

White-Collar Jobs at Risk

On October 11, 2009, the Los Angeles Angels prevailed over the Boston Red Sox in baseball's American League play-offs and earned the right to face the New York Yankees for the league championship and entry into the World Series. It was an especially emotional win for the Angels because just six months earlier one of their most promising players, pitcher Nick Adenhart, had been killed by a drunk driver in a car accident. One sportswriter began an article describing the game like this:

Things looked bleak for the Angels when they trailed by two runs in the ninth inning, but Los Angeles recovered thanks to a key single from Vladimir Guerrero to pull out a 7-6 victory over the Boston Red Sox at Fenway Park on Sunday.

Guerrero drove in two Angels runners. He went 2-4 at the plate.

"When it comes down to honoring Nick Adenhart, and what happened in April in Anaheim, yes, it probably was the biggest hit [of my career]," Guerrero said. "Because I'm dedicating that to a former teammate, a guy that passed away."

Guerrero has been good at the plate all season, especially in day games. During day games Guerrero has a .794 OPS [on-base plus slugging]. He has hit five home runs and driven in 13 runners in 26 games in day games.¹

The author of that text is probably in no immediate danger of receiving any awards for his writing. The narrative is nonetheless a remarkable achievement: not because it is readable, grammatically correct, and an accurate description of the baseball game, but because the author is a computer program.

The software in question, called "StatsMonkey," was created by students and researchers at Northwestern University's Intelligent Information Laboratory. StatsMonkey is designed to automate sports reporting by transforming objective data about a particular game into a compelling narrative. The system goes beyond simply listing facts; rather, it writes a story that incorporates the same essential attributes that a sports journalist would want to include. StatsMonkey performs a statistical analysis to discern the notable events that occurred during a game; it then generates natural language text that summarizes the game's overall dynamic while focusing on the most important events and the key players who contributed to the story.

In 2010, the Northwestern University researchers who oversaw the team of computer science and journalism students who worked on StatsMonkey raised venture capital and founded a new company, Narrative Science, Inc., to commercialize the technology. The company hired a team of top computer scientists and engineers; then it tossed out the original StatsMonkey computer code and built a far more powerful and comprehensive artificial intelligence engine that it named "Quill."

Narrative Science's technology is used by top media outlets, including *Forbes*, to produce automated articles in a variety of areas, including sports, business, and politics. The company's software generates a news story approximately every thirty seconds, and many

of these are published on widely known websites that prefer not to acknowledge their use of the service. At a 2011 industry conference, *Wired* writer Steven Levy prodded Narrative Science co-founder Kristian Hammond into predicting the percentage of news articles that would be written algorithmically within fifteen years. His answer: over 90 percent.²

Narrative Science has its sights set on far more than just the news industry. Quill is designed to be a general-purpose analytical and narrative-writing engine, capable of producing high-quality reports for both internal and external consumption across a range of industries. Quill begins by collecting data from a variety of sources, including transaction databases, financial and sales reporting systems, websites, and even social media. It then performs an analysis designed to tease out the most important and interesting facts and insights. Finally, it weaves all this information into a coherent narrative that the company claims measures up to the efforts of the best human analysts. Once it's configured, the Quill system can generate business reports nearly instantaneously and deliver them continuously—all without human intervention.³ One of Narrative Science's earliest backers was In-Q-Tel, the venture capital arm of the Central Intelligence Agency, and the company's tools will likely be used to automatically transform the torrents of raw data collected by the US intelligence community into an easily understandable narrative format.

The Quill technology showcases the extent to which tasks that were once the exclusive province of skilled, university-educated professionals are vulnerable to automation. Knowledge-based work, of course, typically calls upon a wide range of capabilities. Among other things, an analyst may need to know how to retrieve information from a variety of systems, perform statistical or financial modeling, and then write understandable reports and presentations. Writing—which, after all, is at least as much art as it is science—might seem like one of the least likely tasks to be automated. Nevertheless, it

has been, and the algorithms are improving rapidly. Indeed, because knowledge-based jobs can be automated using only software, these positions may, in many cases, prove to be more vulnerable than lower-skill jobs that involve physical manipulation.

Writing also happens to be an area in which employers consistently complain that university graduates are deficient. One recent survey of US employers found that about half of newly hired two-year university graduates and over a quarter of those with four-year degrees were found to have poor writing—and in some cases even reading—skills.⁴ As for the UK, an international study carried out by the Organisation for Economic Co-operation and Development in 2014 found that only a quarter of today's graduates possess the literacy skills required for succeeding in the workplace.⁵ If intelligent software can, as Narrative Science claims, begin to rival the most capable human analysts, the future growth of knowledge-based employment is in doubt for all university graduates, especially the least prepared.

Big Data and Machine Learning

The Quill narrative-writing engine is just one of many new software applications being developed to leverage the enormous amounts of data now being collected and stored within businesses, organizations, and governments across the global economy. By one estimate, the total amount of data stored globally is now measured in thousands of exabytes (an exabyte is equal to a billion gigabytes), and that figure is subject to its own Moore's Law-like acceleration, doubling roughly every three years.⁶ Nearly all of that data is now stored in digital format and is therefore accessible to direct manipulation by computers. Google's servers alone handle about 24 petabytes (equal to a million gigabytes)—primarily information about what its millions of users are searching for—each and every day.⁷

All this data arrives from a multitude of different sources. On the Internet alone, there are website visits, search queries, emails, social

media interactions, and advertising clicks, to name just a few examples. Within businesses, there are transactions, customer contacts, internal communications, and data captured in financial, accounting, and marketing systems. Out in the real world, sensors continuously capture real-time operational data in factories, hospitals, cars, planes, and countless other consumer devices and industrial machines.

The vast majority of this data is what a computer scientist would call "unstructured." In other words, it is captured in a variety of formats that can often be difficult to match up or compare. This is very different from traditional relational database systems where information is arranged neatly in consistent rows and columns that make search and retrieval fast, reliable, and precise. The unstructured nature of big data has led to the development of new tools specifically geared toward making sense of information that is collected from a variety of sources. Rapid improvement in this area is just one more example of the way in which computers are, at least in a limited sense, beginning to encroach on capabilities that were once exclusive to human beings. The ability to continuously process a stream of unstructured information from sources throughout our environment is, after all, one of the things for which humans are uniquely adapted. The difference, of course, is that in the realm of big data, computers are able to do this on a scale that, for a person, would be impossible. Big data is having a revolutionary impact in a wide range of areas including business, politics, medicine, and nearly every field of natural and social science.

Major retailers are relying on big data to get an unprecedented level of insight into the buying preferences of individual shoppers, allowing them to make precisely targeted offers that increase revenue while helping to build customer loyalty. Police departments across the globe are turning to algorithmic analysis to predict the times and locations where crimes are most likely to occur and then deploying their forces accordingly. The city of London's data portal, London Datastore, allows residents to see both historical trends and

contemporary data in a range of areas that capture the ebb and flow of life in the capital. It includes data on energy consumption, crime, air quality, performance metrics for transportation, schools, health-care, and even the number of bikes hired each day in the Santander-sponsored hire scheme. Tools that provide new ways to visualize data collected from social media interactions as well as sensors built into doors, turnstiles, and escalators offer urban planners and city managers graphic representations of the way people move, work, and interact in urban environments, a development that may lead directly to more efficient and livable cities.

There is a potential dark side, however. The US superstore chain, Target, Inc., provided a far more controversial example of the ways in which vast quantities of extraordinarily detailed customer data can be leveraged. A data scientist working for the company found a complex set of correlations involving the purchase of about twenty-five different health and cosmetic products that were a powerful early predictor of pregnancy. The company's analysis could even estimate a woman's due date with a high degree of accuracy. Target began bombarding women with offers for pregnancy-related products at such an early stage that, in some cases, the women had often not yet shared the news with their immediate families. In an article published in early 2012, the *New York Times* reported one case in which the father of a teenage girl actually complained to store management about mail sent to the family's home—only to find out later that Target, in fact, knew more than he did.⁸ Some critics fear that this rather creepy story is only the beginning and that big data will increasingly be used to generate predictions that potentially violate privacy and perhaps even freedom.

The insights gleaned from big data typically arise entirely from correlation and say nothing about the causes of the phenomenon being studied. An algorithm may find that if A is true, B is likely also true. But it cannot say whether A causes B or vice versa—or if perhaps both A and B are caused by some external factor. In many

[Q45]

it doesn't
"throw"
↓
K.D. ish
androids?

cases, however, and especially in the realm of business where the ultimate measure of success is profitability and efficiency rather than deep understanding, correlation alone can have extraordinary value. Big data can offer management an unprecedented level of insight into a wide range of areas: everything from the operation of a single machine to the overall performance of a multinational corporation can potentially be analyzed at a level of detail that would have been impossible previously.

The ever-growing mountain of data is increasingly viewed as a resource that can be mined for value—both now and in the future. Just as extractive industries like oil and gas continuously benefit from technical advances, it's a good bet that accelerating computer power and improved software and analysis techniques will enable corporations to unearth new insights that lead directly to increased profitability. Indeed, that expectation on the part of investors is probably what gives data-intensive companies like Facebook such enormous valuations.

Machine learning—a technique in which a computer churns through data and, in effect, writes its own program based on the statistical relationships it discovers—is one of the most effective means of extracting all that value. Machine learning generally involves two steps: an algorithm is first trained on known data and is then unleashed to solve similar problems with new information. One ubiquitous use of machine learning is in email spam filters. The algorithm might be trained by processing millions of emails that have been pre-categorized as either spam or not. No one sits down and directly programs the system to recognize every conceivable typographic butchery of the word "Viagra." Instead, the software figures this out by itself. The result is an application that can automatically identify the vast majority of junk email and can also continuously improve and adapt over time as more examples become available. Machine learning algorithms based on the same basic principles recommend books on Amazon, movies on Netflix, and potential dates on Match.com.

One of the most dramatic demonstrations of the power of machine learning came when Google introduced its online language translation tool. Its algorithms used what might be called a “Rosetta Stone” approach to the problem by analyzing and comparing millions of pages of text that had already been translated into multiple languages. Google’s development team began by focusing on official documents prepared by the United Nations and then extended their effort to the Web, where the company’s search engine was able to locate a multitude of examples that became fodder for their voracious self-learning algorithms. The sheer number of documents used to train the system dwarfed anything that had come before. Franz Och, the computer scientist who led the effort, noted that the team had built “very, very large language models, much larger than anyone has ever built in the history of mankind.”

Q4S
On Google’s
Translations tool
Rosetta Stone

In 2005, Google entered its system in the annual machine translation competition held by the National Bureau of Standards and Technology, an agency within the US Commerce department that publishes measurement standards. Google’s machine learning algorithms were able to easily outperform the competition—which typically employed language and linguistic experts who attempted to actively program their translation systems to wade through the mire of conflicting and inconsistent grammatical rules that characterize languages. The essential lesson here is that, when datasets are large enough, the knowledge encapsulated in all that data will often trump the efforts of even the best programmers. Google’s system is not yet competitive with the efforts of skilled human translators, but it offers bidirectional translation between more than five hundred language pairs. That represents a genuinely disruptive advance in communication capability: for the first time in human history, nearly anyone can freely and instantly obtain a rough translation of virtually any document in any language.

While there are a number of different approaches to machine learning, one of the most powerful, and fascinating, techniques

Q4S on
How neural network
and forced learning
works

involves the use of artificial neural networks—or systems that are designed using the same fundamental operating principles as the human brain. The brain contains as many as 100 billion neuron cells—and many trillions of connections between them—but it’s possible to build powerful learning systems using far more rudimentary configurations of simulated neurons.

An individual neuron operates somewhat like the plastic pop-up toys that are popular with very young children. When the child pushes the button, a colorful figure pops up—perhaps a cartoon character or an animal. Press the button gently and nothing happens. Press it a bit harder and still nothing. But exceed a certain force threshold, and up pops the figure. A neuron works in essentially the same fashion, except that the activation button can be pressed by a combination of multiple inputs.

To visualize a neural network, imagine a machine in which a number of these pop-up toys are arranged on the floor in rows. Three mechanical fingers are poised over each toy’s activation button. Rather than having a figure pop up, the toys are configured so that when a toy is activated it causes several of the mechanical fingers in the next row of toys to press down on their own buttons. The key to the neural network’s ability to learn is that the force with which each finger presses down on its respective button can be adjusted.

To train the neural network, you feed known data into the first row of neurons. For example, imagine inputting visual images of handwritten letters. The input data causes some of the mechanical fingers to press down with varying force depending on their calibration. That, in turn, causes some of the neurons to activate and press down on buttons in the next row. The output—or answer—is gathered from the last row of neurons. In this case, the output will be a binary code identifying the letter of the alphabet that corresponds to the input image. Initially, the answer will be wrong, but our machine also includes a comparison and feedback mechanism. The output is compared to the known correct answer, and this automatically

results in adjustments to the mechanical fingers in each row, and that, in turn, alters the sequence of activating neurons. As the network is trained with thousands of known images, and then the force with which the fingers press down is continuously recalibrated, the network will get better and better at producing the correct answer. When things reach the point where the answers are no longer improving, the network has effectively been trained.

This is, in essence, the way that neural networks can be used to recognize images or spoken words, translate languages, or perform a variety of other tasks. The result is a program—essentially a list of all the final calibrations for the mechanical fingers poised over the neuron activation buttons—that can be used to configure new neural networks, all capable of automatically generating answers from new data.

Artificial neural networks were first conceived and experimented with in the late 1940s and have long been used to recognize patterns. However, the last few years have seen a number of dramatic breakthroughs that have resulted in significant advances in performance, especially when multiple layers of neurons are employed—a technology that has come to be called “deep learning.” Deep learning systems already power the speech recognition capability in Apple’s Siri and are poised to accelerate progress in a broad range of applications that rely on pattern analysis and recognition. A deep learning neural network designed in 2011 by scientists at the University of Lugano in Switzerland, for example, was able to correctly identify more than 99 percent of the images in a large database of traffic signs—a level of accuracy that exceeded that of human experts who competed against the system. Researchers at Facebook have likewise developed an experimental system—consisting of nine levels of artificial neurons—that can correctly determine whether two photographs are of the same person 97.25 percent of the time, even if lighting conditions and orientation of the faces vary. That compares with 97.53 percent accuracy for human observers.¹⁰

Geoffrey Hinton of the University of Toronto, one of the leading

researchers in the field, notes that deep learning technology “scales beautifully. Basically you just need to keep making it bigger and faster, and it will get better.”¹¹ In other words, even without accounting for likely future improvements in their design, machine learning systems powered by deep learning networks are virtually certain to see continued dramatic progress simply as a result of Moore’s Law.

Big data and the smart algorithms that accompany it are having an immediate impact on workplaces and careers as employers, particularly large corporations, increasingly track a myriad of metrics and statistics regarding the work and social interactions of their employees. Companies are relying ever more on so-called people analytics as a way to hire, fire, evaluate, and promote workers. The amount of data being collected on individuals and the work they engage in is staggering. Some companies capture every keystroke typed by every employee. Emails, phone records, web searches, database queries and accesses to files, entry and exit from facilities, and untold numbers of other types of data may also be collected—with or without the knowledge of workers.¹² While the initial purpose of all this data collection and analysis is typically more effective management and assessment of employee performance, it could eventually be put to other uses—including the development of software to automate much of the work being performed.

The big data revolution is likely to have two especially important implications for knowledge-based occupations. First, the data captured may, in many cases, lead to direct automation of specific tasks and jobs. Just as a person might study the historical record and then practice completing specific tasks in order to learn a new job, smart algorithms will often succeed using essentially the same approach. Consider, for example, that in November 2013 Google applied for a patent on a system designed to automatically generate personalized email and social media responses.¹³ The system works by first analyzing a person’s past emails and social media interactions. Based on what it learned, it would then automatically write responses to

future emails, Tweets, or blog posts, and it would do so employing the person's usual writing style and tone. It's easy to imagine such a system eventually being used to automate a great deal of routine communication.

Google's automated cars, which it first demonstrated in 2011, likewise provide important insight into the path that data-driven automation is likely to follow. Google didn't set out to replicate the way a person drives—in fact, that would have been beyond the current capabilities of artificial intelligence. Rather, it simplified the challenge by designing a powerful data processing system and then putting it on wheels. Google's cars navigate by relying on precision location awareness via GPS together with vast amounts of extremely detailed mapping data. The cars also, of course, have radars, laser range finders, and other systems that provide a continuous stream of real-time information and allow the car to adapt to new situations, such as a pedestrian stepping off the kerb. Driving may not be a white-collar profession, but the general strategy used by Google can be extended into a great many other areas. First, employ massive amounts of historical data in order to create a general "map" that will allow algorithms to navigate their way through routine tasks. Next, incorporate self-learning systems that can adapt to variations or unpredictable situations. The result is likely to be smart software that can perform many knowledge-based jobs with a high degree of reliability.

The second, and probably more significant, impact on knowledge jobs will occur as a result of the way big data changes organizations and the methods by which they are managed. Big data and predictive algorithms have the potential to transform the nature and number of knowledge-based jobs in organizations and industries across the board. The predictions that can be extracted from data will increasingly be used to substitute for human qualities such as experience and judgment. As top managers increasingly employ data-driven decision making powered by automated tools, there will be an ever-

shrinking need for an extensive human analytic and management infrastructure. Whereas today there is a team of knowledge workers who collect information and present analysis to multiple levels of management, eventually there may be a single manager and a powerful algorithm. Organizations are likely to flatten. Layers of middle management will evaporate, and many of the jobs now performed by both clerical workers and skilled analysts will simply disappear.

WorkFusion, a start-up company based in the New York City area, offers an especially vivid example of the dramatic impact that white-collar automation is likely to have on organizations. The company offers large corporations an intelligent software platform that almost completely manages the execution of projects that were once highly labor-intensive through a combination of crowd sourcing and automation.

The WorkFusion software initially analyzes the project to determine which tasks can be directly automated, which can be crowd sourced, and which must be performed by in-house professionals. It can then automatically post job listings to websites like Elance or Craigslist and manage the recruitment and selection of qualified freelance workers. Once the workers are on board, the software allocates tasks and evaluates performance. It does this in part by asking freelancers to answer questions to which it already knows the answer as an ongoing test of the workers' accuracy. It tracks productivity metrics like typing speed, and automatically matches tasks with the capabilities of individuals. If a particular person is unable to complete a given assignment, the system will automatically escalate that task to someone with the necessary skills.

While the software almost completely automates management of the project and dramatically reduces the need for in-house employees, the approach does, of course, create new opportunities for freelance workers. The story doesn't end there, however. As the workers complete their assigned tasks, WorkFusion's machine learning algorithms

Uber self-driving car
continuously look for opportunities to further automate the process. In other words, even as the freelancers work under the direction of the system, they are simultaneously generating the training data that will gradually lead to their replacement with full automation.

One of the company's initial projects involved retrieving the information necessary to update a collection of about 40,000 records. Previously, the corporate client had performed this process annually using an in-house staff at a cost of nearly \$4 per record. After switching to the WorkFusion platform, the client was able to update the records monthly at a cost of just 20 cents each. WorkFusion has found that, as the system's machine learning algorithms incrementally automate the process further, costs typically drop by about 50 percent after one year and still another 25 percent after a second year of operation.¹⁴

Cognitive Computing and IBM Watson

In the autumn of 2004, IBM executive Charles Lickel had dinner with a small team of researchers at a steakhouse in New York. Members of the group were taken aback when, at precisely seven o'clock, people suddenly began standing up from their tables and crowding around a television in the bar area. It turned out that Ken Jennings, who had already won more than fifty straight matches on the TV game show *Jeopardy!*, was once again attempting to extend his historic winning streak. Lickel noticed that the restaurant's patrons were so engaged that they abandoned their dinners, returning to finish their steaks only after the match concluded.¹⁵

That incident, at least according to many recollections, marked the genesis of the idea to build a computer capable of playing—and

beating the very best human champions at—*Jeopardy!* IBM had a long history of investing in high-profile projects called “grand challenges” that have showcased the company's technology while delivering the kind of organic marketing buzz that just can't be purchased at any price. In a previous grand challenge, more than seven years earlier, IBM's Deep Blue computer had defeated world chess champion Garry Kasparov in a six-game match—an event that forever anchored the IBM brand to the historic moment when a machine first achieved dominance in the game of chess. IBM executives wanted a new grand challenge that would captivate the public and position the company as a clear technology leader—and, in particular, combat any perception that the information technology innovation baton had passed from Big Blue to Google or to start-up companies emerging out of Silicon Valley.

As the idea for a *Jeopardy!*-based grand challenge that would culminate in a televised match between the best human competitors and an IBM computer began to gain traction with the company's top managers, the computer scientists who would have to actually build such a system initially pushed back aggressively. A *Jeopardy!* computer would require capabilities far beyond anything that had been demonstrated previously. Many researchers feared that the company risked failure or, even worse, embarrassment on national television.

Indeed, there was little reason to believe that Deep Blue's triumph at chess would be extensible to *Jeopardy!* Chess is a game with precise rules that operate within a strictly limited domain; it is almost ideally suited to a computational approach. To a significant extent, IBM succeeded simply by throwing powerful, customized hardware at the problem. Deep Blue was a refrigerator-sized system packed with processors that were designed specifically for playing chess. “Brute force” algorithms leveraged all that computing power by considering every conceivable move given the current state of the game. Then for each of those possibilities, the software looked many moves ahead, weighing potential actions by both players and iterating

* Stephen Baker's 2011 book, *Final Jeopardy: Man vs. Machine and the Quest to Know Everything*, offers a detailed account of the fascinating story that ultimately led to IBM's Watson.

through countless permutations—a laborious process that ultimately nearly always produced the optimal course of action. Deep Blue was fundamentally an exercise in pure mathematical calculation; all the information the computer needed to play the game was provided in a machine-friendly format it could process directly. There was no requirement for the machine to engage with its environment like a human chess player.

Jeopardy! presented a dramatically different scenario. Unlike chess, it is essentially open-ended. Nearly any subject that would be accessible to an educated person—science, history, film, literature, geography, and popular culture, to name just a few—is fair game. A computer would also face an entire range of daunting technical challenges. Foremost among these was the need to comprehend natural language: the computer would have to receive information and provide its responses in the same format as its human competitors. The hurdle for succeeding at *Jeopardy!* is especially high because the show has to be not just a fair contest but also an engaging form of entertainment for its millions of television viewers. The show's writers often intentionally weave humor, irony, and subtle plays on words into the clues—in other words, the kind of inputs that seem almost purposely designed to elicit ridiculous responses from a computer.

As an IBM document describing the Watson technology points out: “We have noses that run, and feet that smell. How can a slim chance and a fat chance be the same, but a wise man and a wise guy are opposites? How can a house burn up as it burns down? Why do we fill in a form by filling it out?”¹⁶ A *Jeopardy!* computer would have to successfully navigate routine language ambiguities of that type while also exhibiting a level of general understanding far beyond what you'd typically find in computer algorithms designed to delve into mountains of text and retrieve relevant answers. As an example, consider the clue “Sink it & you've scratched.” That clue was presented in a show televised in July 2000 and appeared on the top row of the game board—meaning that it was considered to

be very easy. Try searching for that phrase using Google, and you'll get page after page of links to web pages about removing scratches from stainless-steel kitchen sinks. (That's assuming you exclude the exact match on a website about past *Jeopardy!* matches.) The correct response—“What is the cue ball?”—completely eludes Google's keyword-based search algorithm.*

All these challenges were well understood by David Ferrucci, the artificial intelligence expert who eventually assumed leadership of the team that built Watson. Ferrucci had previously managed a small group of IBM researchers focused on building a system that could answer questions provided in natural language format. The team entered their system, which they named “Piquant,” in a contest run by the National Bureau of Standards and Technology—the same government agency that sponsored the machine language contest in which Google prevailed. In the contest, the competing systems had to churn through a defined set of about a million documents and come up with the answers to questions, and they were subject to no time limit at all. In some cases, the algorithms would grind away for several minutes before returning an answer.¹⁷ This was a dramatically easier challenge than playing *Jeopardy!*, where the clues could draw on a seemingly limitless body of knowledge and where the machine would have to generate consistently correct responses within a few seconds in order to have any chance against top human players.

Piquant (as well as its competitors) was not only slow; it was inaccurate. The system was able to answer questions correctly only about 35 percent of the time—not an appreciably better success rate than you could get by simply typing the question into Google's search engine.¹⁸ When Ferrucci's team tried to build a prototype *Jeopardy!*-playing system based on the Piquant project, the results were

* In *Jeopardy!* the clues are considered to be answers and the response must be phrased as a question for which the provided answer would be correct.

uniformly dismal. The idea that Piquant might someday take on a top *Jeopardy!* competitor like Ken Jennings seemed laughable. Ferrucci recognized that he would have to start from scratch—and that the project would be a major undertaking spanning as much as half a decade. He received the green light from IBM management in 2007 and set out to build, in his words, “the most sophisticated intelligence architecture the world has ever seen.”¹⁹ To do this, he drew on resources from throughout the company and put together a team consisting of artificial intelligence experts from within IBM as well as at top universities, including MIT and Carnegie Mellon.²⁰

Ferrucci's team, which eventually grew to include about twenty researchers, began by building a massive collection of reference information that would form the basis for Watson's responses. This amounted to about 200 million pages of information, including dictionaries and reference books, works of literature, newspaper archives, web pages, and nearly the entire content of Wikipedia. Next they collected historical data for the *Jeopardy!* quiz show. Over 180,000 clues from previously televised matches became fodder for Watson's machine learning algorithms, while performance metrics from the best human competitors were used to refine the computer's betting strategy.²¹ Watson's development required thousands of separate algorithms, each geared toward a specific task—such as searching within text; comparing dates, times, and locations; analyzing the grammar in clues; and translating raw information into properly formatted candidate responses.

Watson begins by pulling apart the clue, analyzing the words, and attempting to understand what exactly it should look for. This seemingly simple step can, in itself, be a tremendous challenge for a computer. Consider, for example, a clue that appeared in a category entitled “Lincoln Blogs” and was used in training Watson: “Secretary Chase just submitted this to me for the third time; guess what, pal. This time I'm accepting it.” In order to have any chance at responding correctly, the machine would first need to understand that the initial

instance of the word “this” acts as a placeholder for the answer it should seek.²²

human: 20 ?
Once it has a basic understanding of the clue, Watson simultaneously launches hundreds of algorithms, each of which takes a different approach as it attempts to extract a possible answer from the massive corpus of reference material stored in the computer's memory. In the example above, Watson would know from the category that “Lincoln” is important, but the word “blogs” would likely be a distraction: unlike a human, the machine wouldn't comprehend that the show's writers were imagining Abraham Lincoln as a blogger.

As the competing search algorithms reel in hundreds of possible answers, Watson begins to rank and compare them. One technique used by the machine is to plug the potential answer into the original clue so that it forms a statement, and then go back out to the reference material and look for corroborating text. So if one of the search algorithms manages to come up with the correct response “resignation,” Watson might then search its dataset for a statement something like “Secretary Chase just submitted resignation to Lincoln for the third time.” It would find plenty of close matches, and the computer's confidence in that particular answer would rise. In ranking its candidate responses, Watson also relies on reams of historical data; it knows precisely which algorithms have the best track records for various types of questions, and it listens far more attentively to the top performers. Watson's ability to rank correctly worded natural language answers and then determine whether or not it has sufficient confidence to press the *Jeopardy!* buzzer is one of the system's defining characteristics, and a quality that places it on the frontier of artificial intelligence. IBM's machine “knows what it knows”—something that comes easily to humans but eludes nearly all computers when they delve into masses of unstructured information intended for people rather than machines.

Watson prevailed over *Jeopardy!* champions Ken Jennings and Brad Rutter in two matches televised in February 2011, giving IBM

the massive publicity surge it hoped for. Well before the media frenzy surrounding that remarkable accomplishment began to fade, a far more consequential story began to unfold: IBM launched its campaign to leverage Watson's capabilities in the real world. One of the most promising areas is in medicine. Repurposed as a diagnostic tool, Watson offers the ability to extract precise answers from a staggering amount of medical information that might include textbooks, scientific journals, clinical studies, and even physicians' and nurses' notes for individual patients. No single doctor could possibly approach Watson's ability to delve into vast collections of data and discover relationships that might not be obvious—especially if the information is drawn from sources that cross boundaries between medical specialties.* By 2013, Watson was helping to diagnose problems and refine patient treatment plans at major medical facilities, including the Cleveland Clinic and the University of Texas's MD Anderson Cancer Center.

As a part of their effort to turn Watson into a practical tool, IBM researchers confronted one of the primary tenets of the big data revolution: the idea that prediction based on correlation is sufficient, and that a deep understanding of causation is usually both unachievable and unnecessary. A new feature they named "WatsonPaths" goes beyond simply providing an answer and lets researchers see the specific sources Watson consulted, the logic it used in its evaluation, and the inferences it made on its way to generating an answer. In other words, Watson is gradually progressing toward offering more

* According to Stephen Baker's 2011 book *Final Jeopardy*, the Watson project leader, David Ferrucci, struggled with intense pain in one of his teeth for months. After multiple visits to dentists and what ultimately proved to be a completely unnecessary root canal, Ferrucci was finally—largely by happenstance—referred to a doctor in a medical specialty unrelated to dentistry, and the problem was solved. The specific condition was also described in a relatively obscure medical journal article. It was not lost on Ferrucci that a machine like Watson might have produced the correct diagnosis almost instantly.

insight into why something is true. WatsonPaths is also being used as a tool to help train medical students in diagnostic techniques. Less than three years after a team of humans succeeded in building and training Watson, the tables have—at least to a limited extent—been turned, and people are now learning from the way the system reasons when presented with a complex problem.²³

Other obvious applications for the Watson system are in areas like customer service and technical support. In 2013, IBM announced that it would work with Fluid, Inc., a major provider of online shopping services and consulting. The project aims to let online shopping sites replicate the kind of personalized, natural language assistance you would get from a knowledgeable sales assistant in a retail shop. If you're going camping and need a tent, you'd be able to say something like "I am taking my family camping in the Lake District in October and I need a tent. What should I consider?" You'd then get specific tent recommendations, as well as pointers to other items that you might not have considered.²⁴ As I suggested in Chapter 1, it is only a matter of time before capability of that type becomes available via smartphones and shoppers are able to access conversational, natural language assistance while in high street shops.

MD Buyline, Inc., a company that specializes in providing information and research about the latest healthcare technology to hospitals, likewise plans to use Watson to answer the far more technical questions that come up when hospitals need to purchase new equipment. The system would draw on product specifications, prices, and clinical studies and research to make specific and instant recommendations to doctors and procurement managers.²⁵ Watson is also looking for a role in the financial industry, where the system may be poised to provide personalized financial advice by delving into a wealth of information about specific customers as well as general market and economic conditions. The deployment of Watson in customer service call centers is perhaps the area with the most disruptive near-term potential, and it is likely no coincidence that within a year of Watson's triumph on *Jeopardy!*,

IBM was already working with Citigroup to explore applications for the system in the company's massive retail banking operation.²⁶

IBM's new technology is still in its infancy. Watson—as well as the competing systems that are certain to eventually appear—have the potential to revolutionize the way questions are asked and answered, as well as the way information analysis is approached, both internal to organizations and in engagements with customers. There is no escaping the reality, however, that a great deal of the analysis performed by systems of this type would otherwise have been done by human knowledge workers.

Building Blocks in the Cloud

In November 2013, IBM announced that its Watson system would move from the specialized computers that hosted the system for the *Jeopardy!* matches to the cloud. In other words, Watson would now reside in massive collections of servers connected to the Internet. Developers would be able to link directly to the system and incorporate IBM's revolutionary cognitive computing technology into custom software applications and mobile apps. This latest version of Watson was also more than twice as fast as its *Jeopardy!*-playing predecessor. IBM envisions the rapid emergence of an entire ecosystem of smart, natural language applications—all carrying the “Powered by Watson” label.²⁷

The migration of leading-edge artificial intelligence capability into the cloud is almost certain to be a powerful driver of white-collar automation. Cloud computing has become the focus of intense competition among major information technology companies, including Amazon, Google, and Microsoft. Google, for example, offers developers a cloud-based machine learning application as well as a large-scale computational engine that lets developers solve huge, computationally intensive problems by running programs on massive supercomputer-like networks of servers. Amazon is the industry leader in providing

cloud computing services. Cycle Computing, a small company that specializes in large-scale computing, was able to solve a complex problem that would have taken over 260 years on a single computer in just 18 hours by utilizing tens of thousands of the computers that power Amazon's cloud service. The company estimates that prior to the advent of cloud computing, it would have cost as much as \$68 million to build a supercomputer capable of taking on the problem. In contrast, it's possible to rent 10,000 servers in the Amazon cloud for about \$90 per hour.²⁸

Just as the field of robotics is poised for explosive growth as the hardware and software components used in designing the machines become cheaper and more capable, a similar phenomenon is unfolding for the technology that powers the automation of knowledge work. When technologies like Watson, deep learning neural networks, or narrative-writing engines are hosted in the cloud, they effectively become building blocks that can be leveraged in countless new ways. Just as hackers quickly figured out that Microsoft's Kinect could be used as an inexpensive way to give robots three-dimensional machine vision, developers will likewise find unforeseen—and perhaps revolutionary—applications for cloud-based software building blocks. Each of these building blocks is in effect a “black box”—meaning that the component can be used by programmers who have no detailed understanding of how it works. The ultimate result is sure to be that groundbreaking AI technologies created by teams of specialists will rapidly become ubiquitous and accessible even to amateur coders.

While innovations in robotics produce tangible machines that are often easily associated with particular jobs (a hamburger-making robot or a precision assembly robot, for example), progress in software automation will likely be far less visible to the public; it will often take place deep within corporate walls, and it will have more holistic impacts on organizations and the people they employ. White-collar automation will very often be the story of information

technology consultants descending on large organizations and building completely custom systems that have the potential to revolutionize the way the business operates, while at the same time eliminating the need for potentially hundreds or even thousands of skilled workers. Indeed, one of IBM's stated motivations for creating the Watson technology was to offer its consulting division—which, together with software sales, now accounts for the vast majority of the company's revenues—a competitive advantage. At the same time, entrepreneurs are already finding ways to use the same cloud-based building blocks to create affordable automation products geared toward small or medium-sized businesses.

Cloud computing has already had a significant impact on information technology jobs. During the 1990s' tech boom, huge numbers of well-paying jobs were created as businesses and organizations of all sizes needed IT professionals to administer and install personal computers, networks, and software. By the first decade of the twenty-first century, however, the trend began to shift as companies were increasingly outsourcing many of their information technology functions to huge, centralized computing hubs.

The massive facilities that host cloud computing services benefit from enormous economies of scale, and the administrative functions that once kept armies of skilled IT workers busy are now highly automated. Facebook, for example, employs a smart software application called "Cyborg" that continuously monitors tens of thousands of servers, detects problems, and in many cases can perform repairs completely autonomously. A Facebook executive noted in November 2013 that the Cyborg system routinely solves thousands of problems that would otherwise have to be addressed manually, and that the technology allows a single technician to manage as many as 20,000 computers.²⁹

Cloud computing data centers are often built in relatively rural areas where land and, especially, electric power are plentiful and cheap. In the US, states and local governments compete intensively for the facilities, offering companies like Google, Facebook,

and Apple generous tax breaks and other financial incentives. Their primary objective, of course, is to create lots of jobs for local residents—but such hopes are rarely realized. In 2011, the *Washington Post*'s Michael Rosenwald reported that a colossal, billion-dollar data center built by Apple, Inc., in the town of Maiden, North Carolina, had created only fifty full-time positions. Disappointed residents couldn't "comprehend how expensive facilities stretching across hundreds of acres can create so few jobs."³⁰ The explanation, of course, is that algorithms like Cyborg are doing the heavy lifting.

The impact on employment extends beyond the data centers themselves to the companies that leverage cloud computing services. In 2012, Roman Stanek, the CEO of Good Data, a San Francisco company that uses Amazon's cloud services to perform data analysis for about 6,000 clients, noted that "[b]efore, each [client] company needed at least five people to do this work. That is 30,000 people. I do it with 180. I don't know what all those other people will do now, but this isn't work they can do anymore. It's a winner-takes-all consolidation."³¹

The evaporation of thousands of skilled information technology jobs is likely a precursor for a much more wide-ranging impact on knowledge-based employment. As Netscape co-founder and venture capitalist Marc Andreessen famously said, "Software is eating the world." More often than not, that software will be hosted in the cloud. In 2014, 18 percent of the UK's small or medium sized enterprises were using cloud services. That percentage is expected to rise to nearly 50 percent by the end of 2015 as the UK cloud market grows to over £6 billion.³² From that vantage point it will eventually be poised to invade virtually every workplace and swallow up nearly any white-collar job that involves sitting in front of a computer manipulating information.

Algorithms on the Frontier

If there is one myth regarding computer technology that ought to be swept into the dustbin it is the pervasive belief that computers can do only what they are specifically programmed to do. As we've seen, machine learning algorithms routinely churn through data, revealing statistical relationships and, in essence, writing their own programs on the basis of what they discover. In some cases, however, computers are pushing even further and beginning to encroach into areas that nearly everyone assumes are the exclusive province of the human mind: machines are starting to demonstrate curiosity and creativity.

In 2009, Hod Lipson, the director of the Creative Machines Lab at Cornell University, and PhD student Michael Schmidt built a system that has proved capable of independently discovering fundamental natural laws. Lipson and Schmidt started by setting up a double pendulum—a contraption that consists of one pendulum attached to, and dangling below, another. When both pendulums are swinging, the motion is extremely complex and seemingly chaotic. Next they used sensors and cameras to capture the pendulum's motion and produce a stream of data. Finally, they gave their software the ability to control the starting position of the pendulum; in other words, they created an artificial scientist with the ability to conduct its own experiments.

They turned their software loose to repeatedly release the pendulum and then sift through the resulting motion data and try to figure out the mathematical equations that describe the pendulum's behavior. The algorithm had complete control over the experiment; for each repetition, it decided how to position the pendulum for release, and it did not do this randomly—it performed an analysis and then chose the specific starting point that would likely provide the most insight into the laws underlying the pendulum's motion. Lipson notes that the system “is not a passive algorithm that sits back, watching. It asks questions. That's curiosity.”³³ The program, which they later named “Eureqa,” took only a few hours to come up with a number of

physical laws describing the movement of the pendulum—including Newton's Second Law—and it was able to do this without being given any prior information or programming about physics or the laws of motion.

Eureqa uses genetic programming, a technique inspired by biological evolution. The algorithm begins by randomly combining various mathematical building blocks into equations and then testing to see how well the equations fit the data.* Equations that fail the test are discarded, while those that show promise are retained and recombined in new ways so that the system ultimately converges on an accurate mathematical model.³⁴ The process of finding an equation that describes the behavior of a natural system is by no means a trivial exercise. As Lipson says, “[P]reviously, coming up with a predictive model could take a [scientist's] whole career.”³⁵ Schmidt adds that “[p]hysicists like Newton and Kepler could have used a computer running this algorithm to figure out the laws that explain a falling apple or the motion of the planets with just a few hours of computation.”³⁶

When Schmidt and Lipson published a paper describing their algorithm, they were deluged with requests for access to the software from other scientists, and they decided to make Eureqa available over the Internet in late 2009. The program has since produced a number of useful results in a range of scientific fields, including a simplified equation describing the biochemistry of bacteria that scientists are still struggling to understand.³⁷ In 2011, Schmidt founded Nutonian, Inc., a Boston-area start-up company focused on commercializing

* This is significantly more advanced than the commonly used statistical technique known as “regression.” With regression (either linear or nonlinear), the form of the equation is set in advance, and the equation's parameters are optimized so as to fit the data. The Eureqa program, in contrast, is able to independently determine equations of any form using a variety of mathematical components including arithmetic operators, trigonometric and logarithmic functions, constants, etc.

Eureqa as a big data analysis tool for both business and academic applications. One result is that Eureqa—like IBM's Watson—is now hosted in the cloud and is available as an application building block to other software developers.

X Most of us quite naturally tend to associate the concept of creativity exclusively with the human brain, but it's worth remembering that the brain itself—by far the most sophisticated invention in existence—is the product of evolution. Given this, perhaps it should come as no surprise that attempts to build creative machines very often incorporate genetic programming techniques. Genetic programming essentially allows computer algorithms to design themselves through a process of Darwinian natural selection. Computer code is initially generated randomly and then repeatedly shuffled using techniques that emulate sexual reproduction. Every so often, a random mutation (X) is thrown in to help drive the process in entirely new directions. As new algorithms evolve, they are subjected to a fitness test that leads to either their survival, or—far more often—their demise. Computer scientist and consulting Stanford professor John Koza is one of the leading researchers in the field and has done extensive work using genetic algorithms as “automated invention machines.” Koza has isolated at least seventy-six cases where genetic algorithms have produced designs that are competitive with the work of human engineers and scientists in a variety of fields, including electric circuit design, mechanical systems, optics, software repair, and civil engineering. In most of these cases, the algorithms have replicated existing designs, but there are at least two instances where genetic programs have created new, patentable inventions.³⁸ Koza argues that genetic algorithms may have an important advantage over human designers

* In addition to his work in genetic programming, Koza is the inventor of the scratch-off lottery ticket and the originator of the “constitutional workaround” idea to elect US presidents by popular vote by having the states agree to award electoral-college votes based on the country's overall popular-vote outcome.

because they are not constrained by preconceptions; in other words, they may be more likely to result in an “outside-the-box” approach to the problem.³⁹

Lipson's suggestion that Eureqa exhibits curiosity and Koza's argument about computers acting without preconceptions suggest (X) that creativity may be something that is within reach of a computer's capabilities. The ultimate test of such an idea might be to see if a computer could create something that humans would accept as a work of art. Genuine artistic creativity—perhaps more so than any other intellectual endeavor—is something we associate exclusively with the human mind. As *Time's* Lev Grossman says, “Creating a work of art is one of those activities we reserve for humans and humans only. It's an act of self-expression; you're not supposed to be able to do it if you don't have a self.”⁴⁰ Embracing the possibility that a computer could be a legitimate artist would require a fundamental reevaluation of our assumptions about the nature of machines. key =

In the 2004 film *I, Robot*, the protagonist, played by Will Smith, asks a robot, “Can a robot write a symphony? Can a robot turn a canvas into a beautiful masterpiece?” The robot's reply “Can you?” is meant to suggest that, well, the vast majority of people can't do those things either. In the real world of 2015, however, Smith's question would elicit a more forceful answer: “Yes.”

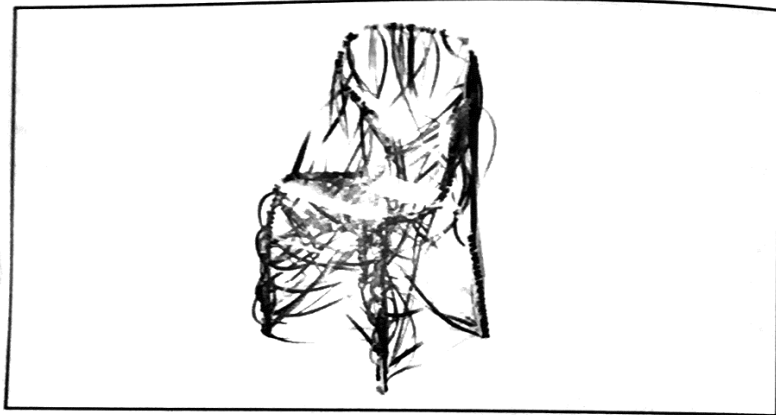
X In July 2012, the London Symphony Orchestra performed a composition entitled *Transits—Into an Abyss*. One reviewer called it “artistic and delightful.”⁴¹ The event marked the first time that an elite orchestra had played music composed entirely by a machine. The composition was created by Iamus, a cluster of computers running a musically inclined artificial intelligence algorithm. Iamus, which is named after a character from Greek mythology who was said to understand the language of birds, was designed by researchers at the University of Malaga in Spain. The system begins with minimal information, such as the type of instruments that will play the music, and then, with no further human intervention, creates a

highly complex composition—which can often evoke an emotional response in audiences—within minutes. Iamus has already produced millions of unique compositions in the modernist classical style, and is likely to be adapted to other musical genres in the future. Like Eureka, Iamus has resulted in a start-up company to commercialize the technology. Melomics Media, Inc., has been set up to sell the music from an iTunes-like online store. The difference is that compositions created by Iamus are offered on a royalty-free basis, allowing purchasers to use the music in any way they wish.

Music is not the only art form being created by computers. Simon Colton, a professor of creative computing at the University of London, has built an artificial intelligence program called “The Painting Fool” that he hopes will someday be taken seriously as a painter (see Figure 4.1). “The goal of the project is not to produce software that can make photos look like they’ve been painted; Photoshop has done that for years,” Colton says. “The goal is to see whether software can be accepted as creative in its own right.”⁴²

Colton has built a set of capabilities he calls “appreciative and imaginative behaviors” into the system. The Painting Fool software

Figure 4.1. An Original Work of Art Created by Software



can identify emotions in photographs of people and then paint an abstract portrait that attempts to convey their emotional state. It can also generate imaginary objects using techniques based on genetic programming. Colton's software even has the ability to be self-critical. It does this by incorporating another software application called “Darci” that was built by researchers at Brigham Young University in Utah. The Darci developers started with a database of paintings that had been labeled by humans with adjectives like “dark,” “sad,” or “inspiring.” They then trained a neural network to make the associations and turned it loose to label new paintings. The Painting Fool is able to use feedback from Darci to decide whether or not it is achieving its objectives as it paints.⁴³

My point here is not to suggest that large numbers of artists or musical composers will soon be out of a job. Rather, it is that the techniques used to build creative software—many of which, as we have seen, rely on genetic programming—can be repurposed in countless new ways. If computers can create musical compositions or design electronic components, then it seems likely that they will soon be able to formulate a new legal strategy or perhaps come up with a new way to approach a management problem. For the time being, the white-collar jobs at highest risk will continue to be those that are the most routine or formulaic—but the frontier is advancing quickly.

Nowhere is the rapid pace of that advance more evident than on the stock exchange. Where once financial trading was highly dependent on direct communication between people, either in bustling trading pits or via telephone, it has now come to be largely dominated by machines communicating over fiber-optic links. By some estimates, automated trading algorithms are now responsible for at least a third of transactions in the UK. However, in the US the proportion is much higher, with some estimates putting it at 70 percent. These sophisticated robotic traders—many of which are powered by techniques on the frontier of artificial intelligence research—go far beyond simply executing routine trades. They attempt to profit

by detecting and then snapping up shares in front of huge transactions initiated by mutual funds and pension managers. They seek to deceive other algorithms by inundating the system with decoy bids that are then withdrawn within tiny fractions of a second. In the US, both Bloomberg and Dow News Service offer special machine-readable products designed to feed the algorithms' voracious appetites for financial news that they can—perhaps within milliseconds—turn into profitable trades. The news services also provide real-time metrics that let the machines see which items are attracting the most attention.⁴⁴ Twitter, Facebook, and the blogosphere are likewise all fodder for these competing algorithms. In a 2013 paper published in the scientific journal *Nature*, a group of physicists studied global financial markets and identified “an emerging ecology of competitive machines featuring ‘crowds’ of predatory algorithms,” and suggested that robotic trading had progressed beyond the control—and even comprehension—of the humans who designed the systems.⁴⁵

In the realm inhabited by these continuously battling algorithms, the action unfolds at a pace that would be incomprehensible to the fastest human trader. Indeed, speed—in some cases measured in millionths or even billionths of a second—is so critical to algorithmic trading success that Wall Street firms have collectively invested billions of dollars to build computing facilities and communications paths designed to produce tiny speed advantages. In 2009, for example, a company called Spread Networks spent as much as \$200 million to lay down a new fiber-optic cable link stretching 825 miles in a straight line from Chicago to New York. The company operated in stealth mode so as not to alert the competition even as it blasted its way through the Allegheny Mountains. When the new fiber-optic path came online, it offered a speed advantage of perhaps three or four thousandths of a second compared with existing communications routes. That was enough to allow any algorithmic trading systems employing the new route to effectively dominate their competition. Wall Street firms, faced with algorithmic decimation, lined up to lease

bandwidth—reportedly at a cost as much as ten times that of the original, slower cable. A similar cable stretching across the Atlantic between London and New York is currently in progress, and is expected to shave about five thousandths of a second off current execution times.⁴⁶

The impact of all this automation is clear: even as the stock market continued on its upward trajectory in 2012 and 2013, large Wall Street banks announced massive layoffs, often resulting in the elimination of tens of thousands of jobs. At the turn of the twenty-first century, Wall Street firms employed nearly 150,000 financial workers in New York City; by 2013, the number was barely more than 100,000—even as both the volume of transactions and the industry's profits soared.⁴⁷ Against the backdrop of that overall collapse in employment, Wall Street did create at least one very high-profile job: in late 2012, David Ferrucci, the computer scientist who led the effort to build Watson, left IBM for a new gig at a Wall Street hedge fund, where he'll be applying the latest advances in artificial intelligence to modeling the economy—and, presumably, trying to gain a competitive advantage for his firm's trading algorithms.⁴⁸

Over in the UK, a similar trend has been seen. Despite rising profits in recent years, banks are making significant lay-offs with estimates suggesting that by the end of 2015, Lloyds, HSBC, RBS and Barclays will have cut 189,000 jobs between them over a five year period. These job losses are, in part, a result of regulations brought in after the financial crisis—some in the banking world have claimed that these reforms have restricted what they can do in certain sectors and so these cuts have been imposed upon them. However, they are also almost certainly part of the increase in automation of the industry. Analysts only expect more job losses to follow in the near future.⁴⁹

Offshoring and High-Skill Jobs

While the trend toward increased automation of white-collar jobs is clear, the most dramatic onslaught—especially for truly skilled

professions—still lies in the future. The same cannot necessarily be said for the practice of offshoring, where knowledge jobs are moved electronically to lower-wage countries. Highly educated and skilled professionals such as lawyers, radiologists, and especially computer programmers and information technology workers have already felt a significant impact. In India, for example, there are armies of call center workers and IT professionals, as well as tax preparers versed in the American and British tax systems, and solicitors specifically trained not in their own country's legal system but in US and UK law, standing ready to perform low-cost legal research for firms engaged in domestic litigation. While the offshoring phenomenon may seem completely unrelated to the jobs lost to computers and algorithms, the precise opposite is true: offshoring is very often a precursor of automation, and the jobs it creates in low-wage nations may prove to be short-lived as technology advances. What's more, advances in artificial intelligence may make it even easier to offshore jobs that can't yet be fully automated.

Most economists view the practice of offshoring as just another example of global trade and argue that it invariably makes both parties to the transaction better off. Harvard professor N. Gregory Mankiw, for example, while serving as George W. Bush's chairman of the White House Council of Economic Advisers, said in 2004 that offshoring is "the latest manifestation of the gains from trade that economists have talked about at least since Adam Smith."⁵⁰ Abundant evidence argues to the contrary. Trade in tangible goods creates a great many peripheral jobs in areas like shipping, distribution, and retail. There are also natural forces that tend to mitigate the impact of globalization to some degree; for example, a company that chooses to move a factory to China incurs both shipping costs and a significant delay before completed products reach consumer markets. Electronic offshoring, in contrast, is almost completely frictionless and subject to none of these penalties. Jobs are moved to low-wage locations instantly and at minimal cost. If peripheral jobs are created,

it is much more likely to be in the country where the workers reside.

I would argue that "free trade" is the wrong lens through which to view offshoring. Instead, it is much more akin to virtual immigration. Suppose, for example, that a huge customer service call center were to be built south of San Diego, just across the border from Mexico. Thousands of low-wage workers are issued "day worker" passes and are bused across the border to staff the call center every morning. At the end of the workday, the buses travel in the opposite direction. What is the difference between this situation (which would certainly be viewed as an immigration issue) and moving the jobs electronically to India or the Philippines? In both cases, workers are, in effect, "entering" the United States to offer services that are clearly directed at the domestic US economy. The biggest difference is that the Mexican day worker plan would probably be significantly better for the California economy. There might be jobs for bus drivers, and there would certainly be jobs for people to maintain the huge facility located on the US side of the border. Some of the workers might purchase lunch or even a cup of coffee while at work, thus injecting consumer demand into the local economy. The company that owned the California facility would pay property tax. When the jobs are offshored, and the workers enter the United States virtually, the domestic economy receives none of these benefits. It's somewhat ironic that, in the West, those on the political right speak of securing borders against immigrants who will likely take jobs that few local people want, while at the same time they express little concern that the virtual border is left completely open to higher-skilled workers who take jobs that local people definitely *do* want.

The argument put forth by economists like Mankiw, of course, measures in the aggregate and glosses over the highly disproportionate impact that offshoring has on the groups of people who either suffer or benefit from the practice. On the one hand, a relatively small but still significant group of people—potentially measured in the millions—may be subjected to a substantial downgrade in their

income, quality of life, and future prospects. Many of these people may have made substantial investments in education and training. Some workers may lose their income entirely. Mankiw would likely argue that the aggregate benefit to consumers makes up for these losses. Unfortunately, although consumers may benefit from lower prices as a result of the offshoring, these savings may be spread across a population of tens or even hundreds of millions of people, perhaps resulting in a cost reduction that amounts to mere pennies and has a negligible effect on any one individual's well-being. And, needless to say, not all the gains will flow to consumers; a significant fraction will end up in the pockets of a few already-wealthy executives, investors, and business owners. This asymmetric impact is, perhaps not surprisingly, intuitively grasped by most average workers but seemingly lost on many economists.

One of the few economists to recognize offshoring's disruptive potential in the US is the former vice chairman of the Federal Reserve's Board of Governors, Alan Blinder, who wrote a 2007 op-ed in the *Washington Post* entitled "Free Trade's Great, but Offshoring Rattles Me."⁵¹ Blinder has conducted a number of surveys aimed at assessing the future impact of offshoring and has estimated that 30–40 million US jobs—positions employing roughly a quarter of the workforce—are potentially offshorable. As he says, "We have so far barely seen the tip of the offshoring iceberg, the eventual dimensions of which may be staggering."⁵²

In Europe, the UK is more affected by offshoring than its continental counterparts. Research by the Forrester Group has previously predicted that of the million European jobs expected to be offshored by the end of 2015, over three-quarters will be from the UK.⁵³ Meanwhile, the Organisation for Economic Co-operation and Development (OECD) estimates that 20 percent of existing jobs in the EU could potentially be offshored.⁵⁴

Virtually any occupation that primarily involves manipulating information and is not in some way anchored locally—for example,

with a requirement for face-to-face interaction with customers—is potentially at risk from offshoring in the relatively near future and then from full automation somewhat further out. Full automation is simply the logical next step. As technology advances, we can expect that more and more of the routine tasks now performed by offshore workers will eventually be handled entirely by machines. This has already occurred with respect to some call center workers who have been replaced by voice automation technology. As truly powerful natural language systems like IBM's Watson move into the customer service arena, huge numbers of offshore call center jobs are poised to be vaporized.

As this process unfolds, it seems likely that those companies—and nations—that have invested heavily in offshoring as a route to profitability and prosperity will have little choice but to move up the value chain. As more routine jobs are automated, higher-skill, professional jobs will be increasingly in the sights of the offshorers. One factor that is, I think, underappreciated is the extent to which advances in artificial intelligence as well as the big data revolution may act as a kind of catalyst, making a much broader range of high-skill jobs potentially offshorable. As we've seen, one of the tenets of the big data approach to management is that insights gleaned from algorithmic analysis can increasingly substitute for human judgment and experience. Even before advancing artificial intelligence applications reach the stage where full automation is possible, they will become powerful tools that encapsulate ever more of the analytic intelligence and institutional knowledge that give a business its competitive advantage. A smart young offshore worker wielding such tools might soon be competitive with far more experienced professionals in developed countries who command very high salaries.

When offshoring is viewed in combination with automation, the potential aggregate impact on employment is staggering. In 2013, researchers at the University of Oxford's Martin School conducted a detailed study of over seven hundred job types and came to the conclusion

that nearly 50 percent of US jobs will ultimately be susceptible to full machine automation.⁵⁵ In February 2015 a parliamentary report to the UK's House of Lords estimated 35 percent of UK jobs will fall victim to automation within the next twenty years.⁵⁶ Now recall that OECD study that said 20 percent of jobs in the EU were potentially offshorable. Let's hope there's significant overlap between those two estimates! Indeed, in all likelihood there is plenty of overlap when the estimates are viewed in terms of job titles or descriptions. The story is different along the time dimension, however. Offshoring will often arrive first; to a significant degree, it will accelerate the impact of automation even as it drags higher-skill jobs into the threat zone.

As powerful AI-based tools make it easier for offshore workers to compete with their higher-paid counterparts in developed countries, advancing technology is also likely to upend many of our most basic assumptions about which types of jobs are potentially offshorable. Nearly everyone believes, for example, that occupations that require physical manipulation of the environment will always be safe. Yet military pilots in the US and UK routinely operate drone aircraft in Afghanistan. Indeed, the British Ministry of Defence intends unmanned aerial vehicles to make up a third of the Royal Air Force's fleet by as soon as 2030.⁵⁷ By the same token, it is easy to envision remote-controlled machinery being operated by offshore workers who provide the visual perception and dexterity that, for the time being, continues to elude autonomous robots. A need for face-to-face interaction is another factor that is assumed to anchor a job locally. However, telepresence robots are pushing the frontier in this area and have already been used to offshore English language instruction from Korean schools to the Philippines. In the not too distant future, advanced virtual reality environments will likewise make it even easier for workers to move seamlessly across national borders and engage directly with customers or clients.

As offshoring accelerates, university graduates in advanced countries like the US and UK may face daunting competition based

not just on wages but also on cognitive capability. The combined population of India and China amounts to roughly 2.6 billion people—that's eight times the population of the United States or over forty times the population of the United Kingdom. The top 5 percent in terms of cognitive ability amounts to about 130 million people—double the entire UK population. In other words, the inescapable reality of the bell-curve distribution stipulates that there are far more very smart people in India and China than in the United States or the United Kingdom. That will not necessarily be a cause for concern, of course, as long as the domestic economies in those countries are capable of creating opportunities for all those smart workers. The evidence so far, however, suggests otherwise. India has built a major, nationally strategic industry specifically geared toward the electronic capture of American and European jobs. And China, even as the growth rate of its economy continues to be the envy of the world, struggles year after year to create sufficient white-collar jobs for its soaring population of new university graduates. In mid-2013, Chinese authorities acknowledged that only about half of the country's current crop of university graduates had been able to find jobs, while more than 20 percent of the previous year's graduates remained unemployed—and those figures are inflated when temporary and freelance work, as well as enrollment in postgraduate degrees and government-mandated "make work" positions, are regarded as full employment.⁵⁸

Thus far, a lack of proficiency in English and other European languages has largely prevented skilled workers in China from competing aggressively in the offshoring industry. Once again, however, technology seems likely to eventually demolish this barrier. Technologies like deep learning neural networks are poised to transport instantaneous machine voice translation from the realm of science fiction into the real world—and this could happen within the next few years. In June 2013, Hugo Barra, Google's top Android executive, indicated that he expects a workable "universal translator"

that could be used either in person or over the phone to be available within several years. Barra also noted that Google already has “near perfect” real-time voice translation between English and Portuguese.⁵⁹ As more and more routine white-collar jobs fall to automation in countries throughout the world, it seems inevitable that competition will intensify to land one of the dwindling number of positions that remain beyond the reach of the machines. The very smartest people will have a significant advantage, and they won’t hesitate to look beyond national borders. In the absence of barriers to virtual immigration, the employment prospects for nonelite university-educated workers in developed economies could turn out to be pretty grim.

Education and Collaboration with the Machines

As technology advances and more jobs become susceptible to automation, the conventional solution has always been to offer workers more education and training so that they can step into to new, higher-skill roles. As we saw in Chapter 1, millions of lower-skill jobs in areas like fast food and retail are at risk as robots and self-service technologies begin to encroach aggressively in these areas. We can be sure that more education and training will be the primary proffered solution for these workers. Yet, the message of this chapter has been that the ongoing race between technology and education may well be approaching the endgame: the machines are coming for the higher-skill jobs as well.

Among economists who are tuned in to this trend, a new flavor of conventional wisdom is arising: the jobs of the future will involve collaborating with the machines. Erik Brynjolfsson and Andrew McAfee of the Massachusetts Institute of Technology have been especially strong proponents of this idea, advising workers that they should learn to “race with the machines”—rather than against them.

While that may well be sage advice, it is nothing especially new. Learning to work with the prevailing technology has always been a good career strategy. We used to call it “learning computer skills.” Nevertheless, we should be very skeptical that this latest iteration will prove to be an adequate solution as information technology continues on its relentless exponential path.

The poster child for the machine-human symbiosis idea has come to be the relatively obscure game of freestyle chess. More than a decade after IBM’s Deep Blue computer defeated world chess champion Garry Kasparov, it is generally accepted that, in one-on-one contests between computers and humans, the machines now dominate absolutely. Freestyle chess, however, is a team sport. Groups of people, who are not necessarily world-class chess players individually, compete against each other and are allowed to freely consult with computer chess programs as they evaluate each move. As things stand in 2014, human teams with access to multiple chess algorithms are able to outmatch any single chess-playing computer.

There are a number of obvious problems with the idea that human-machine collaboration, rather than full automation, will come to dominate the workplaces of the future. The first is that the continued dominance of human-machine teams in freestyle chess is by no means assured. To me, the process that these teams use—evaluating and comparing the results from different chess algorithms before deciding on the best move—seems uncomfortably close to what IBM Watson does when it fires off hundreds of information-seeking algorithms and then succeeds in ranking the results. I don’t think it is much of a stretch to suggest that a “meta” chess-playing computer with access to multiple algorithms may ultimately defeat the human teams—especially if speed is an important factor.

Secondly, even if the human-machine team approach does offer an incremental advantage going forward, there is an important question as to whether employers will be willing to make the

Q4S on
the usual solution
of offering
more training

Q4S on
the idea of
learning to
collaborate with
the machines

investment necessary to leverage that advantage. In spite of the mottos and slogans that corporations direct at their employees, the reality is that most businesses are not prepared to pay a significant premium for "world-class" performance when it comes to the bulk of the more routine work required in their operations. If you have any doubts about this, I'd suggest trying to call your broadband operator's technical support. Businesses *will* make the investment in areas that are critical to their core competency—in other words, the activities that give the business a competitive advantage. Again, this scenario is nothing new. And, more importantly, it doesn't really involve any new people. The individuals that businesses are likely to hire and then couple with the best available technology are the same people who are largely immune to unemployment today. It is a small population of elite workers. Economist Tyler Cowen's 2013 book *Average Is Over* quotes one freestyle chess insider who says that the very best players are "genetic freaks."⁶⁰ That hardly makes the machine collaboration idea sound like a systemic solution for masses of people pushed out of routine jobs. And, as we have just seen, there is also the problem of offshoring. A great many of those 2.6 billion people in India and China are going to be pretty eager to grab one of those elite jobs.

There are also good reasons to expect that many machine collaboration jobs will be relatively short-lived. Recall the example of WorkFusion and how the company's machine learning algorithms incrementally automate the work performed by freelancers. The bottom line is that if you find yourself working with, or under the direction of, a smart software system, it's probably a pretty good bet that—whether you're aware of it or not—you are also training the software to ultimately replace you.

Yet another observation is that, in many cases, those workers who seek a machine collaboration job may well be in for a "be careful what you wish for" epiphany. As one example, consider the current trends in legal discovery. When corporations engage in litigation, it

becomes necessary to sift through enormous numbers of internal documents and decide which ones are potentially relevant to the case at hand. The rules require these to be provided to the opposing side, and there can be substantial legal penalties for failing to produce anything that might be pertinent. One of the paradoxes of the paperless office is that the sheer number of such documents, especially in the form of emails, has grown dramatically since the days of typewriters and paper. To deal with this overwhelming volume, law firms are employing new techniques.

The first approach involves full automation. So-called e-Discovery software is based on powerful algorithms that can analyze millions of electronic documents and automatically tease out the relevant ones. These algorithms go far beyond simple key-word searches and often incorporate machine learning techniques that can isolate relevant concepts even when specific phrases are not present.⁶¹ One direct result has been the evaporation of large numbers of jobs for lawyers and paralegals who once would have sorted laboriously through cardboard boxes full of paper documents.

There is also a second approach in common use: law firms may outsource this discovery work to specialists who hire legions of recent law school graduates. These graduates are typically victims of the bursting law school enrollment bubble. Unable to find employment as full-fledged lawyers—and often burdened with enormous student loans—they instead work as document reviewers. Each lawyer sits in front of a monitor where a continuous stream of documents is displayed. Along with the document, there are two buttons: "Relevant" and "Not Relevant." The law school graduates scan the document on the screen and click the proper button. A new document then appears.⁶² They may be expected to categorize up to eighty documents per hour.⁶³ For these young lawyers, there are no courtrooms, no opportunity to learn or to grow in their profession, and no opportunity for advancement.

handing / training the machine - that does the basic tasks

craft is very concepts

no learning of the craft

Instead, there are—hour after hour—the “Relevant” and “Not Relevant” buttons.*

One obvious question regarding these two competing approaches is whether the collaboration model is sustainable. Even at the relatively low wages (for lawyers) commanded by these workers, the automated approach seems far more cost-effective. As to the low quality of these jobs, you might assume that I've simply cherry-picked a rather dystopian example. After all, won't most jobs that involve collaboration with machines put people in control—so that workers supervise the machines and engage in rewarding work, rather than simply acting as gears and cogs in a mechanized process?

The problem with this rather wishful assumption is that the data does not support it. In his 2007 book *Super Crunchers*, Yale University professor Ian Ayres cites study after study showing that algorithmic approaches routinely outperform human experts. When people, rather than computers, are given overall control of the process, the results almost invariably suffer. Even when human experts are given access to the algorithmic results in advance, they still produce outcomes that are inferior to the machines acting autonomously. To the extent that people add value to the process, it is better to have them provide specific inputs to the system instead of giving them overall control. As Ayres says, “Evidence is mounting in favor of a different and much more demeaning, dehumanizing mechanism for combining expert and [algorithmic] expertise.”⁶⁴

* If you find this type of work appealing but lack the requisite legal training, be sure to check out Amazon's “Mechanical Turk” service, which offers many similar opportunities. “BinCam,” for example, places cameras in your garbage bin, tracks everything you throw away, and then automatically posts the record to social media. The idea is, apparently, to shame yourself into not wasting food and not forgetting to recycle. As we've seen, visual recognition (of types of garbage, in this case) remains a daunting challenge for computers, so people are employed to perform this task. The very fact that this service is economically viable should give you some idea of the wage level for this kind of work.

My point here is that while human-machine collaboration jobs will certainly exist, they seem likely to be relatively few in number and often short-lived. In a great many cases, they may also be unrewarding or even dehumanizing. Given this, it seems difficult to justify suggesting that we ought to make a major effort to specifically educate people in ways that will help them land one of these jobs—even if it were possible to pin down exactly what such training might entail. For the most part, this argument strikes me as a way to patch the tyres on a very conventional idea (give workers still more vocational training) and keep it rolling for a bit longer. We are ultimately headed for a disruption that will demand a far more dramatic policy response.

SOME OF THE FIRST JOBS TO fall to white-collar automation are sure to be the entry-level positions taken by new university graduates. As we saw in Chapter 2, there is already evidence to suggest that this process is well under way. Between 2007 and 2012, average starting salaries for UK graduates with bachelor's degrees fell in real terms by 11 percent, from £24,293 to £21,701 (at 2007 prices).⁶⁵ Over the same period, the total outstanding student debt in the UK more than doubled, from nearly £22 billion in 2007/08 to over £46 billion in 2012/13—the government estimates that this figure will be over £100 billion come 2018.⁶⁶

* In *Average Is Over*, Tyler Cowen estimates that perhaps 10–15 percent of the American workforce will be well equipped for machine collaboration jobs. I think that in the long run, even that estimate might be optimistic, especially when you consider the impact of offshoring. How many machine collaboration jobs will also be anchored locally? (One exception to my skepticism about machine collaboration jobs may be in healthcare. As discussed in Chapter 6, I think it might eventually be possible to create a new type of medical professional with far less training than a doctor who would work together with an AI-based diagnostic and treatment system. Healthcare is a special case, however, because doctors require an extraordinary amount of training and there is likely to be a significant shortage of physicians in the future.)

Underemployment among recent graduates is rampant, and nearly every university student seemingly knows someone whose degree has led to a career working at a coffee shop. In March 2013, Canadian economists Paul Beaudry, David A. Green, and Benjamin M. Sand published an academic paper entitled "The Great Reversal in the Demand for Skill and Cognitive Tasks."⁶⁷ That title essentially says it all: the economists found that around the year 2000, overall demand for skilled labor in the United States peaked and then went into precipitous decline. In the UK, a 2014 report from the CIPD, a professional body for HR development, accused the government of adopting a policy direction which has led to the under-use of highly skilled workers. Their research indicates that 30 percent of British employees are overqualified for the job they are presently in.⁶⁸ Increasingly, graduates in both the US and the UK have been forced into relatively unskilled jobs, often displacing nongraduates in the process.

Even those graduates with degrees in scientific and technical fields have been significantly impacted. As we've seen, the information technology job market, in particular, has been transformed by the increased automation associated with the trend toward cloud computing as well as by offshoring. The widely held belief that a degree in engineering or computer science guarantees a job is largely a myth. An April 2013 analysis by the Economic Policy Institute found that at universities in the United States, the number of new graduates with engineering and computer science degrees exceeds the number of graduates who actually find jobs in these fields by 50 percent. The study concludes that "the supply of graduates is substantially larger than the demand for them in industry."⁶⁹ In the UK it's a similar story for computer science graduates who, for a number of years, have consistently been the most likely to be unemployed six months after graduation. However the picture for the UK's young engineers isn't as bleak as it is for their US counterparts—70 percent of 2013's civil engineering graduates have a job related to their degree, as do nearly

two-thirds of mechanical engineering graduates.⁷⁰ Nonetheless, it is becoming increasingly clear that a great many people will do all the right things in terms of pursuing an advanced education, and yet they will be unable to find a foothold in the economy of the future.

While some of the economists who focus their efforts on sifting through reams of historical data are finally beginning to discern the impact that advancing technology is having on higher-skill jobs, they are typically quite cautious about attempting to project that trend into the future. Researchers working in the field of artificial intelligence are often far less reticent. Noriko Arai, a mathematician with Japan's National Institute of Informatics, is leading a project to develop a system capable of passing the Tokyo University entrance examination. Arai believes that if a computer can demonstrate the combination of natural language aptitude and analytic skill necessary to gain entrance to Japan's highest-ranked university, then it will very likely also be able to eventually perform many of the jobs taken by university graduates. She foresees the possibility of massive job displacement within the next ten to twenty years. One of the primary motivations for her project is to try to quantify the potential impact of artificial intelligence on the job market. Arai worries that 10 to 20 percent of skilled workers replaced by automation would be a "catastrophe" and says she "can't begin to think what 50 percent would mean." She then adds that it would be "way beyond a catastrophe and such numbers can't be ruled out if AI performs well in the future."⁷¹

The higher-education industry itself has historically been one of the primary employment sectors for highly skilled workers. Especially for those who aspire to a doctoral degree, a typical career path has been to arrive on campus as an undergraduate and then never really leave. In the next chapter we'll look at how that industry, and a great many careers, may also be on the verge of a massive technological disruption.