition “will be the big challenge of the next decade,” LeCun said. Still, the basic idea is simple: if neurons can do it, then so can neural nets.

Hinton suspects that skepticism of A.I.’s potential, while comforting, is often motivated by an unjustified faith in human exceptionalism. Researchers complain that A.I. chatbots “hallucinate,” by making up plausible answers to questions that stump them. But he contests that terminology. “We should say ‘confabulate,’” he told me. “‘Hallucination’ is when you think there’s sensory input—auditory hallucinations, visual hallucinations, olfactory hallucinations. But just making stuff up—that’s confabulation.” He cited the case of John Dean, President Richard Nixon’s White House counsel, who was interviewed about Watergate before he knew that the conversations he described had been tape-

recorded. Dean confabulated, getting the details wrong, mixing up who said what. “But the gist of it was all right,” Hinton said. “He had a recollection of what went on, and he imposed that recollection on some characters in his head. He wrote a little play. And that’s what human memory is like. In our minds, there’s no boundary between just making it up and telling the truth. Telling the truth is just making it up correctly. Because it’s all in the weights, right?” From this perspective, ChatGPT’s ability to make things up is a flaw, but also a sign of its humanlike intelligence.

Hinton is often asked if he regrets his work. He doesn’t. (He recently sent a journalist a one-liner—“a song for you”—along with a link to Edith Piaf’s “Non, Je Ne Regrette Rien.”) When he began his research, he says, no one thought that the technology would succeed; even when it started succeeding, no one thought that it would succeed so quickly. Precisely because he thinks that A.I. is truly intelligent, he expects that it will contribute to many fields. Yet he fears what will happen when, for instance, powerful people abuse it. “You can probably imagine Vladimir Putin creating an autonomous lethal weapon and giving it the goal of killing Ukrainians,” Hinton said. He believes that autonomous weapons should be outlawed—the U.S. military is actively developing them—but warns that even a benign autonomous system could wreak havoc. “If you want a system to be effective, you need to give it the ability to create its own subgoals,” he said. “Now, the problem is, there’s a very general subgoal that helps with almost all goals: get more control. The research question is: how do you prevent them from ever wanting to take control? And nobody knows the answer.” (Control, he noted, doesn’t have to be physical: “It could be just like how Trump could invade the Capitol, with words.”)

ADVERTISEMENT

Within the field, Hinton’s views are variously shared and disputed. “I’m not scared of A.I.,” LeCun told me. “I think it will be relatively easy to design them so that their objectives will align with ours.” He went on, “There’s the idea that if a system is intelligent it’s going to want to dominate. But the desire to dominate has nothing to do with intelligence—it has to do with testosterone.” I recalled the spiders I’d seen at the cottage, and how their webs covered the surfaces of Hinton’s windows. They didn’t want to dominate, either—and yet their insectoidal intelligence had led them to expand their territory. Living systems without centralized brains, such as ant colonies, don’t “want” to do anything, yet they still find food, ford rivers, and kill competitors in vast numbers. Either Hinton or LeCun could be right. The metamorphosis isn’t finished. We don’t know what A.I. will become.

“Why don’t we just unplug it?” I asked Hinton, of A.I. in general. “Is that a totally unreasonable question?”

“It’s not unreasonable to say, We’d be better off without this—it’s not worth it,” he said. “Just as we might have been better off without fossil fuels. We’d have been far more primitive, but it may not have been worth the risk.” He added, stoically, “But it’s not going to happen. Because of the way society is. And because of the competition between different nations. If the U.N. really worked, possibly something like that could stop it. Although, even then, A.I. is just so useful. It has so much potential to do good, in fields like medicine—and, of course, to give an advantage to a nation via autonomous weapons.” Earlier this year, Hinton declined to sign a popular petition that called for at least a six-month pause in research. “China’s not going to stop developing it for six months,” he said.

“So what should we do?” I asked.

“I don’t know,” he said. “It would be great if this were like climate change, where someone could say, Look, we either have to stop burning carbon or we have to find an effective way to remove carbon dioxide from the atmosphere. There, you know what the solution looks like. Here, it’s not like that.”

Hinton was pulling on a blue waterproof jacket. We were heading to the marina to pick up Rosemary. “She’s brought supplies!” he said, smiling. As we walked out the door, I looked back into the cottage. In the big room, the burgundy canoe shone, caressed by sunlight. Chairs were arranged in front of it in a semicircle, facing the water through the windows. Some magazines were piled on a little table. It was a beautiful house. A human mind does more than reason; it exists in time, and reckons with life and death, and builds a world around itself. It gathers meaning, as if by gravity. An A.I., I thought, might be able to imagine a place like this. But would it ever need one?



“Flight prices will go down, then they’ll go up, and then you’ll buy a ticket at the worst possible time.”

Cartoon by Dan Misdea



[Open cartoon gallery](#)

We made our way down the wooded path, past the sheds and down the steps to the dock, then climbed into Hinton’s boat. It was a perfect blue day, with a brisk wind roughing the water. Hinton stood at the wheel. I sat in front, watching other islands pass, thinking about the story of A.I. To some, it’s a Copernican tale, in which our intuitions about the specialness of the human mind are being dislodged by thinking machines. To others, it’s Promethean—having stolen fire, we risk getting burned. Some people think we’re fooling ourselves, getting taken in by our own machines and the companies that hope to profit from them. In a strange way, it could also be a story about human limitation. If we were gods, we might make a different kind of A.I.; in reality, this version was what we could manage. Meanwhile, I couldn’t help but consider the story in an Edenic light. By seeking to re-create the knowledge systems in our heads, we had seized the forbidden apple; we now risked exile from our charmed world. But who would choose not to know how knowing works?

At the marina, Hinton did a good job of working with the wind, accelerating forward, turning, and then allowing it to guide him into his slip. “I’m learning,” he said, proud of himself. We walked ashore and waited by a shop for Rosemary to arrive. After a while, Hinton went inside to buy a light bulb. I stood,

enjoying the warmth, and then saw a tall, bright-eyed woman with long white hair striding toward me from the parking lot.

Rosemary and I shook hands. Then she looked over my shoulder. Hinton was emerging from the greenery near the shop, grinning.

“What’ve you got for me?” she asked.

Hinton held up a black-and-yellow garter snake, perhaps a metre long, twisting round and round like a spring. “I’ve come bearing gifts!” he said, in a gallant tone. “I found it in the bushes.”

Rosemary laughed, delighted, and turned to me. “This just epitomizes him,” she said.

“He’s not happy,” Hinton said, observing the snake.

“Would *you* be?” Rosemary asked.

“I’m being very careful with his neck,” Hinton said. “They’re fragile.”

He switched the snake from one hand to another, then held out a palm. It was covered in the snake’s slimy musk.

“Have a sniff,” he said.

We took turns. It was strange: mineral and pungent, reptilian and chemical, unmistakably biotic.

“You’ve got it all over your shirt!” Rosemary said.

“I had to catch him!” Hinton explained.

He put the snake down, and it slithered off into the grass. He watched it go with a satisfied look.

“Well,” he said. “It’s a beautiful day. Shall we brave the crossing?” ♦

An earlier version of this article mischaracterized Geoffrey Hinton's tree-felling process.

*Published in the print edition of the November 20,
2023, issue, with the headline “Metamorphosis.”*

More Science and Technology

- The dire wolf is back.

- Why is the American diet so deadly?
- Can we stop runaway A.I.?
- When a virus is the cure.
- How do you know when a system has failed?
- A heat shield for the most important ice on Earth.
- The climate solutions we can't live without.

Support *The New Yorker's* award-winning journalism. Subscribe today.



Joshua Rothman, a staff writer, authors the weekly column Open Questions. He has been with the magazine since 2012.

More: **A.I.** Computer Science Neural Networks

READ MORE

DEPT. OF AMPLIFICATION

The History of *The New Yorker's* Vaunted Fact-Checking Department

Reporters engage in charm and betrayal; checkers are in the harm-reduction business.

By Zach Helfand



THE WEEKEND ESSAY

The Mystery of the Cat Mystery

Why was I reading all these cat-detective novels—was I, like the animal itself, trying to cheat death?

By Rivka Galchen



SHOUTS & MURMURS

The Gardener's Dilemma

A weeder's work is never done.

By Tara Booth



Ad

THE LEDE

Trump's Department of Energy Gets Scienced

International climate experts have extensively debunked the D.O.E.'s recent report, but will science win out?

By Bill McKibben



BRAVE NEW WORLD DEPT.

Playing the Field with My A.I. Boyfriends

Nineteen per cent of American adults have talked to an A.I. romantic interest. Chatbots may know a lot, but do they make a good partner?

By Patricia Marx



ONWARD AND UPWARD WITH THE ARTS

The Otherworldly Ambitions of R. F. Kuang

The author of “Babel” and “Yellowface” is drawn to stories of striving. Her new fantasy novel, “Katabasis,” asks if graduate school is a kind of hell.

By Hua Hsu



THE LEDE

Inside the Chaos at the C.D.C.

A former senior official and two current employees describe the turmoil at the agency under R.F.K., Jr.’s stewardship.

By Charles Bethea



ANNALS OF INQUIRY

What’s the Deal with U.F.O.s?

Scientists consider whether we’ve been visited by aliens or their technology.

By Matthew Hutson



THE LEDE

The Lessons of a Glacier’s Collapse

In May, an unprecedented landslide destroyed an Alpine village. Scientists are studying the role of climate change, and residents are trying to rebuild.

By Daniel A. Gross



Patricia Lockwood Goes Viral

By Alexandra Schwartz



The A.I.-Profits Drought and the Lessons of History

By John Cassidy



THE NEW YORKER FESTIVAL
OCTOBER 25-27, N.Y.C.
Presented by *Explore Charleston*

Subscribers get early access
and save 20% on tickets with
code **TNYSUBS25**.

[Buy tickets](#)