

A Brief History of AI

Artificial Intelligence: "The conjecture that every aspect of learning or any other feature of intelligence can in principle be so precisely described that a machine can be made to simulate it."

– First definition coming from the Dartmouth Summer Research Project proposal in 1956

Defining Artificial Intelligence (AI) is not easy as AI is a young discipline. AI is undoubtedly a bold, exciting new world, where the lines between humans and machines blur, leading us to question the very nature of intelligence itself. In addition, AI is recognized as a collection of scientific disciplines including mathematical logic, statistics, probabilities, computational neurobiology, and computer science that aims to perform tasks commonly associated with the human cognitive abilities such as the ability to reason, discover meaning, generalize, or learn from past experiences.

Interestingly, AI founders weren't just computer scientists. They were philosophers, mathematicians, neuroscientists, logicians, and economists. Shaping the course of AI required them to integrate a wide range of problem-solving techniques. These tools spanned from formal logic and statistical models to artificial neural networks and even operations research. This multidisciplinary approach became the key to solving intricate problems that AI posed.

The Precursors of the Mechanical or “Formal” Reasoning

One of these precursors was the French philosopher and theologist René Descartes (1596–1650), who wrote in 1637 his *Discourse on the Method*,¹ one of the most influential works in the history of modern philosophy, and important to the development of natural science. He discussed in Part V the conditions required for an animal or a machine to demonstrate an intelligent being. This was one of the earliest examples of philosophical discussion about artificial intelligence. He envisioned later in his *Meditations on First Philosophy*² (1639) the possibility of having machines being composed like humans but with no mind.

In 1666, the German polymath Gottfried Wilhelm Leibniz (1646–1716) published a work entitled *On the Combinatorial Art*³ in which he expressed his strong belief that all human reasoning can be represented mathematically and reduced to a calculation. To support this vision, he conceptualized and

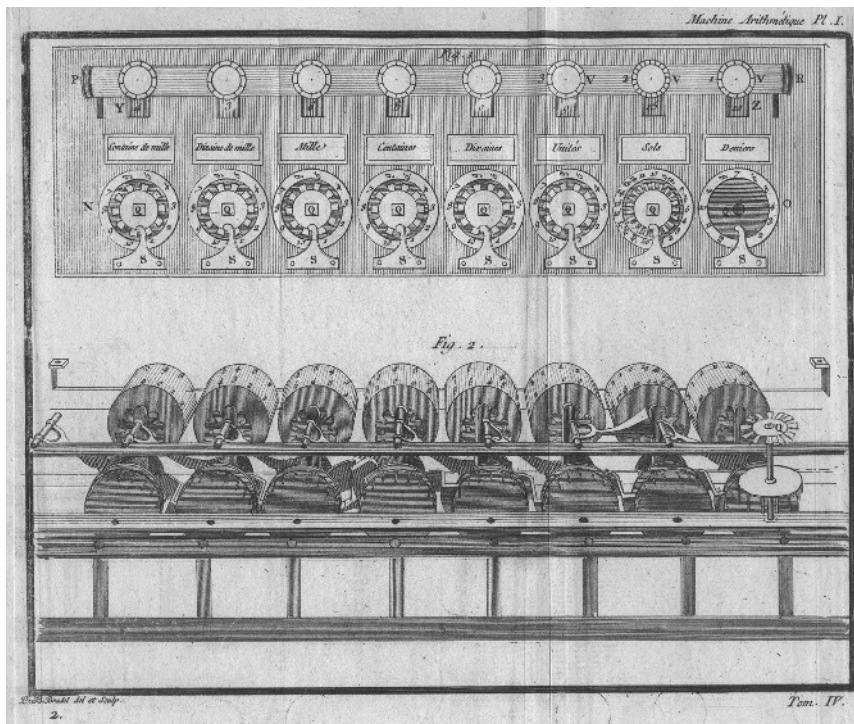


Figure 1-1: Drawing of the top view of the Pascaline and overview of its mechanism.

1779, *Oeuvres de Blaise Pascal*, Chez Detune, La Haye, Public Domain.

¹https://en.wikisource.org/wiki/Discourse_on_the_Method

²https://en.wikisource.org/wiki/Meditations_on_First_Philosophy

³<https://gallica.bnf.fr/ark:/12148/bpt6k625780>

described in his writings a *Calculus ratiocinator*: a theoretical universal logical calculation method to make these calculations feasible and a *Characteristica Universalis*: a universal and formal language to express mathematical, scientific, and metaphysical concepts.

At the same time, Blaise Pascal (1623–1662) built in 1641 one of the first working calculating machines called the “Pascaline,” shown in Figure 1.1, which could perform additions and subtractions. Inspired by this work, Leibniz built his “Stepped reckoner” (1694), as shown in Figure 1.2, a more sophisticated mechanical calculator that could not only add and subtract but also multiply and take the square root of a number.

After the initial developments in mechanical calculation, further advancements were made in the early nineteenth century. In 1822, English mathematician Charles Babbage (1791–1871) designed the Difference Engine, an automatic mechanical calculator intended to tabulate polynomial functions. The Difference Engine was conceived as a room-sized machine, but it was never constructed in Babbage’s lifetime.

Building on the foundations laid by Leibniz, George Boole published in 1854 *The Laws of Thought*⁴ presenting the concept that logical reasoning could

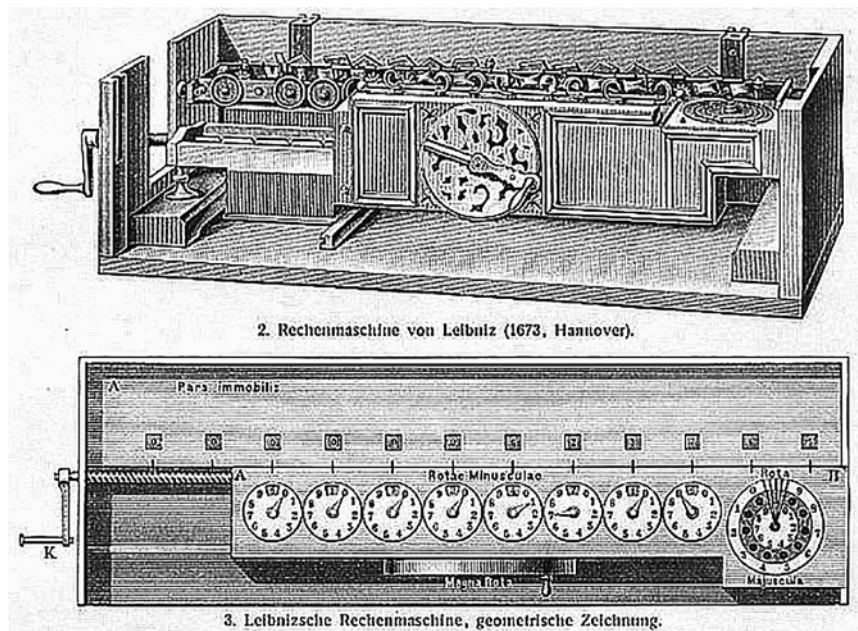


Figure 1-2: Drawing of the Stepped reckoner.

Hermann Julius Meyer /Wikimedia Commons/Public domain.

⁴<https://www.gutenberg.org/files/15114/15114-pdf.pdf>

be expressed mathematically through a system of equations. Now known as Boolean algebra, Boole's breakthrough established the basis for computer programming languages. Additionally, in 1879 German mathematician Gottlob Frege (1848–1925) put forth his *Begriffsschrift*,⁵ which established a formal system for logic and mathematics. The pioneering work of Boole and Frege on formal logic laid essential groundwork that enabled subsequent developments in computation and computer science.

The Digital Computer Era

In 1936, mathematician Alan Turing (Figure 1.3) published his landmark paper “On Computable Numbers,”⁶ conceptually outlining a hypothetical *universal machine* capable of computing any solvable mathematical problem encoded symbolically. This theoretical Turing machine established a framework for designing computers using mathematical logic and introduced the foundational notion of an *algorithm* for programming sequences.

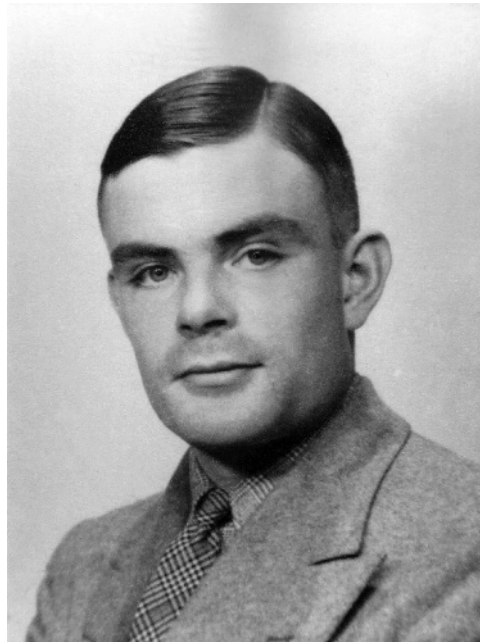


Figure 1-3: Alan Turing.

GL Archive / Alamy Stock Photo

⁵<https://gallica.bnf.fr/ark:/12148/bpt6k65658c>

⁶https://www.cs.virginia.edu/~robins/Turing_Paper_1936.pdf

Around the same time, Claude Shannon's 1937 master's thesis, *A Symbolic Analysis of Relay and Switching Circuits*,⁷ demonstrated Boolean algebra's applicability for optimizing electric relay arrangements: the core components of electromechanical telephone routing systems. Shannon thus paved the way for applying logical algebra to circuit design.

Concurrently in 1941, German engineer Konrad Zuse developed the world's first programmable, fully functional computer, the Z3. Built from 2,400 electro-mechanical relays, the Z3 embodied the theoretical computer models proposed by Turing and leveraged Boolean logic per Shannon's insights. Zuse's pioneering creation was destroyed during World War II.

In the aftermath of World War II, mathematician John von Neumann made vital contributions to emerging computer science. He consulted on the ENIAC (Figure 1.4), the pioneering programmable electronic digital computer built for the US Army during the war. In 1945, von Neumann authored a hugely influential

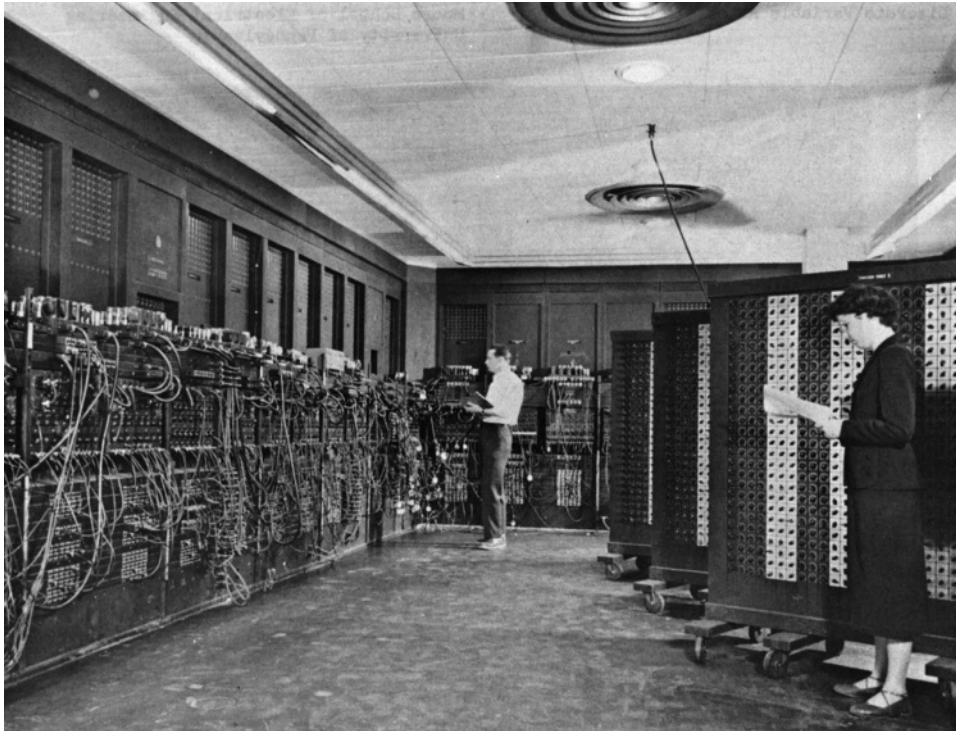


Figure 1-4: ENIAC in building 328 at the Ballistic Research Laboratory (BRL).

Ballistic Research Laboratory, 1947–1955, US Army.

⁷<https://www.cs.virginia.edu/~evans/greatworks/shannon38.pdf>

report on the proposed EDVAC computer, outlining the “stored-program concept”: separating a computer’s task-specific programming from its general-purpose hardware that sequentially executes instructions. This conceptual distinction enabled the adoption of software programs without reconfiguring physical components, thereby allowing a single machine to readily perform different sequences of operations. Von Neumann’s architectural vision profoundly shaped modern computing as the standard *von Neumann architecture*.

Following the first functioning computer constructions, Alan Turing reflected on the capabilities afforded by these theoretically “universal machines.” In a 1948 report, he argued that one such general-purpose device should be sufficient to carry out any computable task, rather than needing infinite specialized machines. Developing this thread further in his landmark 1951 paper “Computing Machinery and Intelligence,”⁸ Turing considered whether machines might mimic human capacities. To examine this, he proposed what later became known as the *Turing test*, an “imitation game” evaluating whether people could distinguish a concealed computer from a human respondent based solely on typed conversation.

In the conventional form of the Turing test, there are three participants: a human interrogator, interacting through written questions and answers with two players: a computer and a person. The interrogator needs to understand which one is the human solely based on the text-based responses to his or her inquiries. Removing other perceptual cues forces the player to rely entirely on the linguistic content and reasoning within the typed replies when attempting to distinguish human from machine. This restriction highlights Turing’s interest in assessing intelligence manifested in communication; if responses are sufficiently comparable between candidates, it suggests the computer can display capacities approaching human-level understanding and dialogue, at least conversationally. Thus, passing this verbal imitation game constituted Turing’s proposed measure for demonstrated machine intelligence.

Cybernetics and the Beginning of the Robotic Era

The term *robot* entered the lexicon in 1920 with Czech writer Karel Capek’s play *R.U.R.* (Rossum’s Universal Robots), which featured artificial workers created to serve humans. The specific concept of *robotics* as a field of study then emerged in science fiction over the following decades. Notably, the 1942 short story “Runaround” by Isaac Asimov introduced his influential Three Laws of Robotics, ethical constraints programmed into the fictional robots to govern their behavior:

⁸<https://academic.oup.com/mind/article-pdf/LIX/236/433/9866119/433.pdf>

1. A robot may not injure a human being, or, through inaction, allow a human being to come to harm.
2. A robot must obey the orders given it by human beings except where such orders would conflict with the First Law.
3. A robot must protect its own existence as long as such protection does not conflict with the First or Second Law.

These simple yet profound guidelines shaped many subsequent fictional depictions and philosophical discussions around machine intelligence safety and control. Along with introducing seminal ideas, Asimov's writings disseminated the term *robotics* into broader usage to describe the engineering discipline focused on constructing automated mechanical beings.

The term *cybernetics* refers to a theory of control mechanisms applied to regulate complex systems. It originates from the ancient Greek word *kybernetikos* meaning "skilled in steering," referring to a ship helmsman's ability to steer and navigate vessels. The mathematician Norbert Wiener (1894–1964) first coined the term cybernetics in his 1948 book of the same name *Cybernetics*,⁹ where he laid the foundations of this new interdisciplinary field. Through his pioneering work, Wiener sought to model the principles behind self-governing behavior, drawing parallels between mechanical, biological, and social systems. Ultimately, he envisioned cybernetics as a way to design autonomous, self-correcting control processes across organizations and machines alike, from servomechanisms to human organizations. The connections Wiener made between self-regulation in mechanical, organic, and social contexts provided an insightful new lens through which complex systems could be analyzed and understood.

The emergence of cybernetics was heavily influenced by concurrent research in neuroscience revealing that the brain operates as an electrical network of neurons that transmit signals. Central to Wiener's theory of cybernetics were feedback loops: mechanisms that enable systems to self-adjust their behavior. This concept of self-governing control through circular causality was analogous to the impulse propagation in neural circuits. At the same time, Claude Shannon's 1948 work *A Mathematical Theory of Communication*¹⁰ laid the theoretical foundations for the digitization of analog signals. Additionally in 1943, Warren S. McCulloch (1898–1969) and Walter Pitts (1929–1969) demonstrated how neural networks could implement logical propositional calculus. The confluence of these key discoveries in cybernetics, information theory, and neurocomputation suggested the possibility of constructing an "electronic

⁹<https://direct.mit.edu/books/oa-monograph/4581/Cybernetics-or-Control-and-Communication-in-the>

¹⁰<https://people.math.harvard.edu/~ctm/home/text/others/shannon/entropy/entropy.pdf>

brain.” This prompted attempts in the 1950s to build the first experimental robots guided entirely by analog circuitry designed to mimic biological neural networks. While primitive, these early successes fed momentum toward developing sophisticated systems that could replicate more complex human decision-making and learning. The interdisciplinary fusion of neuroscience and computational controls remains at the heart of efforts to construct adaptable, self-governing artificial intelligence.

In 1950, British inventor Tony Sale built one of the first human-like robots, nicknamed George, which stood six feet tall and was able to walk and talk. Sale constructed George entirely from analog electronics to carry out conversations and movements activated through pulleys and air pressure. While clunky and primitive by modern standards, robots like George and the analog-based machines developed concurrently represented the pioneering first steps toward life-like intelligent automation. The creation of an anthropomorphic machine, or *android*, capable of human activities like self-directed motion and dialogue opened the door both technologically and psychologically to the widespread cultural adoption of robots. No longer theoretical, these interactive machines inspired sci-fi visions of the future now seemingly within reach. George embodied cybernetics founder Wiener’s vision of building artificial systems that could emulate and perhaps enhance human capabilities. This new generation of thinking machines brought us measurably closer toward that goal.

In the same pioneering year of 1950, Claude Shannon constructed an electro-mechanical mouse named Theseus, capable of solving a maze on its own – one of the earliest examples of a learning machine. Shannon built Theseus entirely from relay circuits and switches to enable it to navigate through a maze by trial and error, “learning” the correct route based on past experience. The mouse would traverse the labyrinth, remembering each turn as either a reward or failure depending on whether it led toward an end goal. Over time, favored paths emerged through this basic feedback learning. Though extraordinarily simple by modern AI standards, Theseus embodied one of the founding principles of cybernetics: that complex goal-oriented behavior could arise from interconnected networks and repetition of simple rules.

In line with early cybernetics principles, French inventor Albert Ducrocq (1921–2001) built an electromechanical fox between 1950 and 1953 that demonstrated autonomous goal-directed behavior through sensory-feedback loops. Nicknamed Job, the robot fox navigated unfamiliar environments largely on its own using an array of analog sensors akin to animal senses. Job was equipped with two photoelectric cells mounted in his head that acted as organs of sight, a microphone that constituted his ear, while his touch came from sensors that reacted to contact with obstacles. Contacts placed in the neck gave him a sense of direction. He also had “capacitive flair,” which

allowed him to recognize an obstacle from a distance. In addition, Job was able to learn through a “memory” and to express himself via two lamps that lit up on the top of his head.

In 1961, Unimate (Figure 1.5), an industrial robot invented by George Devol in the 1950s, became the first to work in a General Motors assembly line in New Jersey. Its responsibilities included transporting die castings from the assembly line and welding the parts on to cars – a task deemed dangerous for humans. Being adopted by the popular culture, it appeared on *The Tonight Show* hosted by Johnny Carson, knocking a golf ball into a cup, pouring a beer, conducting an orchestra, and grasping an accordion.

In the late 1960s, Shakey the Robot (Figure 1.6) was developed as the first mobile robot that had the capacity to apprehend and to interact with its surroundings. It was developed by a group of engineers at Stanford Research Institute (SRI) supervised by Charles Rosen under supervision and funding of the Defense Advanced Research Projects Agency (DARPA).



Figure 1-5: Unimate pouring coffee for a human, 1967.

Wikipedia, CC-BY-SA 4.0, https://commons.wikimedia.org/wiki/File:Unimate_pouring_coffee_for_a_woman_at_Biltmore_Hotel,_1967.jpg



Figure 1-6: Shakey the Robot at the Computer History Museum.

Wikipedia, CC-BY-SA 2.0, <https://commons.wikimedia.org/wiki/File:Shakey.png>, no change.

Birth of AI and Symbolic AI (1955–1985)

As transistor and then integrated circuit technology enabled computers to perform increasingly complex functions, the potential for artificial intelligence vastly expanded. Tasks that previously only the human brain could accomplish, such as playing chess or recognizing images, gradually became possible for AI systems to tackle as well. This interdependent advancement of hardware and software drove rapid progress in the capabilities of artificial intelligence.

In 1950, Claude Shannon published a groundbreaking paper entitled “Programming a Computer for Playing Chess.”¹¹ This article was the first to explore the creation of a computer program capable of playing chess.

¹¹<https://vision.unipv.it/IA1/ProgramminaComputerforPlayingChess.pdf>

However, the field of Artificial Intelligence got officially defined in 1956 with a workshop organized by John McCarthy at the Dartmouth Summer Research Project. The goal of this conference was to investigate ways in which machines could be made to simulate aspects of intelligence. McCarthy coauthored the proposal for the workshop with Marvin Minsky, Nathaniel Rochester, and Claude Shannon. McCarthy was credited for the first use of the term *artificial intelligence* (see Figure 1.7).

Since its inception, AI research has followed two distinct yet competing approaches: the symbolic (or “top-down”) method and the connectionist (or “bottom-up” or subsymbolic) method. The top-down approach aims to mimic intelligence by examining cognition independently of the brain’s biological architecture. It frames thinking in terms of manipulating abstract symbols via rules. On the other hand, the bottom-up methodology focuses on building artificial neural networks that imitate the brain’s interconnected architecture. It gets its “connectionist” designation from the emphasis placed on connections between neuronal units. While differing substantially, both the symbolic and connectionist paradigms strive to unlock the secrets of replicating intelligence artificially. The interplay between these rival approaches has largely defined and driven the evolution of AI research over the decades.

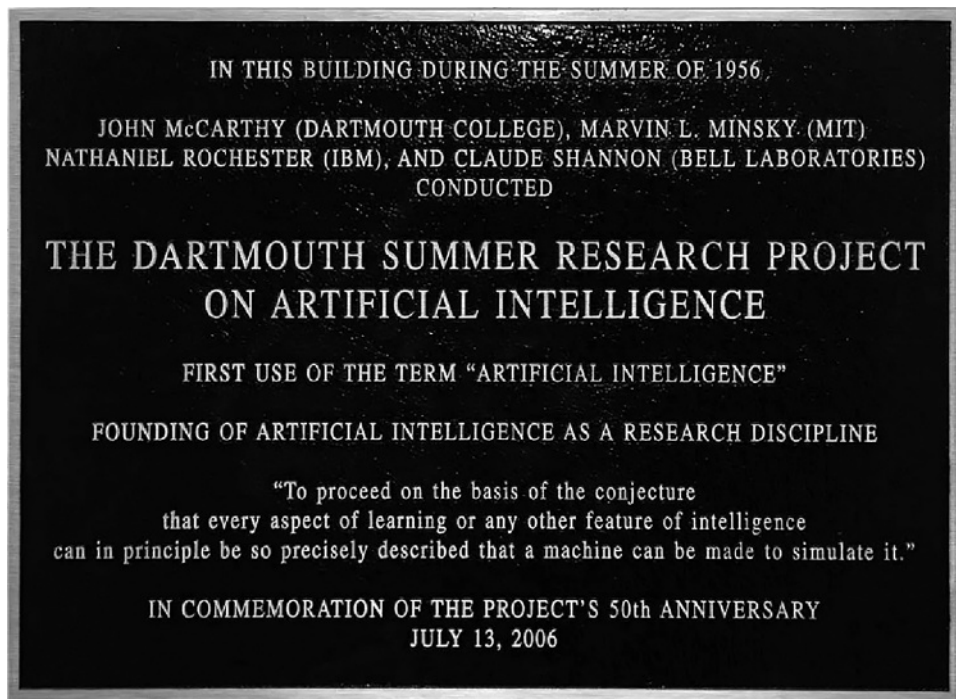


Figure 1-7: Dartmouth workshop commemorative plaque.

In 1952, computer scientist Arthur Samuel created the first self-learning game-playing computer program – one designed to play checkers. His program pioneered the ability to progressively improve its gameplay by accumulating knowledge independently through experience. This marked a significant milestone in developing artificial intelligence that could teach itself skills rather than relying solely on human-programmed rules. Samuel’s checkers application established foundational concepts for AI agents honing strategies through practice over time.

In 1955, researcher Allen Newell and economist Herbert Simon spearheaded the exploration of heuristic search procedures: efficient methods for discovering solutions within massive, multifaceted problem spaces. They implemented this concept in two pioneering AI systems: Logic Theorist, which generated proofs for mathematical theorems, and the more versatile General Problem Solver. Newell and Simon’s adoption of heuristic search techniques enabled their programs to efficiently navigate complex combinatorial possibilities to identify problem-solving pathways. These two programs marked the first implementations of artificial intelligence within computer systems. By leveraging heuristic search, Newell and Simon set the stage for AI applications to practically address intricate real-world problems.

In 1957, Frank Rosenblatt of Cornell University’s Aeronautical Laboratory pioneered research into artificial neural networks, which he termed *perceptrons*. Through rigorous computer simulations and mathematical modeling, Rosenblatt made seminal contributions unraveling the capabilities of neural nets and mechanisms of learning. His connectionist approach stressed the paramount role of modifying connections between neurons to enable knowledge acquisition. Via groundbreaking empirical and theoretical work on network topologies, Rosenblatt cemented himself as a forefather of neural networks within the Artificial Intelligence community. His perceptron findings seeded advancement in adaptive systems that paralleled biological neural functioning.

In 1958, John McCarthy pioneered the creation of Lisp, which has endured as the predominant and preferred programming language utilized in Artificial Intelligence research. McCarthy’s development of Lisp, with its specialized features for manipulating symbolic expressions, provided AI scientists with an essential and adaptable tool for implementing programs geared toward intelligence. To this day, Lisp remains a widely embraced high-level language within the AI community, having withstood the test of time as a staple in the field McCarthy helped establish.

In 1959, Arthur Samuel introduced the seminal notion of “machine learning” while discussing innovations in his chess-playing computer program. He imparted the idea that a computer could be coded to independently enhance its gameplay beyond the skills of even its programmer. Samuel envisioned the possibility for machines to autonomously build expertise by accumulating knowledge through experience, thus “learning” on their own in a way once solely attributed to

human intelligence. His pioneering perspective established machine learning as a fundamental pursuit within Artificial Intelligence research for decades to follow. By conceiving of adaptive programs that exhibit human-like mastery, Samuel set the trajectory toward the advanced AI systems of today.

In 1965, computer scientist Joseph Weizenbaum created ELIZA (Figure 1.8), an early conversational AI program capable of natural language interaction. Though intended to showcase the superficiality of communication between humans and artificial intelligence, ELIZA's uncannily human-like responses led many to curiously anthropomorphize its underlying system. Through his pioneering bot, Weizenbaum inadvertently revealed people's propensity to project human qualities onto AI. In doing so, he surfaced key questions about the fundamental nature of relationships between man and machine that philosophers would grapple with for decades hence. Despite ELIZA's apparent limitations in logic versus human intellect, its historical impact would underscore artificial intelligence's broader potential to profoundly shape human perception and society.

However, by the 1970s, AI faced growing criticism and budget cuts. Scientists had underestimated the depth of the challenges involved, fueling unrealistic public optimism that could not realistically be met. As AI failed to deliver on its lofty promises, private and government funding swiftly evaporated. Concurrently, research into simple, single-layer neural networks stagnated due

```
Welcome to

EEEEEE LL      IIII ZZZZZZZZ  AAAAA
EE      LL      II      ZZ  AA  AA
EEEEEE LL      II      ZZZ  AAAAAA
EE      LL      II      ZZ   AA  AA
EEEEEE LLLLLL  IIII  ZZZZZZZZ  AA  AA

Eliza is a mock Rogerian psychotherapist.
The original program was described by Joseph Weizenbaum in 1966.
This implementation by Norbert Landsteiner 2005.

ELIZA: Please tell me what's been bothering you.
YOU:   The weather is awful.
ELIZA: I'm not sure I understand you fully.
YOU:   I hate the current weather.
ELIZA: You say you hate the current weather ?
YOU:   Yes, that's what I said.
ELIZA: I understand.
YOU:   Great.
ELIZA: Please go on.
YOU:   I would say it's your turn now!
ELIZA: Why are you concerned over my turn now ?
YOU:   █
```

Figure 1-8: A conversation with ELIZA.

Source: Wikimedia Commons/Public domain.

in part to Marvin Minsky's writings spotlighting perceptual limitations. This two-pronged assault on both the symbolic and connectionist approaches left AI reeling from seemingly insurmountable obstacles on all research fronts. For an era once filled with seemingly boundless potential, the 1970s became a bitter period of retreat as the field confronted the harsh realities of its shortcomings. The inflated aspirations for AI would need to be rigorously tempered before progress could resume on more modest, incremental fronts.

Moreover, the disparity between AI's theoretical potential and practical results continued widening in the 1980s. Constraints imposed by restricted computing power, combinatorial expansions, and knowledge representation barriers dashed hopes for near-term progress. Despite a bright horizon, this persistent gap triggered greatly reduced interest and financial backing for Artificial Intelligence. Previously staunch supporters like the British government, DARPA, and NRC grew disillusioned by meager returns on investment. With the field unable to deliver on long-running promises, these agencies aborted funding, casting AI into a paralyzing "winter." This first AI winter was characterized by the mass exodus of scientists due to scarce resources and loss of confidence needed to power human-level intelligence milestones. The period confronted AI with learning to walk again by concentrating on attainable stepping stones before contemplating any rekindled sprint toward human cognition.

Subsymbolic AI Era (1985–2010)

In the 1980s, a type of AI program known as expert systems gained traction by offering focused problem-solving abilities. These programs could field questions and provide solutions within a tightly defined area of expertise using logical rule sets distilled from human specialists. Expert systems found momentum by delivering targeted reasoning skills for particular domains like medicine or engineering without attempting to replicate the full spectrum of human cognition. Their aim was not expansive artificial general intelligence but rather narrowly concentrated artificial niche intelligence. By limiting their capabilities, expert systems attained enough practical utility during the AI winter to regain a foothold for the field. Their success highlighted Artificial Intelligence's path forward: mastering well-scoped applications before progressing to grander human-mimicking aspirations.

Notable pioneering expert systems like DENDRAL (1965) for deducing chemical structures from spectral analyses and MYCIN (1972) for diagnosing bacterial blood infections proved the viability of concentrated knowledge systems. By focusing their algorithmic mastery on specific scientific tasks rather than boundless intelligence, these programs displayed the near-term potential for expert systems. Their demonstrations of feasibility within limited problem spaces

offered a pathway for Artificial Intelligence to rebuild practical value, chipping away at obstacles by domain rather than tackling expansive human cognition all at once. Early expert system success stories thereby reignited pragmatic interest in AI amid the winter by evolving expectations toward achievements in targeted reasoning rather than the elusive quest to mimic mankind's complete intellectual range.

Additional expert system success came in 1980 when Digital Equipment Corporation deployed XCON, an AI system for automatically configuring computer orders per customer specifications. Encoding 2,500 assembly rules, XCON delivered 95–98% precision across 80,000 configurations at DEC's Salem plant, saving \$25 million.

Subsequently in 1984, a bold symbolic AI endeavor named CYC sought to codify extensive common sense knowledge. Developers believed surpassing a critical threshold of human wisdom within CYC would enable self-propelled extraction of further insights from natural language. Although the initiative to systematize reasoning through preprogrammed facts and heuristics proved too ambitious, CYC represented lingering aspirations to manifest more fully general artificial intelligence. But as expectations reset toward more incremental advances, both XCON and CYC highlighted the promise of narrow yet capable AI versus impractical schemes to replicate multifaceted human cognition outright.

The expert systems movement triggered an explosion of commercial interest across industries. Businesses worldwide raced to develop and install domain-specific AI solutions for operational tasks. This surging adoption fueled the rise of affiliated hardware purveyors like Symbolics and Lisp Machines alongside software experts including IntelliCorp and Aion that consulted on specialized system development. As focused artificial intelligence proved economically viable – despite the overall AI winter – robust ecosystems materialized to meet burgeoning corporate demand. The strategic shift from pursuing expansive human mimicry to practical applications revived investment, underscoring the resilient appetite for capable if narrowly intelligent systems. Expert system implementation may not have realized AI's full potential, but it offered a path out of the trough by delivering value scattered across niches.

In 1982, John Hopfield revived enthusiasm for neural networks within the AI community. By devising the seminal *Hopfield net* architecture and mathematical convergence proofs, he substantiated key learning properties applicable to interlocking neuron models. Meanwhile researchers Geoffrey Hinton and David Rumelhart elevated backpropagation techniques aimed at incrementally refining neural processing via error signal feedback. This dual resurrection of connectionist approaches proved timely, as demonstrated by Yann LeCun's backprop-powered neural system at Bell Labs in 1989, which captured complex handwritten digits. Through concerted advancement of neural network algorithms closely paralleling biological brains, this collective renewal reestablished

adaptive networks as a versatile methodology complementary to symbolic methods' logical symbol manipulation.

However, this renaissance soon sputtered as the complexities of neural network programming overwhelmed practical progress. By the early 1990s, this era came to an end despite innovations like IBM's Deep Blue defending its 1997 chess crown against Garry Kasparov – an accomplishment foretold by Herbert Simon that was achieved via systematic brute force rather than sophisticated cognition.

It became apparent that even heavily funded projects like Japan's \$850 million Fifth Generation Computer, UK's £350 million Alvey Initiative, and DARPA's Strategic Computing Program could only produce narrow expert systems ill-suited to general application. As businesses recognized the crippling constraints of commercialization, AI crashed into a second winter by the mid-1990s. The failed expectations of expert systems and inability to capture flexible human thinking in code left investors weary that AI's theories might chronically lack real-world substance without a paradigm shift. While core innovations continued advancing, the field hunkered into a maturity phase marked by modest objectives rather than aspirations of replicating multifaceted cognition.

Deep Learning and LLM (2010–Present)

The possibility to use massive volumes of data changed everything and triggered a new boom supported by the Big Data era, the access to large amounts of data, cheaper and faster computers, and progress on machine learning.

A pivotal 2009 project catalyzed this AI renaissance: Stanford professor Fei-Fei Li's creation of ImageNet, a visual database that curated more than 14 million categorized photos to power machine learning breakthroughs. Li explained, "Our vision was that big data would change machine learning. Data drives learning." This database got used with the ImageNet Large Scale Visual Recognition Challenge (ILSVRC), an annual competition established in 2010 that tasks teams to develop computer vision algorithms to accurately classify and detect objects and scenes within image datasets. By providing a standardized benchmark test for computer vision systems on a massive scale, the ILSVRC accelerated research and development in object recognition technology. Over several contests, error rates by top ILSVRC systems decreased dramatically, even surpassing human-level accuracy by 2015. In 2012, AlexNet got published by Alex Krizhevsky, Ilya Sutskever, and Geoffrey Hinton as one of the first convolutional networks to use a GPU (graphics processing unit) to boost performance. After competing on ILSVRC, AlexNet demonstrated remarkable performance improvements over previous methods and ignited a surge of interest in deep learning.

Still in 2012 a Google research team, led by Stanford's Andrew Ng and Google's Jeff Dean, supervised the creation of a billion-parameter neural network trained

on 10 million random YouTube thumbnail images using an array of 16,000 processors. Without any human labeling or guidance, specific neurons spontaneously developed detectors for concepts like cats and human faces buried within the visual data. This self-organized acquisition demonstrated deep learning's promise for mining big datasets. Their work proved deep neural networks could transcend not just programmed knowledge but also inherent structure through natural observational experience alone.

In 2014 Ian Goodfellow pioneered Generative Adversarial Networks (GANs), a novel AI paradigm using neural networks against each other to boost predictive power. This approach employs two dueling models: one generating hypothetical data mimicking real samples, the other attempting to differentiate artificial from authentic. GANs typically run unsupervised and use a cooperative zero-sum game framework to learn. GANs became a popular ML model for online retail sales because of their ability to understand and recreate visual content with increasingly remarkable accuracy. This was a first step in the direction of Generative AI.

Still in 2014 Google bought UK Artificial Intelligence start-up DeepMind for £400m. Later in 2016 Google DeepMind's AlphaGo (Google's AI specializing in Go games) algorithm won competitions against Go European champion (Fan Hui), Go world champion (Lee Sedol), then itself.

In that same year Facebook activated DeepFace, a pioneering facial recognition system leveraging deep learning techniques, to automatically tag and identify Facebook users.

In December 2015, OpenAI was founded by Sam Altman, Elon Musk, Ilya Sutskever, Greg Brockman, Trevor Blackwell, Vicki Cheung, Andrej Karpathy, Durk Kingma, John Schulman, Pamela Vagata, and Wojciech Zaremba, with Sam Altman and Elon Musk as the co-chairs.

In 2017, Google researchers proposed the transformers architecture, "a novel neural network architecture based on a self-attention mechanism," an improvement of the Seq2Seq technology, which became later widely used in Large Language Models (LLMs).

Starting in 2018 LLMs gained momentum with the releases of some notable LLMs like OpenAI's GPTs (GPT-1 in 2018, GPT-2 in 2019, GPT-3 in 2020, GPT-3.5 and ChatGPT in 2022), Google PaLM and Gemini, Meta LLaMA family, and Anthropic Claude models.

Key Takeaways

From Leibniz's seventeenth-century calculators to Turing's twentieth-century theoretical models, their contributions formulated the essential philosophical, mathematical, and mechanical frameworks upon which AI is constructed. Each incremental innovation, whether Gottfried Wilhelm Leibniz's Step reckoner,

Charles Babbage's mechanical computation, George Boole's mathematical logic, Alan Turing's formalized abstractions of intelligence, or Claude Shannon's game-playing algorithms, allowed AI to become this intelligence of machines or software, backed by mathematical logic, statistics, probabilities, computational neurobiology, and computer science with the aim of performing tasks commonly associated with human cognitive abilities such as the ability to reason, discover meaning, generalize, or learn from past experiences.