

BOOK CLUB SYNOPSIS

Artificial Intelligence: A Guide for Thinking Humans
Melanie Mitchell

Farrar, Straus and Giroux, 2019

CONTENTS

Prologue: Terrified | 3

Part I: Background | 6

1. The Roots of Artificial Intelligence | 7
2. Neural Networks and the Ascent of Machine Learning | 10
3. AI Spring | 12

Part II: Looking and Seeing | 16

4. Who, What, When, Where, Why | 17
5. ConvNets and ImageNet | 19
6. A Closer Look at Machines That Learn | 22
7. On Trustworthy and Ethical AI | 25

Part III: Learning to Play | 28

8. Rewards for Robots | 29
9. Game On | 31
10. Beyond Games | 33

Part IV: Artificial Intelligence Meets Natural Language | 36

11. Words, and the Company They Keep | 37
12. Translation as Encoding and Decoding | 40
13. Ask Me Anything | 42

Part V: The Barrier of Meaning | 44

14. On Understanding | 45
 15. Knowledge, Abstraction, and Analogy in Artificial Intelligence | 47
 16. Questions, Answers, and Speculations | 49
-

About World 50 | 52

Prologue: Terrified

In the prologue, Melanie Mitchell discusses a momentous meeting she had with Google's foremost minds about the topic of artificial intelligence in 2014, and the ways in which it inspired her to write a book about the development of AI.

In the opening, she describes her excitement in anticipation of the meeting since Google was “rapidly becoming an applied AI company” and notes the hiring of Ray Kurzweil, a renowned inventor and futurist who “promotes the idea of an AI Singularity, a time in the near future when computers will become smarter than humans.”

AI and *GEB*: Mitchell briefly provides a bit of her personal history, revealing that her entry into the world of artificial intelligence was largely inspired by AI legend Douglas Hofstadter's book *Gödel, Escher, Bach: An Eternal Golden Braid*—commonly referred to as *GEB*. It's a tome that's considered canon among computer science enthusiasts around the world, and it motivated her to quit her job teaching math at a prep school in the early 1980s and start taking computer science courses in Boston, Massachusetts. Serendipitously, Hofstadter was taking a sabbatical at the Massachusetts Institute of Technology (MIT) that year, and she became his research assistant. Mitchell then continued working for him while she earned a Ph.D. in Computer Science at the University of Michigan. She explains that they've stayed in touch ever since, and he was the one to invite her to Google.

Chess and the First Seed of Doubt: At the Google meeting, Hofstadter surprised the room by expressing tremendous concern about the speed of breakthroughs in AI at Google and around the world. “I am terrified. Terrified,” Hofstadter said at the meeting's opening. He explained how he had underestimated the pace of AI research in his book *GEB*, which was published in 1979. He had predicted, for example, that in the future, humanity would not see computer chess programs that could beat a person. But in 1997, IBM's Deep Blue machine beat the world chess champion Garry

Kasparov. Hofstadter admitted that he was “dead wrong” at the Google meeting.

In *GEB*, Hofstadter had also expressed doubt that a computer would be able to write beautiful music in the near future. But in the mid-1990s, he was shocked by a program called “Experiments in Musical Intelligence, or EMI (pronounced ‘Emmy’),” which was so sophisticated that faculty at the Eastman School of Music mistook one of its pieces for a composition by Chopin.

Hofstadter explained that he still doubted that the Singularity was imminent, but the implications of these developments disturbed him. If the Singularity happens, he said at the meeting, “[W]e will be superseded. We will be relics. We will be left in the dust.”

Why Is Hofstadter Terrified? His main concern wasn’t a rise of malicious machines, or even mass unemployment. “Instead, he was terrified that intelligence, creativity, emotions, and maybe even consciousness itself would be too *easy* to produce—that what he valued most in humanity would end up being nothing more than a ‘a bag of tricks.’”

I Am Confused: Moved by the meeting at Google, Mitchell was forced to consider whether she had underestimated the promise of AI. She notes that prominent thinkers are split on the question of just how dangerous AI is to humanity, with people like physicist Stephen Hawking deeply worried, and others such as roboticist Rodney Brooks saying that people overestimate the speed of AI advancement. Several Google researchers at the meeting predicted the emergence of general human-level AI within the next 30 years.

What This Book Is About: Hofstadter’s remarks were a “wake-up call” for Mitchell. The purpose of this book is to:

- Sort out how far AI has come and explain its disparate goals
- Consider how AI systems actually work and how successful they are

- Make sense of broader questions that have fueled debates about AI, such as:
 - What do “general human” and “superhuman” intelligence mean?
 - What aspects of our own intelligence do we cherish?
 - How would human-level AI challenge how we think about our own humanity?
 - How terrified should we be?

“My aim is for you to share in my own exploration and, like me, to come away with a clearer sense of what the field has accomplished and how much further there is to go before our machines can argue for their own humanity.”

PART

1

BACKGROUND

01

CHAPTER ONE

THE ROOTS OF ARTIFICIAL INTELLIGENCE

In this chapter, Mitchell provides a brief overview of the origins of the field of computer science and its earliest schools of thought.

Two Months and Ten Men at Dartmouth: The founding of AI can be traced back to a workshop at Dartmouth College in 1956, organized by a mathematician named John McCarthy. McCarthy coined the term “artificial intelligence,” and he and his fellow organizers said in their funding proposal for the workshop that they were guided by “the conjecture that every aspect of learning or any other feature of intelligence can be in principle so precisely described that a machine can be made to simulate it.” It was at that workshop that the future “‘big four’ pioneers of the field—McCarthy, [Marvin] Minsky, Allen Newell, and Herbert Simon—met and did some planning for the future.”

Definitions, and Getting On with It: According to Mitchell, AI has largely ignored distinctions between different kinds of intelligence and instead focused on two efforts: one scientific and one practical. “On the scientific side, AI researchers are investigating the mechanisms of ‘natural’ (that is, biological) intelligence by trying to embed it in computers. On the practical side, AI proponents simply want to create computer programs that perform tasks as well as or better than humans, without worrying about whether these programs are actually *thinking* in the way humans think.”

An Anarchy of Methods: Since the inception of AI, researchers have split into different methodological camps—the most prominent gap early on was between those who pursued “symbolic AI” and “subsymbolic AI.”

Symbolic AI: Advocates of symbolic AI argued that it's not necessary to build programs that mimic the brain to attain human intelligence. Instead, "general intelligence can be captured entirely by the right kind of symbol-processing program." These programs consist of "symbols, combinations of symbols, and rules and operations on symbols" that are created by and intelligible to humans. Symbolic AI dominated the field for its first three decades, "most notably in the form of *expert systems*, in which human experts devised rules for computer programs to use in tasks such as medical diagnosis and legal decision-making."

Subsymbolic AI: Perceptrons: Subsymbolic AI, by contrast, took inspiration from neuroscience and "sought to capture the sometimes-*unconscious* thought processes underlying what some have called fast perception, such as recognizing faces or identifying spoken words." A subsymbolic program is "essentially a stack of equations—a thicket of often hard-to-interpret operations on numbers ... designed to learn from data how to perform a task."

"An early example of a subsymbolic, brain-inspired AI program was the perceptron, invented in the late 1950s by the psychologist Frank Rosenblatt." A perceptron is a "simple program that makes a yes-or-no (1 or 0) decision based on whether the sum of its weighted inputs meets a threshold value." Rosenblatt's primary contribution was to develop an algorithm that enabled perceptrons to "learn via *conditioning*" through a concept known as "supervised learning."

But early on, perceptrons were deemed a dead end by the most prominent AI thinkers. "The fact that a perceptron's 'knowledge' consists of a set of numbers—namely, the weights and thresholds it has learned—means that it is hard to uncover the rules the perceptron is using in performing its recognition task." That, in turn, led many researchers to believe perceptrons had a very limited problem-solving capacity.

The Limitations of Perceptrons: By the mid-1970s, many of the general AI breakthroughs that had been promised early on had not come to fruition. This resulted in a drop in government funding for

research and marked the first instance of a now-familiar cycle in AI research: cycles of optimistic booms (“AI spring”) and then pessimistic busts (“AI winter”) during which public funds and venture capital tend to flow generously or dry up.

Easy Things Are Hard: “Marvin Minsky pointed out that in fact AI research had uncovered a paradox: ‘Easy things are hard.’” It was far easier to teach computers complex tasks such as diagnosing diseases than seemingly simple things such as conversing in natural language.

02

CHAPTER TWO

NEURAL NETWORKS AND THE ASCENT OF MACHINE LEARNING

In this chapter, Mitchell explains how the balance between the estimated power of symbolic and subsymbolic AI flipped in the 1980s.

Multilayer Neural Networks: In the late 1970s and early 1980s, several AI research groups rebutted earlier claims that subsymbolic AI had no future as a domain of usable AI. When perceptrons were augmented with a layer of simulated neurons—making them “multilayer neural networks”—and enhanced by a general learning algorithm called “back-propagation,” they became far more sophisticated than originally expected. Multilayer neural networks have been used for tasks such as “speech recognition, stock-market prediction, language translation, and music composition.”

Connectionism: Over the course of the 1980s, an increasingly prominent group of subsymbolic-focused researchers promoted the idea of “connectionism”—the notion that knowledge resides in weighted connections between neural units. They argued that “the key to intelligence was an appropriate computational architecture—inspired by the brain—and the ability of the system to learn on its own from data or from acting in the world.”

Around the same time, the expert systems—“symbolic AI approaches that rely on humans to create rules that reflect expert knowledge of a particular domain”—revealed themselves to be fragile. Researchers realized “how much the human experts writing the rules actually rely on subconscious knowledge—what you might call common sense—in order to act intelligently.” That common sense was not easy to translate into symbolic AI rules. Thus, funding moved away from symbolic AI ventures and instead toward neural networks.

Bad at Logic, Good at Frisbee: Subsymbolic AI has revealed itself to have distinct advantages and disadvantages. “Subsymbolic systems seem much better suited to perceptual or motor tasks for which humans can’t easily define rules. You can’t easily write down rules for identifying handwritten digits, catching a baseball, or recognizing your mother’s voice.” As the philosopher Andy Clark once said, subsymbolic systems are “bad at logic, good at Frisbee.” So far, efforts to combine subsymbolic and symbolic AI in a hybrid system have not been particularly successful.

The Ascent of Machine Learning: Machine learning—the development of algorithms that enable computers to learn from data—also became its own independent subdiscipline of AI, separate from symbolic AI. Following its development, “machine learning had its own cycles of optimism, government funding, start-ups, and overpromising, followed by the inevitable winters.”

03

CHAPTER THREE

AI SPRING

In this chapter, Mitchell provides an assessment of how close AI is to achieving human-level intelligence and casts doubt on predictions that such an accomplishment will happen soon.

Spring Fever: According to Mitchell, the world may very well be in the middle of an AI spring. Examples include Google Translate, Google's self-driving cars, Apple's personal assistant Siri, Facebook's facial recognition in photos and Skype's simultaneous translation between languages on video calls. Companies have become so obsessed with AI technology that the journalist Kevin Kelly has observed: "The business plans of the next 10,000 startups are easy to forecast: Take X and add AI."

AI: Narrow and General, Weak and Strong: The current boom in AI has inspired yet another round of sunny predictions that "general AI"—"AI that equals or surpasses humans in most ways"—is on the brink of arrival. But Mitchell points out that all instances of AI to date fall under the category of "narrow" or "weak" AI. Not nearly as derogatory as they sound, these terms "refer to a system that can perform only one narrowly defined task (or a small set of related tasks)." So far, there are no examples of "*strong, human-level, general, or full-blown* AI (sometimes called AGI, or artificial general intelligence)" in existence. This raises big questions, namely: How long will it take to integrate narrow intelligences, and would they amount to a "thinking machine"?

The Turing Test: The question of what makes for a thinking machine is hotly debated. The most famous test for determining whether or not a machine is thinking is British mathematician Alan

Turing’s “imitation game.” The game was designed to “cut through the Gordian knot of ‘simulated’ versus ‘actual’ intelligence.” In his game, which is now called the Turing test, a judge must interview two contestants—a computer and a human—and determine which is which. The test is handled through typed text. Turing posited that the physical appearance or composition of intelligence is irrelevant to the question of whether we consider it to actually think.

The Singularity: According to Kurzweil, the Singularity will likely transpire in 2045. His optimism about how soon it will happen is predicated on the idea of exponential growth in many areas of science and technology, including computer speed and neuroscience.

That said, Mitchell is highly skeptical of Kurzweil’s precise predictions about the future because most breakthroughs in AI have been narrow. She also points out that, while computer hardware has seen explosive growth in power, “computer *software* has not shown the same exponential progress.” Furthermore, Hofstadter remains ambivalent about Kurzweil’s predictions—skeptical that they’ll come to pass, yet also aware that his own past predictions were incorrect.

Wagering on the Turing Test: Mitchell critiques the very idea of being a futurist, noting that decades must pass before professional predictors can be held accountable for their mistakes; however, she points to one major scenario in which Kurzweil might be held accountable.

Kurzweil and software entrepreneur Mitchell Kapor have made a bet through a website called Long Bets. Kapor’s prediction: “By 2029 no computer—or ‘machine intelligence’—will have passed the Turing Test.” Kurzweil disagrees; he thinks it will happen.

Kapor and Kurzweil devised very specific conditions for the test to ensure that it’s truly rigorous, unlike tests in the past. Kapor believes a computer cannot achieve human intelligence without a physical body and emotion. Kurzweil, however, thinks that technological advancements can compensate for those obstacles.

Mitchell argues that a crucial premise for determining who is more likely to be right is estimating where AI research sits on the curve of exponential growth. “Is the current AI spring ... the first harbinger of a coming explosion? Or is it simply a waypoint on a slow, incremental growth curve that won’t result in human-level AI for at least another century? Or yet another AI bubble, soon to be followed by another AI winter?”

PART

2

LOOKING AND SEEING

04

CHAPTER FOUR

WHO, WHAT, WHEN, WHERE, WHY

Humans are able to process vast amounts of visual information instantly and effortlessly. For this reason, Mitchell argues that the ability of a machine to describe the contents of a photograph in great detail and to make deductions based on the objects within it would be “one of the first things we would require for general human-level AI.” This chapter is about AI researchers’ struggle to make progress toward that goal.

Easy Things Are Hard (Especially in Vision): Since the 1950s, AI researchers have been attempting to get computers to make sense of visual data. But those efforts haven’t progressed too far. “[A] program that can look at and describe photographs in the way humans do still seems far out of reach.” Training computers to look and to see like humans do is immensely difficult.

The Deep-Learning Revolution: Advances in deep learning in recent years are responsible for a “quantum leap” in the ability of machines to recognize objects in images and videos.

“*Deep learning* simply refers to methods for training ‘deep neural networks,’ which in turn refers to neural networks with more than one hidden layer.” Researchers have discovered that “the most successful deep networks are those whose structure mimics parts of the brain’s visual system.”

The Brain, the Neocognitron, and Convolutional Neural Networks: Mitchell then delves into a highly technical discussion of ConvNets, or convolutional neural networks, which were first proposed in the 1980s, but are “*the* driving force behind today’s deep-learning

revolution in computer vision.” A ConvNet consists of a sequence of layers of simulated neurons. “[T]he units in a ConvNet act as detectors for important visual features, each unit looking for its designated feature in a specific part of the visual field.”

Want to try out a well-trained ConvNet yourself? Take a photo of something and upload it to Google’s “search by image” engine. “Google will run a ConvNet on your image and, based on the resulting confidences (over thousands of possible object categories), will tell you its ‘best guess’ for the image.”

Training a ConvNet: Mitchell explains how training a ConvNet works, using the example of teaching one to recognize a given image as a dog or a cat.

“First, collect many example images of dogs and cats—this is your ‘training set.’ Also, create a file that gives a label for each image—that is, ‘dog’ or ‘cat.’”

“Your training program initially sets all the weights in the network to random values.” The program commences training, and “one by one, each image is given as the input to the network; the network performs its layer-by-layer calculations and finally outputs confidence percentages for ‘dog’ and ‘cat.’” For each image, the training program compares these output values to the “correct” values.

The training program then uses the back-propagation algorithm to change the weights throughout the network to make it more accurate. Over the course of many “epochs,” or cycles, of training, the network will get better at the task of recognition.

05

CHAPTER FIVE

CONVNETS AND IMAGENET

ConvNets fell out of favor in the mid-1990s because they struggled to scale up to more complex visual tasks. That changed in 2012 when a major breakthrough at a competition showed the world that they had a unique capacity to process images. Mitchell explains how far they’ve come—and how far they have to go.

Building ImageNet: Fei-Fei Li, a computer-vision professor at Princeton, came up with the idea of creating an image database structured according to the nouns in a database called WordNet, where each noun is linked to a large number of images containing examples of that noun. She then had an idea to use Amazon Mechanical Turk, “a marketplace for work that requires human intelligence,” to get hundreds of thousands of human workers to help build up its database of images corresponding to nouns. That’s how ImageNet was born.

The ImageNet Competitions: “In 2010, the ImageNet project launched the first ImageNet Large Scale Visual Recognition Challenge, in order to spur progress toward more general object-recognition algorithms.” Dozens of groups competed to show their program had the best ability to categorize an image under “a thousand possible categories.” That year, the top program was correct on 72% of the 150,000 test images. In 2011, that number improved to 74%. But in 2012, the winning program achieved a whopping 85% accuracy. What caused the leap in progress? While the first two winning programs used a support vector machine, the third winner used a ConvNet. The program, called AlexNet, “sent a jolt through the computer-vision and broader AI communities, suddenly waking people up to the potential power of ConvNets.” Deep learning became the hottest part of AI overnight.

What allowed ConvNets to dominate the competition, particularly after having been declared obsolete in the 1990s? “It turns out that the recent success of deep learning is due less to new breakthroughs in AI than to the availability of huge amounts of data ... and very fast parallel computer hardware.” Moreover, improvements in training methods allowed networks to be trained far more quickly than in the past.

The ConvNet Gold Rush: Starting in 2012, ConvNets trained with deep learning started springing up everywhere. Image search engines offered by Google and Microsoft were able to improve their “find similar images” features. Google’s Street View could recognize and blur out street addresses and license plates. Facebook could label users’ uploaded photos with the names of their friends. Twitter developed a filter that could screen tweets for pornographic images. Self-driving cars can use them to track pedestrians. Machines can use them to diagnose breast and skin cancer from medical images.

Have ConvNets Surpassed Humans at Object Recognition? Despite huge progress in object recognition, ConvNets have a long way to go before they can rival human ability. In ImageNet competitions, the programs are given five guesses per category; when they’re graded on how often they guess the right category for an image on their first attempt, the accuracy drops dramatically. ConvNets also tend to make many errors that a human wouldn’t. “[U]nlike humans they tend to miss objects that are small in the image, objects that have been distorted by color or contrast filters the photographer applied to the image, and ‘abstract representations’ of objects, such as a painting or statue of a dog, or a stuffed toy dog.” There’s also still much room for improvement in “localization”—the ability of a program to draw a box around the target object to confirm that the machine has “seen” the object rather than just respond to associative cues.

“Object recognition is not yet close to being ‘solved’ by artificial intelligence.” For machines to be able to describe what they see accurately, they’ll have to recognize how objects relate to one another and how they interact with the world. If the “objects” are living beings, they must be able to interpret things like emotions.

Why is this so hard? “It seems that visual intelligence isn’t easily separable from the rest of intelligence, especially general knowledge, abstraction, and language—abilities that, interestingly, involve parts of the brain that have many feedback connections to the visual cortex.”

06

CHAPTER SIX

A CLOSER LOOK AT MACHINES THAT LEARN

This chapter is about the ways that machines—and ConvNets, in particular—learn and how their learning processes differ from those of humans.

Learning on One's Own: ConvNets learn via supervised learning, processing training sets of examples over and over again. “In contrast, even the youngest children learn an open-ended set of categories and can recognize instances of most categories after seeing only a few examples.” Additionally, while ConvNets are passive learners, children ask questions, infer connections between concepts and actively explore the world.

ConvNets require humans to set their “hyperparameters” in order to even begin learning. Hyperparameters include things like the number of layers in a network and other technical details of the training process. Tuning the hyperparameters is no simple task—Mitchell describes it as often requiring a “kind of cabalistic knowledge” derived from years of experience. Demis Hassabis, a co-founder of Google DeepMind, described it, saying, “It’s almost like an art form to get the best out of these systems. ... There’s only a few hundred people in the world that can do that really well.”

Big Data: A ConvNet’s deep learning requires big data, and that tends to come from images, videos and other kinds of data that are uploaded to the internet. When people use services offered by tech companies like Google and Facebook, they’re providing direct examples in the form of images, videos, text and speech that can be used to train these companies’ AI programs. These improved programs attract more users and more data, helping advertisers on

those platforms to target ads more effectively. It's not just tech services; self-driving cars, for example, collect training examples from hours of videos collected from cameras mounted on actual cars.

The Long Tail: One challenge that ConvNets face in their reliance on supervised learning through limited data sets is reacting to unlikely scenarios. “[E]vents in the real world are usually predictable, but there remains a long tail of low-probability, unexpected occurrences.” If the unlikely situations don’t show up very often or at all in a ConvNet’s training set, it is more likely to make errors when faced with unexpected cases—for instance, Tesla vehicles’ Autopilot mode was unable to distinguish between lane markings and salt lines laid out on the highway before a snow storm.

The theoretical solution is for AI systems to learn more through unsupervised learning, which refers to a “broad group of methods for learning categories or actions without labeled data.” But to date, there are no widely successful AI methods for this kind of learning. “[H]umans also have a fundamental competence lacking in all current AI systems: common sense.”

What Did My Network Learn? Machines latch onto different things than humans do when looking at visual data. For example, neural networks might pick up on the fact that the backgrounds of photos of animals are often blurry, and might thus associate animals with blurry backgrounds instead of animals themselves, simply because of statistical associations. That’s called “overfitting,” and it means that a machine “can’t do a good job of applying what it learned to images that differ from those it was trained on.”

Biased AI: A number of AI programs that tag photos have resulted in embarrassing and offensive situations, such as in 2015, when Google’s automated photo-tagging feature mislabeled two African-American people as “gorillas.” According to Mitchell, “biases in AI training data reflect biases in our society.” Part of the problem is that the training sets rely on images that skew heavily white and male, since online images are disproportionately made up of celebrities and powerful people who are themselves

disproportionately white and male. These biases are easy to identify after a mistake, but they can be hard to anticipate.

Show Your Work: Deep neural networks—the bedrock of modern AI systems—cannot easily show their work. The billions of calculations that networks use to arrive at their conclusions are generally unfathomable to humans. “Even the humans who train deep networks generally cannot look under the hood and provide explanations for the decisions their networks make.” This has implications for how much AI can be trusted. “The fear is that if we don’t understand how they work, we can’t really trust them or predict the circumstances under which they will make errors.”

These concerns have been compounded by evidence that it’s fairly easy for AI researchers to dupe sophisticated ConvNets. Some have found, for example, that small changes to the pixels in an image that would be largely indiscernible to a human can trick a machine into thinking that what it had correctly identified as a school bus is an ostrich. “If deep-learning systems, so successful at computer vision and other tasks, can easily be fooled by manipulations to which humans are not susceptible, how can we say that these networks ‘learn like humans’ or ‘equal or surpass humans’ in their abilities?” In order to approach human ability, AI needs to exhibit a level of understanding of visual data that moves beyond responding to superficial cues.

07

CHAPTER SEVEN

ON TRUSTWORTHY AND ETHICAL AI

In this chapter, Mitchell surveys some of the biggest questions surrounding whether AI can be trusted to act responsibly and ethically.

Beneficial AI: Many current-day AI programs already provide a clear public service to humanity, such as GPS navigation, speech transcription, email spam filters, language translation and credit card fraud alerts. In the near future, the possibilities for expanding upon these features seem promising. Mitchell predicts that AI applications will likely be widespread in health care, potentially reducing the prevalence of human errors in diagnosis and treatment. Scientific modeling and data analysis will increasingly rely on AI tools, improving models of climate change, population growth and food science, for example. AI is also positioned to take over jobs that most humans consider overly dangerous, exhausting or boring. “If this actually happens, it could be a true boon for human well-being.”

The Great AI Trade-Off: Mitchell pushes back against AI researcher Andrew Ng’s declaration that “AI is the new electricity.” She agrees that AI could become as integrated into human life as electricity, but she thinks that electricity was far more understood than AI before it was commercialized. “We are good at predicting the behavior of electricity. This is not the case for many of today’s AI systems.”

The question of whether AI is a net good or net harm for society is widely debated. In 2018, the Pew Research Center canvassed “nearly one thousand ‘technological pioneers, innovators, developers,

business and policy leaders, researchers and activists,” and found divided results. While 63% predicted that advances in AI would leave humans better off by 2030, 37% disagreed.

The Ethics of Face Recognition: The benefits of facial recognition include programs that allow people to locate missing children or aid the police in tracking down criminal fugitives. But face recognition also poses many troubling downsides as well. One concern is privacy: There are websites that can tag photos of you with your name without your consent or knowledge. Another concern is liability: Face-recognition technology can make errors. “If your face is matched in error, you might be barred from a store or an airplane flight or wrongly accused of a crime.” Additionally, there is evidence to suggest that people of color are more likely to be the victim of such errors than white people.

For this reason, prominent tech leaders such as Brad Smith, Microsoft’s president and chief legal officer, have called for Congress to regulate facial recognition.

Regulating AI: Mitchell argues that regulation should not be left only in the hands of AI researchers and companies, since the problems surrounding AI “are social and political issues as much as they are technical ones.” She suggests that AI regulation should be modeled on the regulations that apply to the biological and medical sciences and involve “cooperation among government agencies, companies, nonprofit organizations, and universities.” What makes things more complicated is that there is no consensus on what issues should be considered priorities in regulating AI.

Moral Machines: Mitchell discusses one theoretical solution to the issue: “machine morality,” or AI programs with their own sense of morality. Though it has long been conceptualized as a solution, there are numerous reasons to cast doubt on its practicality. For one, AI programmers would need to find a way to ensure their systems’ values align with those of humans. “But what are the values of humans? Does it even make sense to assume that there are universal values that society shares?”

A good example of how difficult it will be to settle on AI morality can be found in debates over self-driving car ethics. A collection of 2016 surveys found that most participants said it would be morally preferable for a self-driving car to sacrifice one passenger rather than kill 10 pedestrians. But, when asked if they would personally buy a car programmed to make such a decision, the overwhelming majority of respondents said they wouldn't.

PART

3

LEARNING TO PLAY

08

CHAPTER EIGHT

REWARDS FOR ROBOTS

This chapter focuses on explaining the concept of reinforcement learning, which stands in contrast to supervised learning as a method for training AI.

Reinforcement learning is inspired by operant conditioning, which has been used for centuries on animals and humans. The idea is that “an *agent*—the learning program—performs *actions* in an *environment* (usually a computer simulation) and occasionally receives *rewards* from the environment.” Intermittent rewards are the only feedback the agent receives when learning. While an animal might receive a reward in the form of praise, a machine can be offered a machine equivalent of appreciation, such as positive numbers added to its memory. Reinforcement learning has been a part of the AI toolbox for decades, but it was often overshadowed by other methods. That changed in 2016, when reinforcement learning played a central role in a machine beating the best human players at the complex game of Go.

Training Your Robo-Dog: Mitchell uses the example of teaching a robo-dog to kick a soccer ball to demonstrate how reinforcement learning works. She dives into a highly technical explanation of how it works, but there are a few broader takeaways. For example: “[T]he more ‘intelligent’ you want your robot to be, the harder it is to manually specify rules for behavior. And of course, it’s impossible to devise a set of rules that will work in every situation.”

Mitchell explains how a robo-dog can learn flexible strategies on its own, “simply by performing actions in the world and occasionally receiving rewards (that is, *reinforcement*) without humans having to

manually write rules or directly teach the agent every possible circumstance.” The robo-dog can’t be taught too many things at once through rewards though, because that will produce behavior that Mitchell tentatively characterizes as “superstition”—that is, a situation where the robo-dog is led into “erroneously believing that a particular action can help cause a particular good or bad outcome.”

Stumbling Blocks in the Real World: Training a self-driving car through reinforcement learning is far, far harder. It faces an infinitely complex set of environments. It also is more difficult to carry out the learning process in real life. Mitchell explains that reinforcement-learning practitioners almost always use simulations for learning episodes to save time and resources, but she points out that simulations are far from perfect. “[T]he more complex and unpredictable the environment, the less successful are the attempts to transfer what is learned in simulation to the real world.”

09

CHAPTER NINE

GAME ON

In this chapter, Mitchell discusses how AI enthusiasts have long obsessed over creating programs that can beat humans at games and what breakthroughs have been made recently.

Video games hold the appeal that they do for AI researchers in part because, as DeepMind co-founder Hassabis explained, they're "like microcosms of the real world, but ... cleaner and more constrained."

In 2013, a group of Canadian AI researchers released a software platform called the Arcade Learning Environment that made it easy to test machine learning systems on 49 Atari video games. The AI group DeepMind used this platform in their work on reinforcement learning.

Deep Q-Learning: DeepMind coupled reinforcement learning—in particular, Q-learning—with deep neural networks to build a system that could learn to play Atari games. DeepMind was so successful at training Deep Q-Networks to beat expert human players that it was acquired by Google for \$650 million, less than a year after sharing its findings.

Deep Blue: Claude Shannon, the inventor of information theory, wrote in 1950 that a machine that surpasses humans at chess "will force us either to admit the possibility of mechanized thinking or to further restrict our concept of thinking." Mitchell believes that the latter has happened in the wake of IBM's Deep Blue defeating Garry Kasparov in 1997. "Superhuman chess playing is now seen as something that doesn't require general intelligence. ... [Deep Blue]

can't do anything but play chess, and it doesn't have any conception of what 'playing a game' or 'winning' means to humans."

The Grand Challenge of Go: The success of AlphaGo, a DeepMind program that learned to play the game Go via deep Q-learning, was a more impressive feat in the eyes of many analysts. Go is considerably more complex than chess: A chess player must choose from, on average, 35 possible moves from a given board position; a Go player has, on average, 250 such possibilities. AlphaGo's victory over the best human Go players in the world was astonishing to fans of the game. Hassabis noted that "the thing that separates out top Go players [is] their intuition" and that "what we've done with AlphaGo is to introduce with neural networks this aspect of intuition, if you want to call it that."

How AlphaGo Works: AlphaGo's "intuition" arises from its combination of deep Q-learning with a clever method called the "Monte Carlo tree search," a family of computer algorithms first used during the Manhattan Project to help design the atomic bomb. "The name comes from the idea that a degree of *randomness*—like that of the iconic spinning roulette wheel in the Monte Carlo Casino—can be used by a computer to solve difficult mathematical problems."

10

CHAPTER TEN

BEYOND GAMES

In this chapter, Mitchell evaluates how much recent AI breakthroughs with games actually matter beyond the world of gaming itself.

Generality and “Transfer Learning”: DeepMind’s programs that play Go and other games all involve separate convolutional neural networks that must be trained from scratch for their particular game. “Unlike humans, none of these programs can ‘transfer’ anything it has learned about one game to help it learn a different game.”

While programs are not yet able to transfer their learning to help them learn different yet related tasks, for humans, “transfer learning” is automatic. For example, learning to play ping pong helps humans develop skills that apply to some degree when learning to play tennis and badminton. Learning how to open one doorknob as a child allows the child to understand how to open most others. “[O]ur ability to generalize what we learn is a core part of what it means for us to *think*.”

“Without Human Examples or Guidance”: DeepMind once made a grand claim about its program AlphaGo: that its reinforcement learning approach required no human examples or guidance, and that it succeeded in the most challenging of domains.

Mitchell doesn’t buy it. “A few aspects of human guidance that were critical to its success include the specific architecture of its convolutional neural network, the use of Monte Carlo tree search, and the setting of the many hyperparameters that both of these

entail.” In other words, AlphaGo definitely benefited from human guidance.

Mitchell also is skeptical that the game Go represents the most challenging domain for AI to succeed in. Charades—a game that requires sophisticated visual, linguistic and social understanding “far beyond the abilities of any current AI system”—is a better example of something that represents the apex of difficulty for AI.

What Did These Systems Learn? There are indications that DeepMind’s programs don’t possess what could be considered a deep understanding of the games in which they excel. For example, DeepMind’s program for the Atari game *Breakout* plunges in performance if the paddle’s position on the screen is shifted by a mere few pixels—something that wouldn’t affect a human much, if at all. Changing the background color is also hugely disruptive for the program. “This hints that the system has not even learned the basic concept of *paddle*.”

While these deep Q-learning systems have achieved superhuman performance in specific narrow domains, they’re “lacking something absolutely fundamental to human intelligence. Whether it is called abstraction, domain generalization, or transfer learning, imbuing systems with this ability is still one of AI’s most important open problems.”

This is a huge reason to argue that programs like AlphaGo do not represent human intelligence. For humans, a crucial part of intelligence is learning to think and then applying that thinking flexibly. “It may sound strange to say, but in this way the lowliest kindergartner in the school chess club is smarter than AlphaGo.”

From Games to the Real World: There are inherent challenges to applying game-specific AI to problems in the real world. As Hofstadter explained, “If you look at situations in the world, they don’t come framed, like a chess game or a Go game. ... A situation in the world is something that has no boundaries at all; you don’t know what’s in the situation, what’s out of the situation.”

Mitchell uses the example of how difficult it would be to train a robot to do a simple real-world task like putting dirty dishes in the dishwasher. The robot would have to identify different objects, reason about the objects that it could and couldn't see, and respond to unexpected scenarios like a pet walking nearby. The kind of abilities that AlphaGo excels at couldn't be transferred to tasks that are inherently more unpredictable and less rule-bound.

PART

4

ARTIFICIAL INTELLIGENCE
MEETS NATURAL LANGUAGE

11

CHAPTER ELEVEN

WORDS, AND THE COMPANY THEY KEEP

This chapter introduces the concept of natural language processing (NLP), or “getting computers to deal with human language.” NLP includes things like speech recognition, web search, automated question answering and machine translation. Mitchell examines the progress that has been made so far on NLP and the challenges to come.

The Subtlety of Language: Question-answering systems are a central focus of NLP research—and this is an immensely difficult focus for AI researchers. Mitchell provides an example of a short story in which it’s easy for any human reader to deduce that a customer leaves a restaurant because he’s dissatisfied with the food, and lays out why it would be extremely difficult for a machine to understand it. “Today’s machines lack the detailed, interrelated concepts and commonsense knowledge that even a four-year-old child brings to understanding language.” The reason? “Language is inherently ambiguous, is deeply dependent on context, and assumes a great deal of background knowledge common to the communicating parties.” For those reasons, symbolic rule-based approaches don’t work nearly as well as statistical approaches, which lately have focused on deep learning.

Speech Recognition and the Last 10 Percent: Automatic speech recognition—the task of transcribing spoken language into text in real time—was deep learning’s first major success in NLP. According to Mitchell, this represents “AI’s most significant success to date in any domain.” In 2012, speech recognition made a huge leap from middling to “very nearly perfect in some circumstances.”

There is still room for improvement, and it won't be easy. There is a famous rule of thumb in engineering circles: "[T]he first 90 percent of the project takes 10 percent of the time and the last 10 percent takes 90 percent of the time." The last 10% for speech recognition means building programs that can deal with background noise, unfamiliar accents, unknown words, and the way that ambiguity and context shape language interpretation. Furthermore, it might require an "actual *understanding* of what the speaker is saying."

Classifying Sentiment: Other NLP tasks, such as question answering and language translation, have proven to be more difficult.

Applying neural networks to tasks involving things like variable-length sentences goes back to the 1980s with the introduction of recurrent neural networks, which were inspired by how the brain interprets sequences. But, there are limits to its power as long as neural networks are unable to capture semantic relationships between words and phrases.

The NLP research community has proposed several methods for encoding words in a way that would capture semantic relationships. All of them are based on an idea expressed by linguist John Firth in 1957: "You shall know a word by the company it keeps." Essentially, the key to teaching a machine to understand language is training it to understand context patterns.

To capture the web of associations between words requires mapping words along many dimensions through geometric concepts. In this method, the meaning of a word can be represented by its location in space relative to another word—that is, "by the coordinates defining its word vector." A great deal of NLP research has gone into figuring out if there's an algorithm that will properly place all of the words in a network's vocabulary in a semantic space that will accurately capture the many dimensions of each word's meaning.

Word2Vec: There have been advances in placing words in a geometric space—most notably through Google's "word2vec" (meaning "word to vector") method—but there is still considerable

work to be done. Analogy exercises demonstrate that machines are often not learning fundamental meanings of words, and they also end up clumsily replicating biases inherent in the language data that they draw from.

12

CHAPTER TWELVE

TRANSLATION AS ENCODING AND DECODING

In this chapter, Mitchell dissects the state of AI translation and explains why she remains skeptical of its approaching human parity.

Automated translation was one of the earliest AI projects, spurred in part by the ambition to translate between English and Russian during the Cold War. Early programs relied on complicated sets of human-specified rules, but they were quite brittle. In the 1990s, statistical machine translation, which relied on learning from data and probability tables, came to dominate the field. But in 2016, Google launched a new “neural machine-translation system,” which the company claims has achieved “the largest improvements to date for machine translation quality.” Mitchell argues that while big strides have been made, “the caliber of machine-translation systems remains far below that of capable human translators.”

Encoder, Meet Decoder: After Google Translate launched its neural machine translation in 2016, other companies were inspired to build similar programs based on the same encoder-decoder architecture. The jump in translation quality prompted lots of media buzz about how Google and its peers now had the ability to translate language as well as humans, and that their programs demonstrated an understanding of the underlying semantic meaning of the languages they were translating. But Mitchell says not to believe the hype.

She points out several issues with the commonly used criteria for determining that new translation services have reached human parity. First, there’s the issue of the standard automated program used to evaluate machine-translation quality, called bilingual

evaluation understudy (BLEU). “BLEU tends to overrate bad translations.” It uses an overly simplistic word-matching system to judge translations. Despite its flaws, it’s used because no better automatic options are available.

There are also issues with human evaluations of translation quality. Human judges only look at translated sentences—not paragraphs or longer passages, where a machine is more likely to fail. And when companies tout how human judges’ *average* ratings of their machine translations compare to human translations, the statistics might obscure how bad some of the worst sentences are.

Lost in Translation: Mitchell walks through an example of her own to illustrate the efficacy and limitations of Google Translate. She puts an English-language short story into Google Translate, and it produces some awkward translations into Chinese, French and Italian. While the gist of her story is captured, key descriptive and idiomatic language is completely—and comically—mistranslated. “I’m skeptical that machine translation will actually reach the level of human translators—except perhaps in narrow circumstances—for a long time to come.” The reason is that comprehending concepts behind the words is crucial to good translation. She quotes Hofstadter once again to emphasize the point: “Translation is far more complex than mere dictionary look-up and word rearranging. ... Translation involves having a mental model of the world being discussed.”

Translating Images to Sentences: “It’s hard not to be dazzled, and maybe a bit stunned, that a machine can take in images in the form of raw pixels” and produce remarkably accurate captions, as in the case of Google’s Show and Tell decoder network. But, automated image captioning also “suffers from the same kind of bipolar performance seen in language translation.” When a caption is well-executed, it’s very accurate, but when a caption misses the mark, it’s sometimes entirely inaccurate, betraying that it “misses the *meaning* of the photo.” Unfortunately, “the fundamental lack of *understanding* in caption-generating networks inevitably means that, as in language translation, these systems will remain untrustworthy.”

13

CHAPTER THIRTEEN

ASK ME ANYTHING

In this chapter, Mitchell discusses question-answering programs at length and then offers big-picture thoughts about why she believes NLP research has a long way to go.

The Story of Watson: Shortly before the advent of Siri and Alexa, IBM's question-answering program Watson famously won a game of *Jeopardy!* against human champions in 2011. Though obviously impressive, Watson made several notable errors during game play that revealed lack of comprehension of the questions posed to it. While the AI community debated just how sophisticated Watson might or might not be, IBM marketed Watson as a program that would be able to exceed human ability in several real-world domains, including medicine, law, finance, customer service and fashion design.

This serves as an important example of how corporations sometimes exaggerate the implications of their AI findings—and often dupe credulous journalists in the process. IBM prevented Watson from being exposed to third-party studies and routinely implied that it had developed one program that could be applied to disparate fields when, in reality, it was simply building new programs from scratch. “It turns out that the skills needed for *Jeopardy!* are not the same as those needed for question answering in, say, medicine or law. Real-world questions and answers in real-world domains have neither the simple short structure of *Jeopardy!* clues nor their well-defined responses. In addition, real-world domains, such as cancer diagnosis, lack a large set of perfect, cleanly labeled training examples, each with a single right answer, as was the case with *Jeopardy!*”

Reading Comprehension: NLP systems have made progress over time on the Stanford Question Answering Dataset (SQuAD), a test in which a computer must answer questions about a set of text. Because in a SQuAD test the answer must appear as a sentence or a phrase in the original text, the skill tested appears to be more akin to “answer extraction” than anything else. Some AI groups have developed programs that exceed human performance on SQuAD tests. But Mitchell is not convinced that those programs are engaging in reading comprehension. She quotes a *Washington Post* article that captures her beliefs: “The real miracle of reading comprehension ... is in reading between the lines—connecting concepts, reasoning with ideas and understanding implied messages that aren’t specifically outlined in the text.”

What Does *It* Mean? NLP systems are far from demonstrating reading comprehension. Programs that excel at SQuAD tests fare very poorly in elementary and middle school-level multiple-choice science questions that require processing and reasoning, and the best-performing NLP programs have only 61% accuracy on a set of 250 Winograd schemas—simple questions that are designed specifically to be easy for humans and hard for computers because they require the respondent to make a conceptual leap. Moreover, NLP systems have shown themselves to be acutely susceptible to adversarial attacks (professional hacking).

“While deep learning has produced some very significant advances in speech recognition, language translation, sentiment analysis, and other areas of NLP, human-level language processing remains a distant goal.”

PART

5

THE BARRIER OF MEANING

14

CHAPTER FOURTEEN

ON UNDERSTANDING

The mathematician and philosopher Gian-Carlo Rota once said, “I wonder whether or when AI will ever crash the barrier of meaning.” It’s a question that Mitchell has returned to throughout her thinking about the future of AI, and it underpins this chapter, in which she explores what philosophers, psychologists and AI researchers believe about how humans understand the world and find meaning within it.

The Building Blocks of Understanding: There are a number of comprehension-related qualities that humans naturally have that machines will need to exhibit if they’re ever to approach general intelligence. One is “*intuitive physics*”—the basic knowledge and beliefs humans share about objects and how they behave. Another is “*intuitive biology*”—knowledge about how living things differ from inanimate objects. And another is “*intuitive psychology*”—or the ability to sense and predict the feelings, beliefs and goals of other people.

“These core bodies of intuitive knowledge constitute the foundation for human cognitive development, underpinning all aspects of learning and thinking, such as our ability to learn new concepts from only a few examples, to generalize these concepts, and to quickly make sense of situations ... and decide what actions we should take in response.”

Predicting Possible Futures: A key component of understanding a situation is the ability to make predictions about what could happen next. “[Y]ou have what psychologists call mental models of important aspects of the world, based on your knowledge of physical

and biological facts, cause and effect, and human behavior.” Those mental models allow you to mentally simulate situations and conjure up future scenarios.

Understanding as Simulation: Psychologist Lawrence Barsalou has argued that mental simulations also play a critical role in understanding situations that we don’t directly participate in—situations we might watch, hear or read about. According to Barsalou, “As people comprehend a text, they construct simulations to represent its perceptual, motor, and affective content. Simulations appear central to the representation of meaning.”

Metaphors We Live By: The ability to perceive metaphors constitutes an integral component of human understanding. In the book *Metaphors We Live By*, written by linguist George Lakoff and philosopher Mark Johnson, the authors argue that “not only is our everyday language absolutely teeming with metaphors that are often invisible to us, but our understanding of essentially *all* abstract concepts comes about via metaphors based on core physical knowledge.”

Abstraction and Analogy: Constructing and using intuitive mental models rely on two fundamental capabilities: abstraction and analogy. “Abstraction is the ability to recognize specific concepts and situations as instances of a more general category.” The abstract effectively underlies *all* of our concepts in some form, whether recognizing your mother’s face as a baby or recognizing a musical style.

Closely linked to abstraction is analogy-making. Hofstadter defines analogy as “the perception of a common essence between two things.” Analogies are important because they are “what underlie our abstraction abilities and the formation of concepts.”

All of these concepts are crucial for building AI that harbors the kind of common sense and reasoning abilities that humans naturally have.

15

CHAPTER FIFTEEN

KNOWLEDGE, ABSTRACTION, AND ANALOGY IN ARTIFICIAL INTELLIGENCE

In this chapter, Mitchell explores current efforts to imbue AI systems with the capacity to think like a human in the ways mentioned in the previous chapter.

Core Knowledge for Computers: AI research groups are increasingly interested in the concept of giving machines common sense. In 2018, Microsoft’s co-founder Paul Allen doubled the budget of his AI research institute specifically to study common sense. And the Defense Advanced Research Projects Agency (DARPA) has published plans to provide substantial funding for research on common sense in AI.

Abstraction, Idealized: “[E]nabling machines to form humanlike conceptual abstractions is still an almost completely unsolved problem.”

A set of visual puzzles called the Bongard problems, which were formulated by the computer scientist Mikhail Bongard in the late 1960s, prompt the problem-solver to construct many highly abstract and subtle analogies and grapple with hard-to-verbalize concepts. While researchers have created AI programs that can solve a subset of specific Bongard problems, “none have shown that their methods could generalize in a humanlike way.”

In one study, ConvNets were trained to solve puzzles somewhat similar to the Bongard problems using 20,000 examples, yet they managed to perform “only slightly better than random guessing.” Meanwhile, humans given the same challenge scored close to 100%.

Metacognition in the Letter-String World: Metacognition is another realm in which AI researchers are attempting to make breakthroughs. The idea is to create a program that has the capacity to reflect on its own thinking or possesses mechanisms for self-perception. Metacognition would enhance a program's problem-solving abilities by allowing it to be aware of its specific avenues of failure and implicitly encouraging it to try new ones. A colleague of Mitchell's created a program that would allow an AI program to perceive patterns in its own actions, but it "only scratched the surface of humanlike self-reflection abilities." In general, metacognition is a realm in which AI has made very little progress.

"We Are Really, Really Far Away": Mitchell concludes the chapter with an extended meditation on a blog post by Andrej Karpathy, an expert in deep learning and computer vision who leads Tesla's AI efforts. The blog post centers on a photograph showing U.S. President Barack Obama furtively stepping on the back of a scale while a colleague attempts to weigh himself on it in a locker room, much to the amusement of onlookers. The post poses the question: "What would it take for a computer to understand this image as you or I do?"

A human looking at the photograph can quickly understand subtleties about the setting, the relationship between the people involved, and the irony and mood of the situation. Notably, the image prompts the viewer to slide into the mindset of the different subjects in the photo and observe their differing perceptions of the situation.

Karpathy wrote of the photo: "You are reasoning about [the] state of mind of people, and their view of the state of mind of another person. That's getting frighteningly meta." His admiration of the ability of humans to do this effortlessly is simultaneously meant to be cautionary about the long road ahead for AI. He concludes his blog post with the thought that perhaps the only way AI can reach general intelligence is through embodiment—having some kind of body that interacts with the world.

Mitchell finds the embodiment argument "increasingly compelling."

16

CHAPTER SIXTEEN

QUESTIONS, ANSWERS, AND SPECULATIONS

In her final chapter, Mitchell lists a series of commonly asked questions about the future of AI and makes her own predictions about them.

Question: How soon will self-driving be commonplace?

“Achieving full autonomy in driving essentially requires general AI, which likely won’t be achieved anytime soon. Cars with partial autonomy exist now, but are dangerous because the humans driving them don’t always pay attention. The most likely solution to this dilemma is to change the definition of *full autonomy*: allowing autonomous cars to drive only in specific areas—those that have created the infrastructure to ensure that the cars will be safe.”

Question: Will AI result in massive unemployment for humans?

There is “no question” that AI systems will replace humans in some jobs—some already have. But since the future abilities of AI technologies are unknown, it’s difficult to predict the overall effect on employment.

Conversely, AI could have beneficial effects on employment. Not only could it take away some undesirable jobs, it could help create more interesting ones. “Several people have pointed out that, historically, new technologies have created as many new kinds of jobs as they replace, and AI might be no exception. Perhaps AI will take away truck-driving jobs, but because of the need to develop AI ethics, the field will create new positions for moral philosophers.”

Question: Could a computer be creative?

AI has already produced works of great beauty, but Mitchell doesn’t

believe they constitute creative work because the programs lack awareness of what they've done. "[B]eing creative entails being able to understand and judge what one has created. In this sense of creativity, no existing computer can be said to be creative."

Question: How far are we from creating general human-level AI?

To quote Oren Etzioni, director of the Allen Institute for AI: "Take your estimate, double it, triple it, quadruple it. That's when."

AI researchers have a long track record of overestimating the speed and power of AI development. "We've seen that over the history of the field well-known AI practitioners have predicted that general AI will arrive in ten years, or fifteen, or twenty-five, or 'in a generation.' However, none of these predictions has come to pass." Most breakthroughs have involved very narrow intelligence, and general intelligence is still far-off.

On the notion of computers reaching "general human-level AI" and then becoming superintelligent, Mitchell is skeptical, primarily because she believes humans' limitations are inextricable from their general intelligence. "The cognitive limitations forced upon us by having bodies that work in the world, along with the emotions and 'irrational' biases that evolved to allow us to function as a social group and all the other qualities sometimes considered cognitive 'shortcomings,' are in fact precisely what enable us to be generally intelligent rather than narrow savants."

Question: How terrified should we be about AI?

According to Mitchell, superintelligent, malevolent AI shouldn't be a concern. Instead, she worries about AI being treated as smarter than it actually is. "I think the most worrisome aspect of AI systems in the short term is that we will give them too much autonomy without being fully aware of their limitations and vulnerabilities."

Question: What exciting problems in AI are still unsolved? "Nearly all of them." To quote MIT's Rodney Brooks: "When AI got started, the clear inspiration was human level performance and human level intelligence. I think that goal has been what attracted most

researchers into the field for the first sixty years. The fact that we do not have anything close to succeeding at those aspirations says not that researchers have not worked hard or have not been brilliant. It says that it is a very hard goal.”

About World 50

Founded in 2004, World 50 consists of private peer communities that enable CEOs and C-level executives at globally respected organizations to discover better ideas, share valuable experiences and build relationships that make a lasting impact. The busiest officer-level executives and their most promising future leaders trust World 50 to facilitate collaboration, conversation and counsel on the topics most crucial to leading, transforming and growing modern enterprises. Membership is by invitation only.

World 50 communities serve every significant enterprise leadership role. World 50 members reside in more than 27 countries on six continents and are leaders at companies that average more than \$30 billion in revenue.

World 50 is a private company that serves no other purpose than to accelerate the success of its members and their organizations. It is composed of highly curious associates who consider it a privilege to help leaders stay ahead.

