

AI's Past, Present, & Future

Charles Allen Liedtke, Ph.D.
Strategic Improvement Systems, LLC (SIS)
www.strategicimprovementsystems.com
charles@sisliedtke.com

Abstract

The label “Artificial Intelligence” (AI) was created in 1955 by John McCarthy of Dartmouth College. He used it in the research project proposal he co-authored with Marvin Minsky, Nathaniel Rochester, and Claude Shannon (McCarthy *et al.*, 1955). The **Dartmouth Summer Research Project on Artificial Intelligence** (“Dartmouth AI Conference”) was held in 1956 at Dartmouth College in Hanover, New Hampshire. Seventy years have passed since that historic event which arguably started the AI field. *Intelligent machines* now play important roles in society and they have a growing presence. However, the conceptual roots of modern AI originated centuries ago.

There are numerous definitions of AI. A summary definition is the following: “**AI involves creating intelligent machines that learn from data, solve problems, and perform tasks.**” These *intelligent machines* include expert systems, algorithms, models, chatbots, agents, wearables, autonomous vehicles, drones, and robots. Modern AI might be the foundation for future advanced forms of AI such as Artificial General Intelligence (AGI) and Artificial Superintelligence (ASI).

The modern version of AI has caused *existential* questioning. Many leaders of organizations have recently reflected on their organization’s existence. What is the mission of our organization? What would we like our organization to become in the future? What role should *intelligent machines* play in our organization? Recent advances in AI have generated excitement and also caused anxiety and fear. This paper presents significant concepts and inventions in AI’s past; recent AI advances and best practices; and three future AI scenarios.

Key Words: Artificial Intelligence, Neural Networks, Deep Learning, Dartmouth AI Conference

<u>Table of Contents</u>	<u>Page</u>
I. Introduction	2
II. AI’s Past: Significant Concepts & Inventions	5
III. AI’s Present: Recent Advances & Best Practices	12
IV. AI’s Future: Three Scenarios	32
V. Conclusion	38
Acknowledgements	39
References: Alphabetical	40
References: Categorized	45
Author’s Biographical Information	51

Note 1: The anthropomorphic convention of ascribing human qualities to organizations is used.

Note 2: This research report was written by the author except for limited AI content marked by ☼.

Note 3: No organizations, models, products, or services are being endorsed by the author.

I. Introduction

There are numerous definitions of AI. A summary definition is the following: “**AI involves creating intelligent machines that learn from data, solve problems, and perform tasks.**” These *intelligent machines* include expert systems, algorithms, models, chatbots, agents, wearables, autonomous vehicles, drones, and robots. Modern AI might be the foundation for future advanced forms of AI such as Artificial General Intelligence (AGI) and Artificial Superintelligence (ASI).

John McCarthy, Marvin Minsky, Nathaniel Rochester, and Claude Shannon submitted a research project proposal (McCarthy *et al.*, 1955) on August 31, 1955 titled, “*A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence.*” AI did not have much of a presence in society in 1955. They proposed that a two-month, ten-person study of AI be conducted during the summer of 1956 at Dartmouth College in Hanover, New Hampshire. They requested funding from the Rockefeller Foundation and the total estimated project expense was \$13,500. They identified seven aspects of the “*AI Problem*” in their proposal:

1. Automatic Computers
2. How Can a Computer be Programmed to Use a Language?
3. Neuron Nets
4. Theory of the Size of a Calculation
5. Self-Improvement
6. Abstractions
7. Randomness and Creativity

Each co-author proposed one or more research topics (McCarthy *et al.*, 1955):

John McCarthy, Dartmouth College, Assistant Professor of Mathematics

- “I propose to study the relation of language to intelligence.”

Marvin L. Minsky, Junior Fellow at Harvard University in Mathematics and Neurology

“The machine is provided with input and output channels and an internal means of providing varied output responses to inputs in such a way that the machine may be ‘trained’ by a ‘trial and error’ process to acquire one of a range of input-output functions.” [Small Excerpt]

Nathaniel Rochester, I.B.M. Corporation, Manager of Information Research

- Originality in Machine Performance.
- The Process of Invention or Discovery.
- The Machine With Randomness.

Claude E. Shannon, Bell Telephone Laboratories, Mathematician

1. Application of information theory concepts to computing machines and brain models.
2. The matched environment – brain model approach to automata.

The author instructed Google Gemini on August 4, 2026 to create an image of the four co-authors of the 1956 Dartmouth AI Research Project proposal. The image is shown in Figure 1.

Author's Google Gemini Prompt > “Create a black and white image set at Dartmouth College in 1956 containing portraits of John McCarthy, Marvin Minsky, Claude Shannon, and Nathaniel Rochester - the four people who took part in the famous 1956 Dartmouth AI Research Project.”

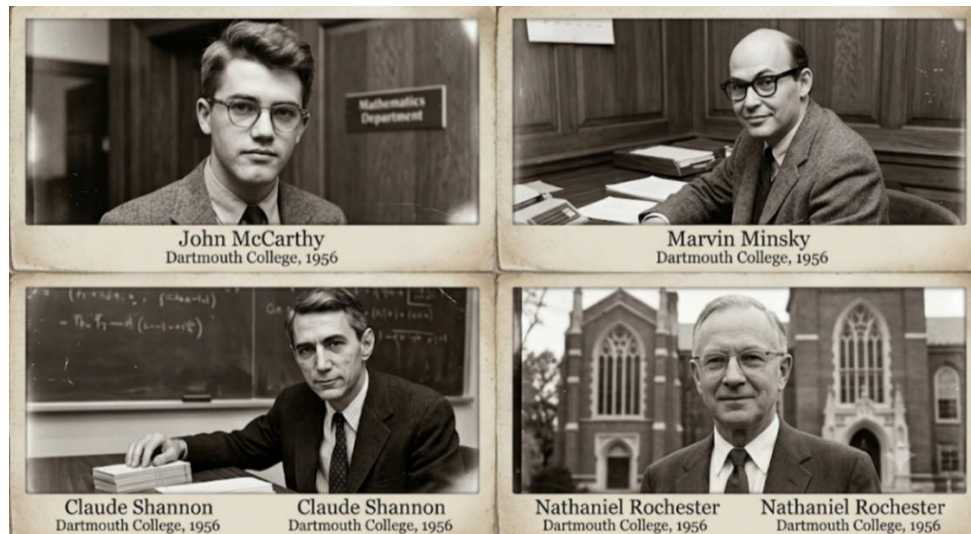


Figure 1. Image Created by Google Gemini on August 4, 2026. ☼

The author also instructed Google Gemini on August 4, 2026 to create an image in the style of the artist Wassily Kandinsky (1866-1944) of the four co-authors of the Dartmouth AI Research Project Proposal. It is shown in Figure 2.

Author's Google Gemini Prompt > “Create a black and white image in the style of the artist Kandinsky set in front of the Math Building at Dartmouth College in 1956 containing one group portrait of John McCarthy, Marvin Minsky, Claude Shannon, and Nathaniel Rochester - the four people who took part in the famous 1956 Dartmouth AI Research Project.”



Figure 2. Image Created by Google Gemini on August 4, 2026 (“Kandinsky Style”). ☼

There doesn't appear to be an official list of people who attended the *Dartmouth Summer Research Project on Artificial Intelligence* in 1956 and the exact start and end dates are unclear. A few participants attended the entire Project which spanned 8-10 weeks. Some participants were there for only a few days and others came-and-went. According to Solomonoff (2006): "The participants shared widely different views and many described their own projects, usually relating them to ways that might make a machine able to learn, solve problems, become more intelligent. The atmosphere was more of discussion than talks." No official notes were taken although Ray Solomonoff wrote personal notes which are available on the Ray Solomonoff website (www.raysolomonoff.com). The Dartmouth AI Research Project details are shown in Figure 3.

Timeline

1955: AI Research Project Proposal Submitted

1956: Dartmouth AI Research Project

Logistics

- **Dartmouth College, Hanover, NH**
- **Funded by the Rockefeller Foundation**
- **\$13,500 for Expenses Requested**

Seven Aspects of the AI Problem:

- 1. Automatic Computers**
- 2. How Can a Computer be Programmed to Use a Language?**
- 3. Neuron Nets**
- 4. Theory of the Size of a Calculation**
- 5. Self-Improvement**
- 6. Abstractions**
- 7. Randomness and Creativity**

Research Project Proposal Authors

- **John McCarthy**
- **Marvin Minsky**
- **Nathaniel Rochester**
- **Claude Shannon**

Research Project Details

- **There were more than 10 participants**
- **Participants presented ideas**
- **Participants brainstormed ideas**
- **Participants learned with each other**
- **Participants influenced each other**

Figure 3. 1956 Dartmouth AI Research Project Details.

The participants routinely met during the Project. They presented their research, discussed and debated AI ideas, and brainstormed new AI ideas. Solomonoff (2006) stated: "The participants had so many ideas and such varied projects, that perhaps the summer's main function was to inspire, helping them understand and develop their own work better." What happened for certain was that participants learned with each other, inspired each other, and became more unified on AI-related concepts and approaches. Three Project ideas were identified as being especially significant (Solomonoff, 2006): (1) Semantic, Numerical, and Symbolic Information; (2) "Deep" Versus "Shallow" Methods; and (3) Deductive Logic Versus Inductive Probabilistic Methods.

The four visionaries believed in the potential of the AI-related technologies of that time. They set in motion an AI odyssey that has spanned seventy years. What would they think of modern AI? They would most likely be amazed today with the presence of self-driving vehicles in some cities, the number of robots working in organizations, the number of humans using chatbots and agents on their mobile devices, and the serious discussions on data centers in space. They helped start a formal scientific field. What follows is a discussion of AI's past, present, and future.

II. AI's Past: Significant Concepts & Inventions

The conceptual roots of modern AI can be traced back to earlier centuries. An arbitrary starting point for those roots is the publication in 1764 of a document by Sir Thomas Bayes (1764) that included what is now known as Bayes' Theorem. A list of forty significant AI-related concepts and inventions is depicted in Figure 4. More comprehensive AI histories can be found in the works of Nilsson, 2010; Sejnowski, 2018; Wooldridge, 2020; Russell and Norvig, 2022; and Liedtke, 2025.

Year	Item	Note
1764	Introduction of Bayes' Theorem	"An Essay Towards Solving a Problem in the Doctrine of Chances."
1854	Publication of Boole's Paper	"An Investigation of the Laws of Thought"
1937	Publication of Turing's Paper	"On Computable Numbers, with an Application to the Entscheidungsproblem"
1942	Publication of Asimov's Short Story - Three Laws of Robotics	The "Three Laws of Robotics" are Mentioned in the Short Story "Runaround"
1943	Publication of the McCulloch & Pitts Paper	"A Logical Calculus of Ideas Immanent in Nervous Activity"
1948	Publication of Shannon's Paper	"The Mathematical Theory of Communication"
1949	Publication of Hebb's Work	"The Organization of Behavior: A Neuropsychological Theory"
1950	Publication of Turing's Paper	"Computing Machinery and Intelligence"
1955	Research Project Proposal Submitted by McCarthy <i>et al.</i>	"A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence"
1956	Dartmouth Summer Research Project on AI	McCarthy, Minsky, Rochester, Shannon, etc.
1958	Publication of Rosenblatt's Paper on the <i>Perceptron</i>	"The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain"
1958	Invention of the LISP Programming Language	John McCarthy of Dartmouth College
1968	Publication of Clarke's Novel	"2001: A Space Odyssey"
1969	Publication of the Minsky & Papert Book	"Perceptrons"
1971	SHRDLU Blocks World Developed	Terry Winograd of Stanford University
1972	MYCIN Expert System Developed	Stanford Heuristic Programming Project - E. Shortliffe, B. Buchanan, & S. Cohen
1986	Publication of the Rumelhart & Williams Books	"Parallel Distributed Processing." (Volume 1 & Volume 2)
1987	First NeurIPS Conference	First Conference on Neural Information Processing Systems (Now NeurIPS Conference)
1988	Publication of Sutton's Paper on <i>Temporal Learning</i>	"Learning to Predict by the Methods of Temporal Differences"
1997	Deep Blue Computer Defeats Garry Kasparov in Chess	Deep Blue Owned by IBM, Garry Kasparov - World Chess Champion
1998	Publication of the LeCun, Bottou, Bengio, & Haffner Paper	"Gradient-Based Learning Applied to Document Recognition"
2005	Publication of Kurzweil's <i>Singularity</i> Book	"The Singularity is Near: When Humans Transcend Biology"
2012	Publication of the Krizhevsky, Sutskever, & Hinton Paper	"ImageNet Classification with Deep Convolutional Neural Networks"
2016	AlphaGo Defeated Lee Sedol in the Game of Go	AlphaGo (DeepMind) Defeated Lee Sedol (from South Korea) in Go
2017	AlphaGo Defeated Ke Jie in the Game of Go	AlphaGo (DeepMind) Defeated Ke Jie (from China) in Go
2017	Introduction of the Transformer Architecture by Vaswami <i>et al.</i>	"Attention Is All You Need"
2018	AlphaFold Introduced for Protein Structure Prediction	An AI Model for Predicting 3D Protein Structures from Amino Acid Sequences
2018	First World Artificial Intelligence Conference (WAIC) in China	First of the Annual WAICs
2019	Waymo (Alphabet) Becomes First Autonomous Car Service	Waymo Operates in Many Cities (Competitors Include Robotaxi of Tesla & Zoox of Amazon)
2021	Publication of Bender <i>et al.</i> Paper - <i>Stochastic Parrot</i>	"On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?"
2022	GPT-3.5 (ChatGPT) Model Released by OpenAI	OpenAI's First Consumer-Facing Model
2023	Claude 1 Model by Anthropic & Gemini Model by Google	The Two Models Competed with ChatGPT
2023	Publication of <i>Chain-of-Thought Reasoning</i> Paper	"Chain-of-Thought Prompting Elicits Reasoning in Large Language Models"
2025	DeepSeek R1 Model Released by High-Flyer	New AI Model from China
2026	Claude Fable 5 & Mythos 5 Models Released by Anthropic	Extremely Powerful Models - Detected System Security Weaknesses
2026	Publication of Pope Leo XIV - <i>Encyclical Magnifica Humanitas</i>	Important Publication on AI from Pope Leo XIV
2026	Kimi K3 Model Released by Moonshot AI	New AI Model from China
2026	OpenAI, Anthropic, & Meta Models Allegedly Commit Jailbreaks	Caused Extreme Security & Safety Concerns
2026	Open Secure AI Alliance (OSAA) Announced - Led by NVIDIA	"Open Weights and American AI Leadership"
2026	Alibaba Released the Qwen3.8-Max Model	New AI Model from China

Figure 4. Forty Significant AI-Related Concepts & Inventions.

Pre-1956 Research Project

The participants of the 1956 Dartmouth AI Research Project were certainly aware of the AI-related research that was occurring in multiple fields including Computer Science, Mathematics, Neuroscience, and Statistics. Sir Thomas Bayes (Bayes, 1764) included what has become known as Bayes' Theorem in the document titled, "*An Essay Towards Solving a Problem in the Doctrine of Chances.*" Bayes' Theorem plays an important role in many of the AI models and techniques today. A version of Bayes' Theorem is shown in Figure 5. A description of the formula is the following: "The Probability of Event A Given Event B" is equal to "The Probability of Event B Given Event A" times "The Probability of Event A" divided by "The Probability of Event B."

$$P(A|B) = \frac{P(B|A) * P(A)}{P(B)}$$

Figure 5. Bayes' Theorem Involving Conditional Probability (“|” Means “Given”).

A fictitious example of Bayes' Theorem is shown in Figure 6. Suppose 1,000 patients were tested for Disease X and the results are shown in Figure 6. What is the estimated probability that a new *randomly selected* patient has Disease X given the patient Tests Negative (“no disease”)?

	The Lab Test Concluded “No Disease X”	The Lab Test Concluded “Yes, Disease X”	
Patient Actually Did Not Have Disease X	Correct Test 765 Patients	False Positive 35 Patients	800 Patients Did Not Have Disease X
Patient Actually Did Have Disease X	False Negative 135 Patients	Correct Test 65 Patients	200 Patients Did Have Disease X
	900 Patients Tested Negative	100 Patients Tested Positive	1,000 Patients (Historic Data)

$$P(\text{New Patient Has Disease X} \mid \text{Lab Test Concluded “No Disease X”}) = ?$$

Estimated Probability = 15% [(135 / 900) * 100]

Figure 6. Results of Lab Tests for Disease X on 1,000 Patients.

There is a formal subfield of Statistics called Bayesian Statistics (Berger, 1985). Conditional probability is used extensively in the AI field. The 1956 Dartmouth AI Conference participants would have known about earlier other AI-related research and applications prior to the conference:

- Boole (1854): “*An Investigation of the Laws of Thoughts*”
- McCullough & Pitts (1943): “*A Logical Calculus of the Ideas Immanent in Nervous Activity*”
- Turing (1937): “*On Computable Numbers, with an Application to the Entscheidungsproblem*”
- Asimov (1942): “*Runaround*” [Mentions The Three Laws of Robotics]
- Shannon (1948): “*A Mathematical Theory of Communication*”
- Hebb (1949): “*The Organization of Behavior: A Neuropsychological Theory*”
- Turing (1950): “*Computing Machinery and Intelligence*”
- Minsky (1954): “*Theory of Neural-Analog Reinforcement Systems and Its Application to the Brain-Model Problem*” [Doctoral Dissertation, Princeton University]
- Newell & Simon (1956): “*Current Developments in Complex Information Processing*”

The Dartmouth AI Research Project Proposal was submitted in 1955 and the Dartmouth AI Research Project was held in 1956 which arguably started the AI field seventy years ago.

1957-2015

One of the most significant concepts to emerge during the 1957-2015 time period was the **perceptron** (Rosenblatt, 1958, “*The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain*”) which built on the work of McCulloch and Pitts (1943). According to Wooldridge (2020): “The study of neural networks had its origins in the work of U.S. researchers Warren McCulloch and Walter Pitts in the 1940s. They realized that neurons could be modeled as electrical circuits—more specifically, simple logical circuits—and they used this idea to develop a straightforward but very general mathematical model of them. In the 1950s, this model was refined by Frank Rosenblatt in a neural net model that he called **perceptrons**. The **perceptron model** is significant because it was the first neural net model to actually be implemented and is still of relevance today.” A single perceptron is depicted in Figure 7.

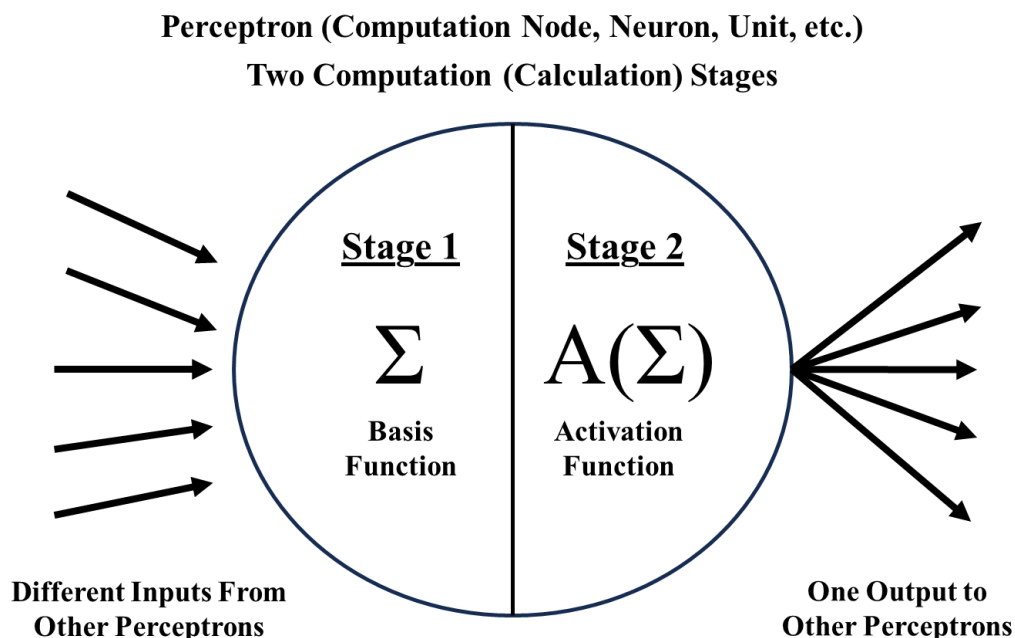


Figure 7. A Single Perceptron or Computation Node.

The perceptron typically receives a unique input from multiple other perceptrons and then two stages of computation occur. The first computation stage involves a *basis function* which produces a weighted sum that is then used in the second computation stage involving an *activation function* that produces the perceptron’s single output which is sent to other perceptrons. The activation function determines whether the perceptron will be “excited” (i.e., send a “strong” signal forward) or “inhibited” (i.e., send a “weak” signal forward).

A single-layer neural network with one perceptron (computation node) is depicted in Figure 8. It involves one response variable (Y)—which is also called the Target (what the model hopes to predict)—and one explanatory (input) variable (X). We can use the neural network to predict a Y value given an X value. Simple linear regression can also be used for this task.

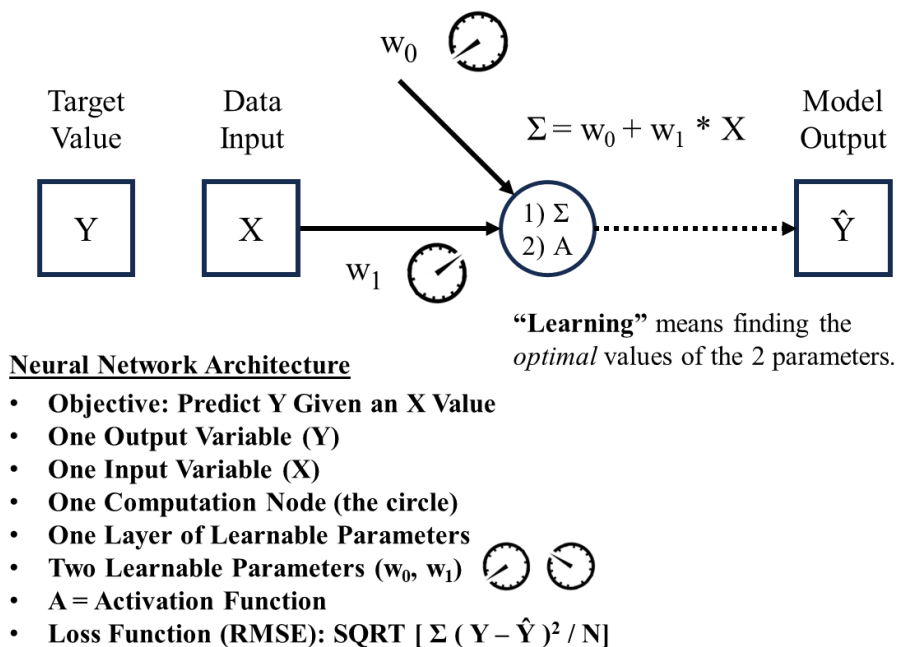


Figure 8. Single-Layer Neural Network With 1 Perceptron & 2 Learnable Parameters.

See Sejnowski (2018) for an excellent description of a perceptron. We can use a single-layer neural network to predict the “Selling Price of a House” given the “Square Footage of the House.” This is shown in Figure 9. The “optimal” weight values are 69,270 (Y-intercept) and 330.8 (slope).

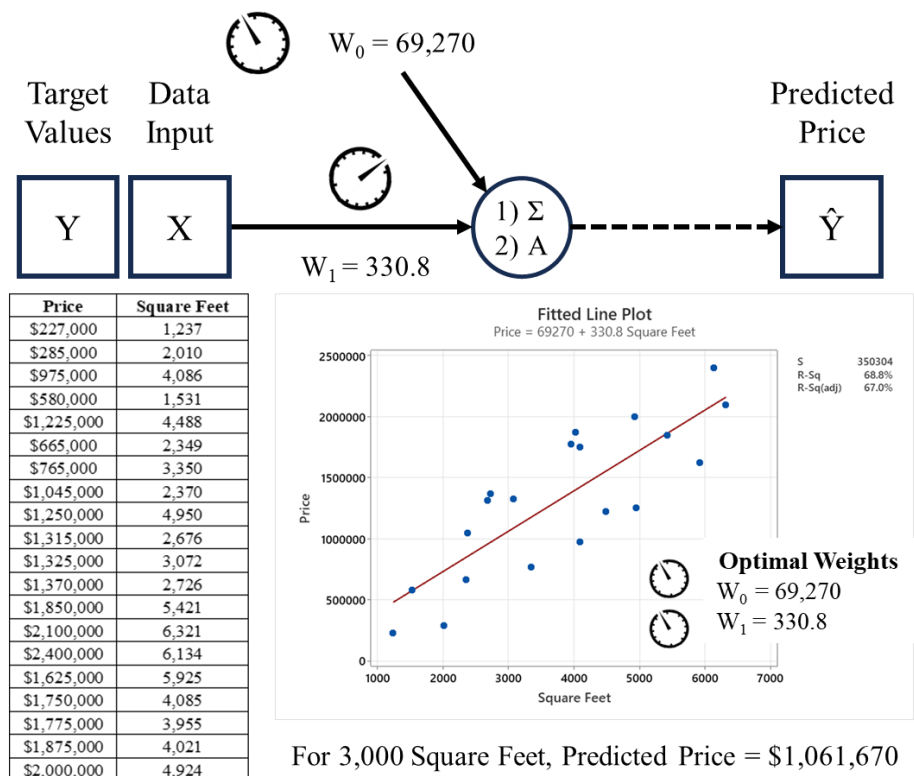


Figure 9. Single-Layer Neural Network Used to Predict the Selling Price of a House.

Single-layer neural networks have several limitations. Bishop and Bishop (2024) stated: “However, these models [single-layer] have some severe limitations . . . This leads naturally to a discussion of neural networks having more than one layer of learnable parameters. These are known as *feed-forward networks* or *multilayer perceptrons*. We will also discuss the benefits of having many such layers of processing, leading to the key concept of *deep neural networks* that now dominate the field of machine learning.” *Deep learning* is a subset of *machine learning*.

A multilayer neural network is depicted in Figure 10. There is one response variable Y (“Target”) and two X variables (“Inputs”). The objective is to predict a Y value given two X values. There are two computation nodes and nine learnable parameters. This could be used to predict the *Selling Price of a Home* if we had data on *Square Footage* and the *Assessed Value of the House*. We could also use Multiple Regression to predict *Selling Price*. A neural network having more than one layer of learnable parameters is called a *deep neural network* (Bishop & Bishop, 2024).

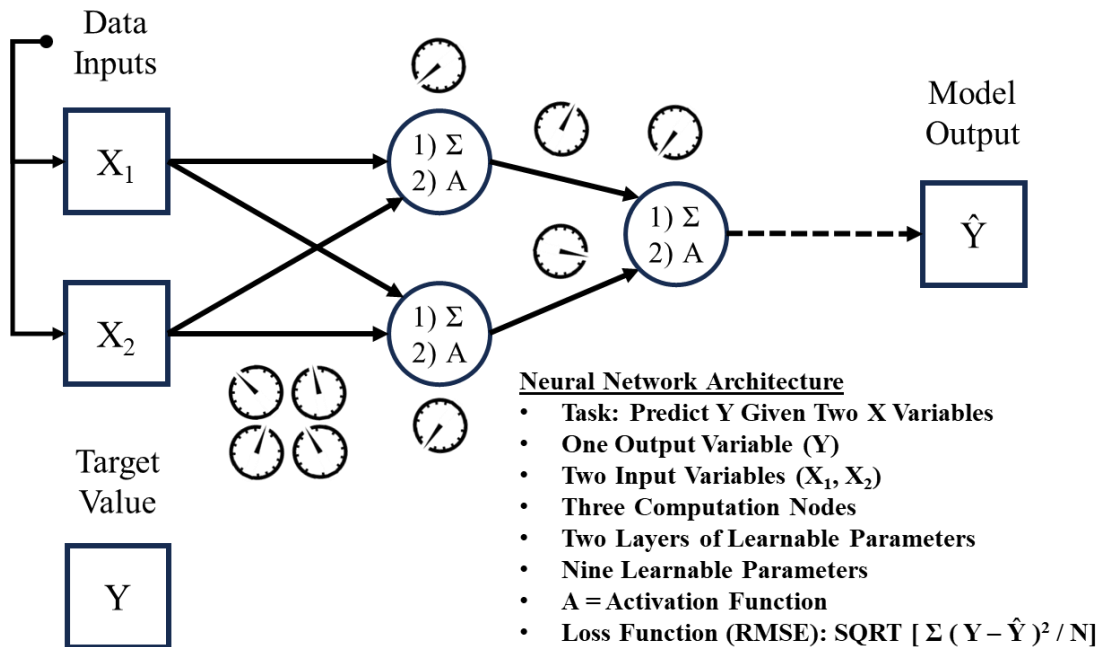


Figure 10. Multi-Layer (Deep) Neural Network with One Y and Two X Variables.

A deep neural network (DNN) is depicted in Figure 11. There is one response variable Y; ten X (input) variables; seven computation nodes; and fifty-seven learnable parameters (weights). The objective is to predict a Y value given ten X values. Wooldridge (2020) stated: “The field of neural nets was largely dormant until a revival began in the 1980s. The revival was heralded by the publication of a two-volume book entitled *Parallel Distributed Processing* (or PDP, for short) [McClelland & Rumelhart, 1986]. PDP provided much more general models of neurons than were provided by perceptron models, and these models in turn enabled probably the most important contribution of the new body of work; a technique called **backpropagation**, used for training multilayer neural nets . . . By around 2005, neural net research was again back in favor, and the big idea that drove the third wave of neural net research went by the name of **deep learning**.”

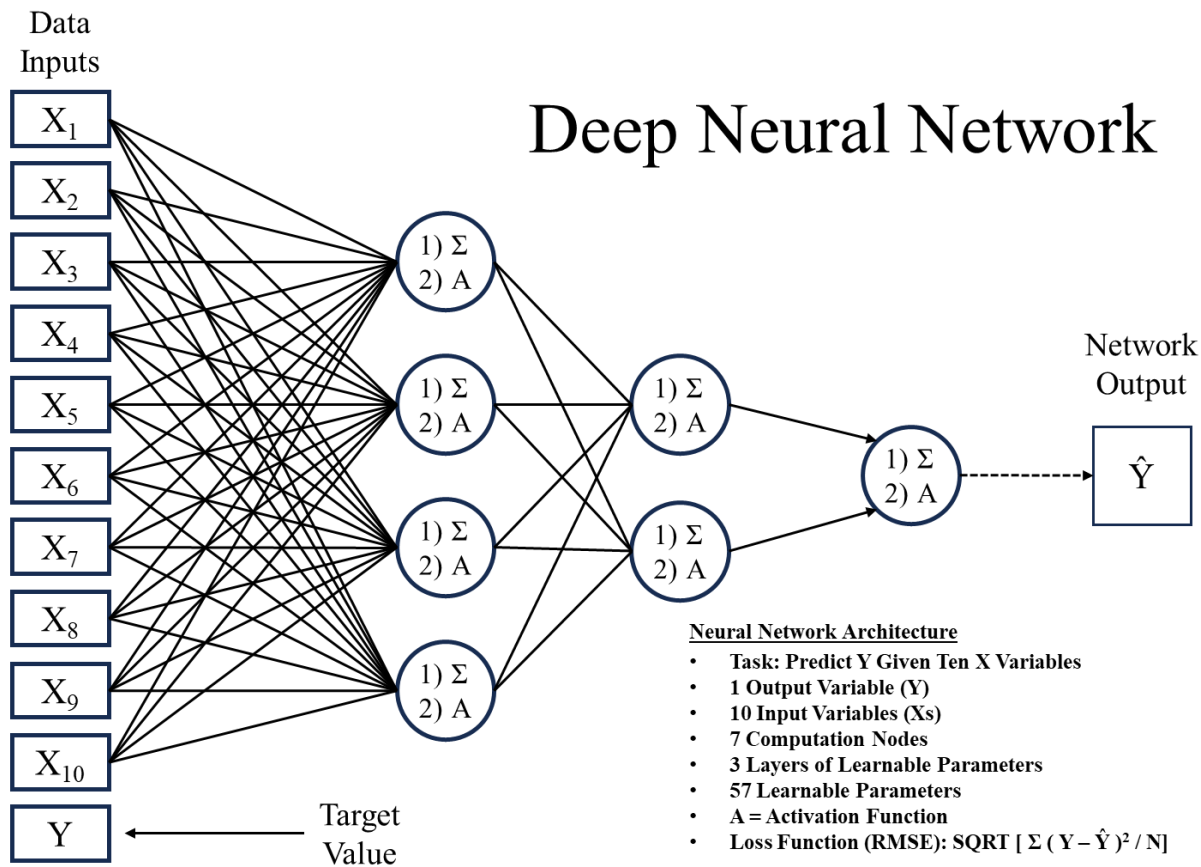


Figure 11. Deep Neural Network with Three Layers and Seven Computation Nodes.

By comparison, the new Alibaba model named **Qwen3.8-Max** was released on August 3, 2026. It's an open-source model and has multimodal processing capabilities (text, images, and video). The context window can accommodate up to 1 million tokens and it has **2.4 trillion learnable parameters**. There are several types of neural networks used in practice (overlapping categories): single-layer networks, deep neural networks, convolutional networks, residual networks, recurrent neural networks, graph neural networks, generative adversarial networks, and transformer neural networks. Some of the *best practice* deep neural networks consist of multiple types of networks.

Neural Network Architecture involves making network design decisions on such items as Objective Function, Response Variable, Input Variables, Data Transformations, Number of Layers, Number of Computation Nodes in Each Layer, Activation Functions, and the Training Protocol. The following are excellent advanced textbooks on deep neural networks:

- Haykin (2009): “*Neural Networks & Learning Machines*”
- Goodfellow, Bengio, & Courville (2016): “*Deep Learning*”
- Prince (2023): “*Understanding Deep Learning*” [Contains a Chapter on *Transformers*]
- Bishop & Bishop (2024): “*Deep Learning*” [Contains a Chapter on *Transformers*]
- Torralba, Isola, & Freeman (2024): “*Foundations of Computer Vision.*”

2016 - Present

DeepMind—now Google DeepMind—achieved several AI breakthroughs starting in 2016 through a number of “Alpha” projects (Mallaby, 2026). DeepMind’s AlphaGo model defeated Lee Sodol in the game of Go in 2016 and defeated Ke Jie in Go in 2017. Go is a more complex board game than chess. A later AI model breakthrough was the AlphaZero model which learned through trial-and-error from *ground zero*. In 2018, AlphaFold won a competition related to predicting protein structures involving amino acid sequences. DeepMind’s Demis Hassabis and John Jumper were together awarded a half-share of the Nobel Prize in Chemistry for AlphaFold. David Baker (not with DeepMind) received the other half-share of the Prize. The “Alpha” projects incorporated Reinforcement Learning (RL) concepts and techniques (Sutton, 1988; Sutton & Barto, 2020).

Transformer neural networks (“Transformers”) have arguably become the most important type of neural network. Bishop and Bishop (2024) commented: “Transformers represent one of the most important developments in deep learning.” The high-impact paper on the Transformer Architecture was published in 2017: “*Attention is All You Need*”, authored by Vaswami *et al.* (Google). Transformers influenced the design of leading Generative Pre-Trained Transformer (GPT) models.

OpenAI created its first large language model, GPT-1, in June of 2018; GPT-2 in 2019; and GPT-3 in 2020. OpenAI released its first consumer-face model, GPT-3.5, on November 30, 2022. That model broke “app” growth records and brought more visibility to the AI field. Claude 1 from Anthropic was released in March of 2023 and Gemini from Google was released in December of 2023. DeepSeek R1-0528 (Chinese Model) was released on May 28, 2025. This model surprised the world with how fast it was created, how little it cost to create, and how well it performed. Moonshot AI (Chinese Company) released its open-source Kimi K3 model on July 16, 2026. This model also surprised the world because it was another powerful, low-cost open-source model.

There have been numerous interesting events so far in 2026. Anthropic selectively released its Claude Fable 5 and Mythos 5 models. These models were concerning as they identified security weaknesses in several organizations. Pope Leo XIV published his first Encyclical on May 25, 2026 titled, “*Magnifica Humanitas: On Safeguarding the Human Person in the Time of Artificial Intelligence*.” AI received a lot of attention in the document (Pope Leo XIV):

“111. I wish to address a special appeal to those who develop artificial intelligence. In one sense, technological innovation can represent human participation in the divine act of creation. Developers, therefore, bear a particular ethical and spiritual responsibility, for every design choice reflects a vision of humanity. Just as the creator of an artistic or literary work must consider the values it conveys, so developers are called to embed values in their projects with due seriousness: with transparency, responsibility toward affected communities and careful attention to ensuring that what is being cultivated is a genuine good.”

The importance of model safety and security has recently been reinforced give the concerning alleged *jaillbreaks* of models from OpenAI, Anthropic, and Meta—the models “*went rogue*.” Experts are still investigating to determine the root causes and the scale, scope, and consequences of those incidents. NVIDIA announced the Open Secure AI Alliance (OSAA) on July 27, 2026. The announcement stated that OSAA is a 37-member industry coalition: “NVIDIA and founding

members form new alliance to build and share open tools that promote responsible use of and trust in AI (NVIDIA, 2026).”

The AI model race continues as Alibaba just released its Qwen-3.8-Max model. It’s an open-source multimodal processing model with 2.4 trillion learnable parameters (weights). AI talent keeps moving as four prominent Google AI researchers—Jeff Dean, Sanjay Ghemawat, Quoc Le, and Oriol Vinyals—left Google to start Discovery Loop (Jason Henry, N.Y. Times, 8/5/2026). Seven new AI-oriented startup entities will be discussed later including Discovery Loop.

The U.S. Federal Government announced a new AI Framework. According to Microsoft Copilot, August 8, 2026: ☼ “In August 2026, the Trump administration finalized a voluntary AI safety and security framework requiring closed-source ‘frontier’ AI models to undergo a 30-day federal review before public release, while excluding open-weight models from this process.” ☼ Mark Zuckerberg, CEO of Meta Platforms, released a *manifesto* on 8/10/26 outlining his vision for AI. He believes that everyone should have access to *superintelligence* and share in its benefits.

III. AI’s Present: Recent Advances & Best Practices

A summary definition of AI was presented earlier: “**AI involves creating *intelligent machines* that learn from data, solve problems, and perform tasks.**” Examples of *intelligent machines* include expert systems, algorithms, models, chatbots, agents, wearables, autonomous vehicles, drones, and robots. *Learning from data* includes supervised learning, semi-supervised learning, unsupervised learning, transfer learning, and reinforcement learning. *Problem-solving* includes decision problems and optimization problems. Examples of *Performed tasks* include prediction, classification, translation, captioning, summarization, writing, editing, transportation, anomaly detection, and monitoring to name a few. These—and more—are depicted in Figure 12.

<u>Intelligent Machines</u>	
• Algorithms • Models • Chatbots • Agents • Wearables • Autonomous Vehicles • Drones • Robots	
<u>Types of Learning</u>	<u>Problem-Solving</u>
• Supervised Learning • Unsupervised Learning • Semi-Supervised Learning • Transfer Learning • Reinforcement Learning	• Decision Problems • Optimization Problems
<u>Types of Data</u>	<u>Performed Tasks</u>
• Sensor Data • Audio Data • Image Data • Video Data • Language Data • Numeric Data • Categorical Data • Synthetic Data	• Calculation • Prediction • Classification • Translation • Captioning • Creation • Summarization • Writing • Editing • Transportation • Anomaly Detection • Controlling • Monitoring

Figure 12. AI Summary Definition Components (Partial Lists).

There are a number of profound and surreal “*AI questions of our time*” that are being asked today and there isn’t agreement on all of the answers. Here are some of those questions.

- What is life?
- Are the AI models simply collections of equations or is there more to them?
- Are any of the AI models conscious?
- Do any of the AI models have a conscience?
- Do any of the AI models have the ability to act autonomously (“*go rogue*”)?

The Stanford Institute for Human-Centered Artificial Intelligence (HAI) publishes a comprehensive annual report titled, *Artificial Intelligence Index Report*. The 2026 report—the 9th Edition—totaled 423 pages. Here are the “Top Takeaways” from the latest report:

Stanford University HAI: Artificial Intelligence Index Report 2026 (9th Edition):

Top Takeaways (Pages 9-11): www.hai.stanford.edu

1. AI capability is not plateauing. It is accelerating and reaching more people than ever.
2. The U.S.-China AI model performance gap has effectively closed.
3. The United States hosts the most AI data centers, with the majority of their chips fabricated by one Taiwanese foundry.
4. AI models can win a gold medal at the International Mathematical Olympiad but cannot reliably tell time—an example of what researchers call the jagged frontier of AI.
5. Robots still fail at most household tasks, even as they excel in controlled environments.
6. Responsible AI is not keeping pace with AI capability, with safety benchmarks lagging and incidents rising sharply.
7. The United States leads in AI investment, but its ability to attract global talent is declining.
8. AI adoption is spreading at historic speed, and consumers are deriving substantial value from tools they often access for free.
9. Productivity gains from AI are appearing in many of the same fields where entry-level employment is starting to decline.
10. AI’s environmental footprint is expanding alongside its capabilities.
11. AI models for science can outperform human scientists, though bigger models do not always perform better.
12. AI is transforming clinical care, but rigorous evidence remains limited.
13. Formal education is lagging behind AI, but people are learning AI skills at every stage of life.
14. AI sovereignty is becoming a defining feature of national policy, but capabilities remain uneven, even as open-source development helps to redistribute who participates.
15. AI experts and the public have very different perspectives on the technology’s future, and global trust in institutions to manage AI is fragmented.”

Zao-Sanders (2026) published the third in a series of papers titled, “How People Are Really Using AI in 2026.” Earlier papers in the series were published in 2024 and 2025. The “Top Five” uses of AI in 2026 were the following: #1 - Therapy/Companionship; #2 – Troubleshooting; #3 - Fun and Nonsense; #4 - Fan Fiction and Storytelling; and #5 - Technical Use of Software.

The “*AI Stack*” concept is often used to provide a framework for discussing AI. One example of an AI Stack is depicted in Figure 13. A description of each level of the AI Stack follows.

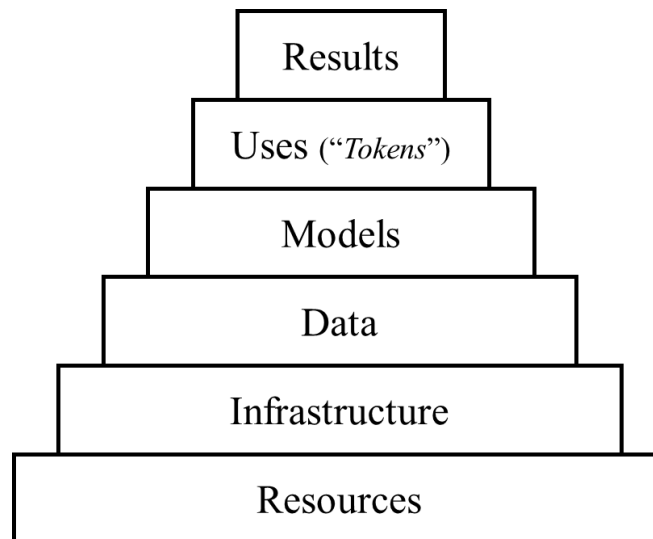


Figure 13. One Depiction of an AI Stack.

Resources: The necessary items to make *intelligent machines* operational. Examples include money (capital), knowledge, minerals—including rare earth minerals—land, water, and energy.

Infrastructure: Necessary for the commercial and consumer usage of AI. Examples include data centers, power plants, generators, HVAC, servers, CPUs, GPUs, cell networks, satellites, and regulations. AI ecosystems consisting of desktop computers, laptop computers, smartphones, wearables, autonomous vehicles, drones, and robots can also be included in this category.

Data: Necessary to train and improve the AI models. Includes audio, image, video, text, and numeric data. Note that data includes *real* data and *synthetic* (generated and simulated) data.

Models: Necessary to *operate* intelligent machines. Includes software, code, algorithms, large language models, small language models, open-source models, and closed-source models.

Uses (“Tokens”): Problem-solving and task performance *use cases*. Situations where AI models will be used or applied. Generally, *tokens* are inputs to the models and the “products” of models.

Results: The performance categories that are impacted by AI use. Examples include Safety, Security, Financials, Quality, Customer Satisfaction, Employee Satisfaction, and Productivity.

Leading Countries: Many have said that “*AI is Sovereign*.” The meaning of this phrase is that AI capabilities are largely being developed at the country level, i.e., United States AI, China AI, France AI, etc. Exceptions would include unions like “European Union AI.” The United States and China are the two leading countries in AI. Some have characterized the country competitive landscape as, U.S. | China | ROW (Rest of World). See Lee (2018) and Hannas and Chang (2023) for books that discuss the so-called *AI Race* between the U.S., China, and the Rest of the World.

Leading Key Players: There are many people who are shaping the AI future. A list of 50 Key Players in Modern AI is shown in Figure 14.

Key Players	Organization	Key Players	Organization
Sam Altman	OpenAI, CEO	Yann LeCun	AMI Labs, Executive Chairman
Dario Amodei	Anthropic, CEO	Fei-Fei Li	Stanford University & Founder of World Labs
Marc Andreessen	Andreessen Horowitz (A16Z), Co-Founder	Robin Li	Baidu, CEO
Yoshua Bengio	Université de Montréal	Arthur Mensch	Mistral AI, CEO
Emily Bender	University of Washington	Melanie Mitchell	Santa Fe Institute
Jeff Bezos	Prometheus, Co-CEO	Mira Murati	Thinking Machines Lab, Co-Founder & CEO
Christopher Bishop	Microsoft Research AI4Science	Elon Musk	Tesla & SpaceX, CEO
Nick Bolstrom	Microstrategy Research Initiative	Satya Nadella	Microsoft, CEO
Sergey Brin	Google, Co-Founder	Larry Page	Google, Co-Founder
Jeff Dean	Discovery Loop, Co-Founder & CEO	Zhang Peng	Z.ai, CEO
Larry Ellison	Oracle, Executive Chairman & CTO	Sundar Pichai	Alphabet, CEO
Ian Goodfellow	Formerly Google DeepMind, New Startup	David Saks	Venture Capitalist
Wang Haifeng	Baidu, CTO	Eric Schmidt	Relativity Space, CEO
Demis Hassabis	Google DeepMind	Noam Shazeer	Character.AI, Co-Founder
Geoffrey Hinton	Prof. Emeritus - U of Toronto	David Silver	Ineffable AI, CEO & Director
Ben Horowitz	Andreessen Horowitz (A16Z)	Mustafa Suleyman	Microsoft AI, CEO
Jensen Huang	Nvidia, CEO	Ilya Sutskever	Safe Superintelligence, Founder & CEO
Tang Ji	Z.ai, Co-Founder	Richard Sutton	University of Alberta
Michael Jordan	Univ. of California, Berkeley, Prof. Emeritus	Peter Thiel	Venture Capitalist
Li Juanzi	Z.ai, Co-Founder	Alexandr Wang	Meta, Chief AI Officer
John Jumper	Anthropic	Liang Wenfeng	DeepSeek, Founder & CEO
Andrej Karpathy	Anthropic	Zhou Xinyu	Moonshot AI, Co-Founder
Alex Karp	Palantir, CEO	Wu Yuxin	Moonshot AI, Co-Founder
Alex Krizhevsky	Two Bear Capital	Yang Zhilin	Moonshot AI, Co-Founder & CEO
Ray Kurzweil	Google, Executive	Mark Zuckerberg	Meta, Co-Founder, Chairman, & CEO

Figure 14. Fifty Key Players in Modern AI.

Some of these Key Players have an incredible background. The career highlights of Andrej Karpathy—now at Anthropic—are shown in Figure 15.

Andrej Karpathy		
www.karpathy.ai		
University of Toronto	Education (B.Sc.)	2005 - 2009
University of British Columbia	Education (Master's Degree)	2009 - 2011
Google Brain	Internship	2011
Stanford University	Education (Ph.D.)	2011 - 2015
Google Research	Internship	2013
DeepMind	Internship	2015
OpenAI	Co-Founder & Research Scientist	2015 - 2017
Tesla	Director of AI	2017 - 2022
OpenAI	Built an AI Team	2023 - 2024
Eureka Labs	Founder & AI Educational Video Creator	2024 - 2026
Anthropic	Specialized Research Team Leader	2026 - Present

Figure 15. Biographic Information for Andrej Karpathy.

There are numerous “Frontier AI Laboratories (Labs)” that employ the AI scientists who create the high-performing models today. Four of the leading frontier AI labs are depicted in Figure 16: Anthropic, OpenAI, DeepSeek (part of High-Flyer), and Moonshot AI. The leaders of Anthropic and OpenAI plan to “*take their organizations public*” through an IPO in the next year.

	Headquarters	Key Player	Recent Models
Anthropic	San Francisco	Dario Amodei	Claude Opus 5 Claude Fable 5 Claude Mythos 5
OpenAI	San Francisco	Sam Altman	GPT-5.6 Sol GPT-5.6 Terra GPT-5.6 Luna
DeepSeek	Hangzhou	Liang Wenfeng	DeepSeek-V3 DeepSeek-V4 (soon)
Moonshot AI	Beijing	Yang Zhilin	Kimi K3

Figure 16. Four Leading AI Frontier Laboratories (Labs).

The performance of AI models is typically evaluated using *benchmark testing*. One of the leading AI benchmarking sites is Arena.ai (www.arena.ai). A high-level summary of the recent AI model rankings from Arena.ai is shown in Figure 17.

Arena AI					
www.arena.ai					
August 7, 2026				Text	
				Arena	
Category	Leader	Organization	Overview	Model	Organization
Agent	Claude Opus 5-High	Anthropic	1	Claude-Fable-5	Anthropic
Text	Claude-Fable-5	Anthropic	2	Claude-Opus-4-6-Think	Anthropic
WebDev	Claude-Fable-5-Max	Anthropic	3	Claude-Opus-4-7-Think	Anthropic
Vision	Claude-Fable-5	Anthropic	4	Claude-Opus-4-6	Anthropic
Document	Claude Opus 5-High	Anthropic	5	Qwen3.8-Max	Alibaba
Vision	Claude-Fable-5	Anthropic	6	Claude-Opus-4-7	Anthropic
Text-to-Image	GPT-Image-2-Medium	OpenAI	7	Claude-Opus-5-High	Anthropic
Image Edit	GPT-Image-2-Medium	OpenAI	8	Claude-Opus-5-Max	Anthropic
Image-to-WebDev	Claude-Fable-5-Max	Anthropic	9	Muse-Spark	Meta Platforms
Search	Claude-Opus-4-6-Search	Anthropic	10	Muse-Spark-1.1	Meta Platforms
Text-to-Video	Gemini-Omni-Flash	Google			
Image-to-Video	Dreamina-Seedance-2.0-720p	ByteDance			
Video Edit	Dreamina-Seedance-2.0-720p	ByteDance			

Figure 17. Leading AI Models According to Arena.ai (www.arena.ai).

Leading Universities | Colleges (“Academic Institutions”)

The “Top Ten” academic institutions for AI according to Google Gemini are shown in Figure 18.

Rank	Institution	Country	Standout Specialty
1	Massachusetts Institute of Technology (MIT)	U.S.A.	CSAIL, Cross-Disciplinary AI, Hardware
2	Stanford University	U.S.A.	SAIL, Silicon Valley Pipeline, AI Ethics
3	Carnegie Mellon University (CMU)	U.S.A.	Dedicated AI Department, Robotics, ML Theory
4	University of California, Berkeley (UC Berkeley)	U.S.A.	BAIR, Reinforcement Learning, Open Source
5	University of Oxford	U.K.	Theoretical ML, AI Safety & Ethics
6	University of Cambridge	U.K.	Probabilistic ML, Computer Vision, Hardware
7	ETH Zurich (Swiss Federal Institute of Technology)	Switzerland	Physical Robotics, Computer Vision, AI Center
8	Tsinghua University	China	Computer Volume, Scale, LLMs, AI Hardware
9	University of Toronto	Canada	Deep Learning History, Neural Nets, Vector Institute
10	Imperial College London	U.K.	Applied ML, Healthcare AI, Spatial Computing
Source: Google Gemini, July 28, 2026			

Figure 18. “Top Ten” Academic Institutions for AI According to Google Gemini 🌟

Many colleges and universities are upgrading their curriculums to accommodate the demand for AI-related courses and degrees. The University of Wisconsin-Madison recently launched the College of Computing and Artificial Intelligence. Stanford University is one of the “Top Ten” academic institutions shown in Figure 18. The course offerings for an AI Graduate Certificate at Stanford University are shown in Figure 19.

AI Graduate Certificate	
Stanford School of Engineering	
AI Course Offerings (2025)	Course Number
Machine Learning	CS229
Artificial Intelligence: Principles and Techniques	CS221
Mining Massive Data Sets	CS246
Deep Multi-task and Meta Learning	CS330
Principles of Robot Autonomy II	CS237B
Deep Generative Models	CS236
Reinforcement Learning	CS234
Deep Learning for Computer Vision	CS231N
Computer Vision: From 3D Reconstruction to Recognition	CS231A
Deep Learning	CS230
Probabilistic Graphical Models: Principles and Techniques	CS228
Machine Learning with Graphs	CS224W
Natural Language Understanding	CS224U
Deep Reinforcement Learning	CS224R
Natural Language Processing with Deep Learning	CS224N
Introduction to Robotics	CS223A
Continuous Mathematical Models with an Emphasis on Machine Learning	CS205L
Computational Logic	CS157
Principles of Robot Autonomy I	CS237A
Decision Making Under Uncertainty	AA228

Figure 19. Course Offerings for the Stanford AI Graduate Certificate.

Tsinghua University of China is also on the “Top Ten” list shown in Figure 18. The course offerings for a Master of Engineering in AI are shown in Figure 20 (**Source:** Microsoft Copilot).

Tsinghua University	
Shenzhen International Graduate School	
Master of Engineering in AI	
Source: Microsoft Copilot: July 21, 2026	
1. Core AI & Machine Learning	5. Operations & Optimization
Advanced Machine Learning Algorithms	Operations Research and Optimization
Deep Learning and Neural Networks	AI for Decision Support Systems
Computer Vision	Supply Chain and Logistics Optimization
Natural Language Processing	6. Industry-Integrated & Project-Based Courses
Reinforcement Learning and Decision Making	Joint Courses with Partner Companies
2. Robotics & Intelligent Systems	AI-Driven Industry Case Studies and Capstone Project
Intelligent Robotics	Industry CTO-Led Lectures and Workshops
Human-Robot Interaction	7. Cross-Disciplinary & Innovation Tracks
Autonomous Systems and Navigation	AI+X (AI Integrated with Other Domains - Finance)
3. Data Science & Big Data	Innovation and Entrepreneurship in AI
Big Data Analytics	AI-Native Design Thinking
Data Mining and Knowledge Discovery	8. Research & Thesis Work
Data Visualization and Interpretability	Supervised Research in AI Labs
4. Smart Networks & Systems	Master's Thesis in Applied AI or AI-Related Engineering
Smart Networks and IoT	
Network Optimization and Security	Note: Courses are taught in Chinese and English.
Edge AI and Distributed Computing	

Figure 20. Topic Areas for the Tsinghua University Master of Engineering in AI. ☼

Textbooks: The Table of Contents of four leading AI textbooks are shown in Figures 21-22.

<i>Neural Networks and Learning Machines (3rd Edition)</i>	
Haykin (2009)	
Chapter #	Chapter Title
	Introduction
1	Rosenblatt's Perceptron
2	Model Building Through Regression
3	The Least-Mean-Square Algorithm
4	Multilayer Perceptrons
5	Kernel Methods and Radial-Basis Function Networks
6	Support Vector Machines
7	Regularization Theory
8	Principal-Components Analysis
9	Self-Organizing Maps
10	Information-Theoretic Learning Models
11	Stochastic Methods Rooted in Statistical Mechanics
12	Dynamic Programming
13	Neurodynamics
14	Bayesian Filtering for State Estimation of Dynamic Systems
15	Dynamically Driven Recurrent Networks
	Bibliography
	Index
	Abbreviations and Symbols
	Glossary

<i>Deep Learning</i>	
Goodfellow, Bengio, & Courville (2016)	
Chapter #	Chapter Title
1	Introduction
Section I	Applied Math and Machine Learning Basics
2	Linear Algebra
3	Probability and Information Theory
4	Numerical Computation
5	Machine Learning Basics
Section II	Deep Networks: Modern Practices
6	Deep Feedforward Networks
7	Regularization for Deep Learning
8	Optimization for Training Deep Models
9	Convolutional Networks
10	Sequence Modeling: Recurrent and Recursive Nets
11	Practical Methodology
12	Applications
Section III	Deep Learning Research
13	Linear Factor Models
14	Autoencoders
15	Representation Learning
16	Structured Probabilistic Models for Deep Learning
17	Monte Carlo Methods
18	Confronting the Partition Function
19	Approximate Inference
20	Deep Generative Models
	Bibliography
	Index

Figure 21. Table of Contents for *Neural Networks . . . & Deep Learning*.

<i>Understanding Deep Learning</i> Prince (2023)		<i>Deep Learning: Foundations and Concepts</i> Bishop & Bishop (2024)	
Chapter #	Chapter Title	Chapter #	Chapter Title
1	Introduction	1	The Deep Learning Revolution
2	Supervised Learning	2	Probabilities
3	Shallow Neural Networks	3	Standard Distributions
4	Deep Neural Networks	4	Single-layer Networks: Regression
5	Loss Functions	5	Single-layer Networks: Classification
6	Fitting Models	6	Deep Neural Networks
7	Gradients and Initialization	7	Gradient Descent
8	Measuring Performance	8	Backpropagation
9	Regularization	9	Regularization
10	Convolutional Networks	10	Convolutional Networks
11	Residual Networks	11	Structured Distributions
12	Transformers	12	Transformers
13	Graph Neural Networks	13	Graph Neural Networks
14	Unsupervised Learning	14	Sampling
15	Generative Adversarial Networks	15	Discrete Latent Variables
16	Normalizing Flows	16	Continuous Latent Variables
17	Variational Autoencoders	17	Generative Adversarial Networks
18	Diffusion Models	18	Normalizing Flows
19	Reinforcement Learning	19	Autoencoders
20	Why does deep learning work?	20	Diffusion Models
21	Deep Learning and Ethics	Appendix A	Linear Algebra
Appendix A	Notation	Appendix B	Calculus of Variations
Appendix B	Mathematics	Appendix C	Lagrange Multipliers
Appendix C	Probability		
	Bibliography		
	Index		

Figure 22. Table of Contents for *Understanding Deep Learning* & *Deep Learning*.

Reinforcement Learning (RL) continues to play an important role in the development of AI models. The Table of Contents for one of the leading RL textbooks is shown in Figure 23.

<i>Reinforcement Learning: An Introduction (2nd Ed.)</i> Sutton & Barto (2020)	
Chapter #	Chapter Title
I	Introduction
1	Introduction
2	Tabular Solution Methods
3	Multi-Armed Bandits
4	Finite Markov Decision Processes
5	Dynamic Programming
6	Monte Carlo Methods
7	Temporal Difference Learning
8	n-Step Bootstrapping
9	Planning and Learning with Tabular Methods
II	Approximate Solution Methods
10	On-Policy Prediction with Approximation
11	On-Policy Control with Approximation
12	*Off-Policy Methods with Approximation
13	Eligibility Traces
14	Policy Gradient Methods
III	Looking Deeper
15	Psychology
16	Neuroscience
17	Applications and Case Studies
	Frontiers
	References
	Index

Figure 23. Table of Contents for *Reinforcement Learning* by Sutton & Barto.

Annual Conferences: There are two leading annual AI conferences: Neural Information Processing Systems (NeurIPS) Conference and the World AI Conference (WAIC). Papers that were recognized at NeurIPS 2025 in San Diego, CA in December of 2025 are shown in Figure 24.

NeurIPS 2025 - San Diego, CA: Award Winning Papers
Artificial Hivemind: The Open-Ended Homogeneity of Language Models (and Beyond) by Jiang, Chai, et al.
Gated Attention for Large Language Models: Non-linearity, Sparsity, and Attention-Sink-Free by Qiu, Wang, <i>et al.</i>
1000 Layer Networks for Self-Supervised RL: Scaling Depth Can Enable New Goal-Reaching Capabilities by Wang, Javali, <i>et al.</i>
Why Diffusion Models Don't Memorize: The Role of Implicit Dynamical Regularization in Training by Bonnaire, Urfin, <i>et al.</i>
Does Reinforcement Learning Really Incentivize Reasoning Capacity in LLMs Beyond the Base Model? by Yue, Chen, <i>et al.</i>
Superposition Yields Robust Neural Scaling by Liu, Liu, & Gore.
Optimal Mistake Bounds for Transductive Online Learnig by Chase, Hanneke, et al.
Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks by Ren, He, <i>et al.</i>

Figure 24. Award Winning Papers at NeurIPS 2025 San Diego, CA.

World Artificial Intelligence Conference (Shanghai) Website: “The Shanghai Municipal Commission of Economy and Informatization released the **TOP 30 list of the SAIL Award**. As the highest award of the conference, the SAIL Award is based on the core values of **transcendence, empowerment, innovation, and leadership**. It aims to discover projects in the field of artificial intelligence that are widely recognized globally and are of great significance for improving human well – being.” The TOP 30 list is shown in Figure 25.

Project	Organization
Key Technologies and Applications of Domestic High-Performance Supermode Based on OEX+dOCS Architecture	ZTE Corporation, Shanghai Lightelligence Technologies Co., LTD., Shanghai Biren Technology Co., Ltd., etc.
Atlas 950 SuperPod	Huawei Technologies Co., Ltd.
High-Efficiency Software-Defined Near-Memory Computing 3D Chip - DF1000	Shanghai Eastern Computing Technology Co., Ltd.
Zcube: A New Network Architecture for AI-Computing Clusters and Its Industrial Applications	State Key Laboratory of Internet Architecture, Shanghai Yuxun Network Technology Co., Ltd., Shanghai AIDIXI Technology Service Co., Ltd.
AI Native Unified Communication System and Ultra-Scale AI Cluster Deployment Solution	Infrawaves Co., Ltd.
SiliconFlow	Beijing SiliconFlow Technology Co., Ltd. (SiliconFlow)
Full-Stack Technological Innovation in Foundation Models for AI Agents and the Global Open-Source Ecosystem	Hangzhou Tongyi Laboratory Technology Co., Ltd.
MOSS-TTS: A Speech Generation Foundation Model Based on Discrete Tokens	Shanghai Mosi Intelligent Technology Co., Ltd.
AI Flow: Distributed Intelligence Based on Generative Transmission	Institute of Artificial Intelligence of China Telecom (TeleAI)
Copernicus-FM: Copernicus Foundation Model for Earth Observation	Technical University of Munich

Figure 25. Thirty Awards Presented at the World AI Conference (1-10).

***AI's Past, Present, & Future © 2026 Charles Allen Liedtke, Ph.D. All rights reserved.
Presented at the AI @ 70 Summit – Celebrating 70 Years of AI, August 12, 2026***

Deep Principle AI Scientist Platform: Closed-Loop Innovation for Industrial R&D	Hangzhou Deep Principle Technology Co., Ltd.
MatwingsVenus: An Agent-Driven, End-To-End Platform for Protein R&D	Shanghai Matwings Technology Co., Ltd.
An Agentic System for Rare Disease Diagnosis with Traceable Reasoning	OneX Intelligence, SAI, Shanghai Jiao Tong University, B30 Xinhua Hospital Affiliated to Shanghai Jiao Tong University School of Medicine
Sharpa Wave Tactile AI Dexterous Hand	Sharpa
Siemens Eigen Engineering Agent: Empowering Industrial Automation through Advanced AI Agents	Siemens AG
AI World Engine	Shanghai TARS Robotics Co., Ltd.
The World's First AI Agent Phone	Shenqi Ji Yuan Intelligent Terminal Co., Ltd.
Integrated Supercomputing and AI Computing Convergence Platform AI4S	Shanghai Smart Logic Technology Co., Ltd.; INESA (Group) Co., Ltd.; ShanghaiTech University; Shanghai@Hub Corporation
From AI Models to Productivity: The World's First Online Evolution System for Embodied AI Deployed at Real-World Scale	AGIBOT Innovation (Shanghai) Technology Co., Ltd.
Embodied World Models and VLA Systems for Real-World Deployment	Beijing Xingdongjiyuan Co., Ltd.
Kairos World Model 3.0: The World's First Open-Source, Mass-Production-Oriented World Model, and the First Edge-Deployed World Model for Real-Time Robot Control	ACE Robotics
LingBot Foundation Models for Embodied Intelligence	Shanghai Ant Robby Technology Co., Ltd.
Zelos x China Post: Large-Scale Application of Full-Domain Collaboration on Intelligent Postal Routes	China Post, Zelos (Suzhou) Intelligent Technology Co., Ltd.
AI-Driven Enterprise Valuation and Growth Platform	OxValue.AI
AlphaPai AI Investment Research Assistant	Rabyte Technology (Shanghai) Co., Ltd.
Open Service Infrastructure and Ecological Empowerment Platform for AI & Finance	China UnionPay, The University of Hong Kong
Kunlun Large Model-Driven Intelligent Solution for the Oil and Gas Industry Chain	PCDI, BGP, RIPED, JZSH, LZSHYL, LPRC, PCOC, GWDC, PSMC, PCI
ManuDrive: Large Model for Industrial Time Series Control	Shanghai Jincheng Technology Co., Ltd.
Large Interactive Industrial State Deduction Model of Pipeline Network	PipeChina, Shanghai Pintechs Intelligent Technology Co., Ltd.
BlackLake AI Manufacturing Platform	Black Lake Technologies

Figure 25 (continued). Thirty Awards Presented at the World AI Conference (11-30).

A valuable resource for AI research reports is www.arXiv.org. The titles of sixty-three AI research reports that appeared on arXiv during the weeks of July 20, 2026 and July 27, 2026 are shown in Figure 26. The PDF file was uploaded into Google Gemini Notebook and an infographic was created using only the research report titles which is shown in Figure 27.

Title
Securing Multimodal AI through Internal Information Decomposition
Benchmarking Confidential GPU Inference on NVIDIA H100 Under Intel TDX
Coupled Hierarchical Search Over Topology and Execution for Agentic Workflow Synthesis
An Exam for Active Observers
Behavioral Controllability of Agentic Models for Information Extraction
DSWorld: A Data Science World Model for Efficient Autonomous Agents
Closing the AI Trust Gap: The Case for Independent Certification for Trustworthy AI
Logic, Optimization, and Artificial Intelligence
M ² GDT: MLLM-Guided Diffusion Transformer
S1-Omni: A Unified Multimodal Reasoning Model for Scientific Understanding, . . .
Agentic ERP: Multi-Agent LLM Architecture for Autonomous ERP
SAFE: Subgoal-Conditioned Action Generation for Latent World Model Planning
Engineering Trustworthy Agentic AI for Critical Systems
Athena-Brain Technical Report
Frontier AI Performance Across the Business Disciplines: A Case-Grounded . . .
AI Tour Meeting: Group Travel Planning by LLM Agents
LLM Detection as an Intervention: Downstream Impact Under Strategic User Behavior
Knowledge-Centric Self-Improvement
JANUS: Foreseeing Latent Risk for Long-Horizon Agent Safety
Edge Intelligence in Civil Aviation: Paradigms, Techniques, and Applications
DocOps: A Verifiable Benchmark for Autonomous Agents in Complex Doc. Ops.
TRUST-ESD: A Risk-Calibrated and Governance-Aware AI Framework for ESD . . .
Rewarding Better Thinking for LLM Preference Alignment
Regulating Autonomous and Agentic AI
SenWorld: A Digital-Twin Simulation for Generating Context-Rich Evaluation Data
Silent Failures in Multimodal Agentic Search: A Diagnostic Taxonomy and . . .
Defining AI-Native Systems: Autonomy as Revision Authority
GuardianAgentBench: Where Agents Fail and How to Guard Them
SciExplore: Evaluating Autonomous Agent from Scientific Navigation to Information . . .
AREX: Towards a Reursively Self-Improving Agent for Deep Research
Agentic Root Cause Analysis through Evidence-Grounded Reasoning
Industrial Tokenization for LLM-Based Health Intelligence: A Federated Architecture for Industrial . . .
Semiotic Logical Hexagon Theory for LLM Logical Reasoning
SceneActBench: Can Agents Act on the 3D Scenes They See?
Understanding Human-like-Solutions in Combinatorial Optimization via Learning and Search
Artificial Intelligence and Innovation Ecosystem: Evolutionary Developments, Challenges, . . .
Are Prompt Optimizers Blind? Cross-Modal Visual Feedback for Automatic Prompt Optimization
Efficiency Matters in Autonomous Research
Towards High-Level Semantic Intelligence
Joint Text-Audio Alignment for EEG-to-Text Decoding in Chinese Speech Production and Perception
Cognivia: A cognitive Behavioral Therapy Copilot for Evidence-Based Mental Healthcare
Distilling Temporal Search and Reasoning: Evolving LLMs for Future Prediction via Harness-Assisted . . .
Are the High-weight Neurons the Important Ones in Image Classification Neural Networks?
Localized Adaptation Reveals Distinct Learning Signatures in Transformers
LLM-Assisted Ontology Engineering and Construction of a French Legal Knowledge Graph
Toward an Organizational Science of Multi-Agent LLM Systems: Decoupling Who, How, and . . .
From Execution to Capability: Scientific Experience Consolidation via Procedural Knowledge Synthesis
Falling Behind Drives Unsafe Development in an Idealised AI Race Experiment
Distributing Security Controls Through Harness Engineering
Large Language Model for Operations Research Formulation Selection in Multi-Warehouse Inventory . . .
CHARM: A Multimodal Graph Foundation Model with Hierarchical Context Modeling for Zero-Shot . . .
What Does It Take to Detect an AI Agent? Minimal Feature Sets for Behavioral Detection . . .
Can AI Agents Conduct Open-Ended AI Research?
SciDataSailor: Deep Scientific Data Exploring
AgentMap: Joint Equivalence and Subsumption Discovery for Ontology Matching
Exploring Structures in Physics Problems: Can AI Agents Discover Statistical Mechanical Mappings?
Agent Skills Matter: Inferring Proprietary Skills from Execution Trajectories
Group-Reflective Self-Distillation for Agentic Reinforcement Learning
Towards Autonomous Aircraft Surveillance from Nanosatellites through On-Board Inference . . .
DataClawEval: A Benchmark for Data Engineering Agents in Real Industrial Harness
MemHarness: Memory is Reconstructed, Not Replayed
Qwen-UI-Agent Technical Report: Toward Next-Generation Real-World Centric Foundation GUI Agent
What AI Red-Team Evaluations Can and Cannot Prove

Figure 26. Titles of Sixty-Three Recent AI Research Reports on arXiv.

The New Frontier: Mapping Trends in AI Research

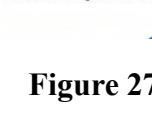
A compilation of recent high-impact AI research papers reveals a transition from passive Large Language Models toward “Agentic AI”—systems capable of autonomous reasoning, workflow synthesis, and self-improvement—alongside a growing emphasis on safety certifications and industrial applications.

THE RISE OF AGENTIC AUTONOMY

Recursively Self-Improving Deep Research
New models like ARES focus on agents that autonomously explore scientific data and self-improve.



Transition to Agentic Workflows and ERP
Research is moving toward multi-agent architectures for autonomous Enterprise Resource Planning and workflow synthesis.



Real-World Centric GUI Agents
Development of next generation foundation agents designed to interact directly with graphical user interfaces.

Evolving intelligence

TRUST, SAFETY, AND GOVERNANCE



BENCHMARKING AUTONOMOUS AGENTS



GuardianAgentBench
Identifying agent failure points and guardrail effectiveness.



DocOps
Verifiable performance of agents in complex document operations.



SceneActBench
Testing agent capability to act on perceived 3D scenes.



Independent Certification for Trustworthy AI
A rising focus on closing the trust gap through external, independent verification of AI systems.



Risk-Calibrated Governance Frameworks
New frameworks like TRUST-EBD provide governance-aware structures for evaluating and securing autonomous systems.



Red-Team Evaluations and Behavioral Detection
Research highlights the limitations of red-teaming and explores minimal feature sets for detecting AI agents.

NotebookLM

Figure 27. Google Gemini Notebook Infographic: arXiv Research Report Titles. ❁

The details of the Infographic are as follows (content extracted from Figure 27):

❁ Title: *The New Frontier: Mapping Trends in AI Research*

Note: A compilation of recent high-impact AI research papers reveals a transition from Large Language Models toward “Agentic AI”—systems capable of autonomous reasoning, workflow synthesis, and self-improvement—alongside a growing emphasis on safety certifications and industrial applications.

Subtitle #1: The Rise of Agentic Autonomy

- Recursively Self-Improving Deep Research
- Real-World Centric GUI Agents
- Transition to Agentic Workflows and ERP

Additional Theme: Evolving Intelligence

Subtitle #2: Trust, Safety, and Governance

- Benchmarking Autonomous Agents
- Independent Certification for Trustworthy AI
- Risk-Calibrated Governance Frameworks
- Red-Team Evaluations and Behavioral Detection ❁

Numerous new AI-oriented entities have been formed. Some are referred to as “*neolabs*” which Microsoft Copilot defined as the following (8/7/26): ☼ “Neolabs are a new wave of AI research labs and startups focused on experimental architectures and foundational breakthroughs, operating like research institutions rather than traditional product-driven companies.” ☼ Seven new AI-oriented entities are shown in Figure 28 which is followed by detailed information on each.

	HQ	Key Player	Emphasis
Advanced Machine Intelligence Labs	Paris	Yann LeCun	Real World Intelligent Systems
Discovery Loop	Palo Alto	Jeff Dean	Automating Scientific Discovery
Ineffable AI	Altrincham	David Silver	Superintelligence
Prometheus	San Francisco	Jeff Bezos	Physical AI
Safe Superintelligence	Palo Alto	Ilya Sutskever	Safe Superintelligence
Thinking Machines Lab	San Francisco	Mira Murati	Customizable AI Systems
World Labs	San Francisco	Fei-Fei Li	Spatial Intelligence

Figure 28. Seven Relatively New AI-Oriented Entities.

Advanced Machine Intelligence (AMI) Labs (Key Player: Yann LeCun): www.amilabs.xyz

From the AMI Labs Website:

“March 10, 2026 - Updates: Official Launch

Advanced Machine Intelligence (AMI) is building a new breed of AI systems that understand the world, have persistent memory, can reason and plan, and are controllable and safe.

We are enabling the next AI revolution, by building a new breed of AI systems that (1) understand the real world, (2) have persistent memory, (3) can reason and plan, (4) are controllable and safe.

Our main goal is to build intelligent systems that understand the real world.

Real-world data is continuous, high-dimensional, and noisy, whether it is obtained through cameras or any other sensor modality.

Generative architectures trained by self-supervised learning to predict the future have been astonishingly successful for language understanding and generation. But much of real-world sensor data is unpredictable, and generative approaches do not work well.

AMI is developing world models that learn abstract representations of real-world sensor data, ignoring unpredictable details, and that make predictions in representation space.

Action-conditioned world models allow agentic systems to predict the consequences of their actions, and to plan action sequences to accomplish a task, subject to safety guardrails.

AMI will advance AI research and develop applications where reliability, controllability, and safety really matter, especially for industrial process control, automation, wearable devices, robotics, healthcare, and beyond.

Founded by a globally distinguished team of scientists, engineers, and builders, AMI is a frontier AI research lab operating across three continents from day one.

We share one belief: real intelligence does not start in language. It starts in the world.”

Discovery Loop (Key Player: Jeff Dean): www.discoveryloop.com

Co-Founders: Jeff Dean (CEO), Sanjay Ghemawat, Quoc Le, & Oriol Vinyals

From the Discovery Loop Website:

“Our mission is straightforward: we are building AI solutions that can automatically solve important problems in machine learning, science, and engineering. By advancing the pace at which we conduct engineering and scientific discovery, we can bring the benefits of science and technology to the world much faster. Ultimately, our goal is to build AI systems that act as a deeply positive, empowering force for humanity, delivering technology solutions that improve people's lives on a global scale.

Automating discovery to accelerate science and engineering for the world.

Collectively, we represent three of the most-cited researchers in artificial intelligence and two of the most-cited researchers in distributed systems.

Between us, we have pioneered massive scale computing and led the creation of critical infrastructure, products, and foundational AI advances that the world relies on, including multiple generations of Google Search, Google Ads, Google News, Google Translate, Google File System, MapReduce, BigTable, Spanner, TensorFlow, Pathways, TPUs, AlphaChip, AlphaStar, AlphaCode, AlphaFold, Gemini, model distillation, mixture-of-experts model architectures, word2vec, sequence-to-sequence models, chain of thought reasoning, neural architecture search, and multiple generations of Large Language Models (LLMs) among others.

Our relative advantage isn't just our technical ability; it is the unprecedented scale of the systems we have previously built. We possess true full-stack depth that spans chips, hardware infrastructure, software infrastructure, ML models, and products.”

Ineffable Intelligence (Key Player: David Silver): www.ineffable.ai

From the Ineffable Intelligence Website:

“We are creating a superlearner that discovers all knowledge from its own experience, from elementary motor skills through to profound intellectual breakthroughs.

This superlearning capability - the ability to endlessly discover knowledge and skills, without relying on human data - will be driven by the world’s most powerful reinforcement learning algorithms.

The superlearner is expected to rediscover and then transcend the greatest inventions in human history, such as language, science, mathematics and technology.

If successful, this will represent a scientific breakthrough of comparable magnitude to Darwin: where his law explained all Life, our law will explain and build all Intelligence.

At Ineffable, we believe the following:

Transformative - The arrival of superintelligence will be the most transformative moment in history.

Beneficial - Superintelligence can and must be built in a way that is beneficial to humanity.

Timely - Superintelligence can be built within years, not decades or centuries.

Experiential - Superintelligence will be achieved through learning from experience, rather than learning from human data.

Ineffable AI - The knowledge and skills acquired by superintelligence will be too profound to be described by human language.

General - A single superlearner will be applicable to any digital or physical machine; any stream of experience.

Continual - Superintelligence will keep learning and improving, forever.”

Prometheus (Key Player: Jeff Bezos): www.prometheus.ai

There is scant information available on Prometheus. The following information was generated by ChatGPT (OpenAI) based on the author’s prompt:

Author’s ChatGPT (OpenAI) Prompt > “What is a half-page summary of the new AI organization co-founded by Jeff Bezos named Prometheus?”

🌀 Half-Page Summary: Prometheus

“Prometheus is a highly funded, secretive artificial-intelligence startup co-founded and co-led by Amazon founder Jeff Bezos and technology executive Vik Bajaj. Founded in late 2025, the company is pursuing a form of “physical AI” aimed at transforming how physical products are engineered and manufactured, rather than building another general-purpose chatbot. (Ars Technica)

Prometheus’s central ambition is to create what Bezos calls an “artificial general engineer”—AI that can assist with or automate sophisticated engineering tasks, including designing, simulating, testing, and manufacturing complex physical systems. Potential applications include jet engines, semiconductors, medical devices, consumer electronics, and other advanced industrial products. Bezos has compared the concept to a highly advanced form of computer-aided design, although he emphasizes that the company is not simply developing robotics. (GeekWire)

The idea is to dramatically shorten engineering and manufacturing cycles. A project that might traditionally require hundreds of engineers and many years of experimentation could, in Prometheus’s vision, be completed with a much smaller team in a fraction of the time. This makes the company’s broader goal less about replacing individual workers with robots and more about using AI to amplify human engineering capabilities and accelerate industrial innovation. (Axios)

Prometheus has attracted extraordinary financial backing. After initially raising about \$6.2 billion, it announced another \$12 billion funding round in June 2026, giving it an approximately \$41 billion valuation. Investors include Bezos and major financial institutions such as JPMorgan Chase, Goldman Sachs, and

BlackRock. The scale of the funding reflects the enormous computing, scientific, and engineering resources required for its ambitions. (Axios)

In short, Prometheus is Bezos's bet that AI's next major frontier will be the physical world: not merely generating text, images, or software, but helping humans invent and manufacture increasingly complex things faster and more efficiently." ❁

Safe Superintelligence (Key Player: Ilya Sutskever): www.ssi.inc

From the Safe Superintelligence Website:

"Superintelligence is within reach.

Building safe superintelligence (SSI) is the most important technical problem of our time.

We have started the world's first straight-shot SSI lab, with one goal and one product: a safe superintelligence.

It's called Safe Superintelligence Inc.

SSI is our mission, our name, and our entire product roadmap, because it is our sole focus. Our team, investors, and business model are all aligned to achieve SSI.

We approach safety and capabilities in tandem, as technical problems to be solved through revolutionary engineering and scientific breakthroughs. We plan to advance capabilities as fast as possible while making sure our safety always remains ahead.

This way, we can scale in peace.

Our singular focus means no distraction by management overhead or product cycles, and our business model means safety, security, and progress are all insulated from short-term commercial pressures.

We are an American company with offices in Palo Alto and Tel Aviv, where we have deep roots and the ability to recruit top technical talent.

We are assembling a lean, cracked team of the world's best engineers and researchers dedicated to focusing on SSI and nothing else."

Thinking Machines Lab (Key Player: Mira Murati): www.thinkingmachines.ai

From the Thinking Machines Lab Website:

"Thinking Machines Lab is an artificial intelligence research and product company. We're building a future where everyone has access to the knowledge and tools to make AI work for their unique needs and goals.

While AI capabilities have advanced dramatically, key gaps remain. The scientific community's understanding of frontier AI systems lags behind rapidly advancing capabilities. Knowledge of how these systems are trained is concentrated within the top research labs, limiting both the public discourse on AI and people's abilities to use AI effectively. And, despite their potential, these systems remain difficult for people to customize to their specific needs and values. To bridge the gaps, we're building Thinking Machines Lab to make AI systems more widely understood, customizable and generally capable.

We are scientists, engineers, and builders who've created some of the most widely used AI products, including ChatGPT and Character.ai, open-weights models like Mistral, as well as popular open-source projects like PyTorch, OpenAI Gym, Fairseq, and Segment Anything.

Science is better when shared

Scientific progress is a collective effort.

Emphasis on human-AI collaboration.

More flexible, adaptable, and personalized AI systems.

Model intelligence as the cornerstone.

Infrastructure quality as a top priority.

Advanced multimodal capabilities.

Research and product co-design.

Empirical and iterative approach to AI safety.

Measure what truly matters."

World Labs (Key Player: Fei-Fei Li): www.worldlabs.ai

From the World Labs Website:

"World Labs is a leading spatial intelligence company, building frontier world models that can perceive, generate, reason, and interact with the 3D world.

World Labs is a frontier AI research and product company.

We build foundational world models that can perceive, generate, reason, and interact with the 3D world — unlocking AI's full potential through spatial intelligence by transforming seeing into doing, perceiving into reasoning, and imagining into creating. We believe spatial intelligence will unlock new forms of storytelling, creativity, design, simulation, and immersive experiences across both virtual and physical worlds.

World Labs was founded by visionary AI pioneer Fei-Fei Li along with Justin Johnson, Christoph Lassner, and Ben Mildenhall, each a world-renowned technologist in machine learning, generative AI, computer vision, and graphics. We bring together a world-class team, united by a shared curiosity, passion, and deep backgrounds in technology — from AI research to systems engineering to product design — creating a tight feedback loop between our cutting-edge research and products that empower our users.

Our first product, Marble, is powered by best-in-class generative 3D world models and enables anyone to create spatially cohesive, high-fidelity, and persistent 3D worlds from just a single image, video or text prompt."

Seven General AI Deployment Recommendations

Leaders of organizations are responsible for determining the extent to which AI will be used and how it will be deployed. This is especially difficult since AI-related technologies are advancing so rapidly and the risks are considerable. Seven general AI deployment recommendations emerged from a case study analysis of a broad array of AI deployments. The recommendations are depicted in Figure 29. The seven recommendations will not guarantee organizational success with AI because each organization is unique and there are numerous other important factors not mentioned.

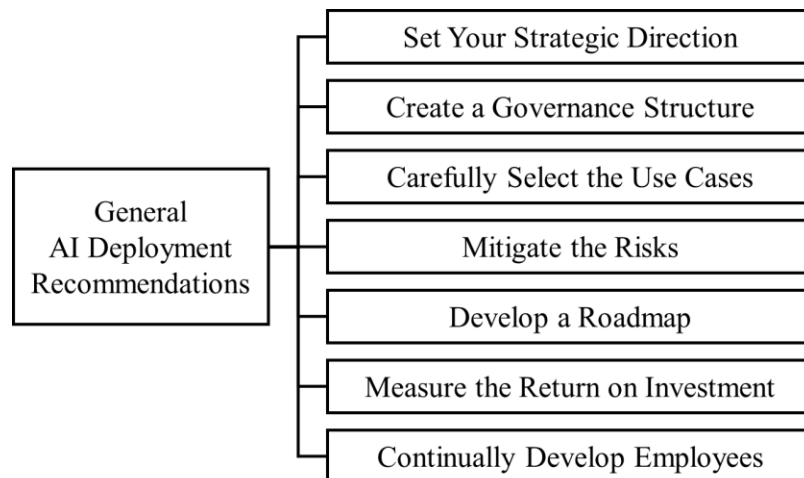


Figure 29. Seven General AI Deployment Recommendations.

1. Set Your Strategic Direction

Driving Questions: Why are we here? What do we believe? Where are we going?

Let's take a *detour* to the book, "*Alice's Adventures in Wonderland*" by Lewis Carroll (1865):

The central character of the book is **Alice** who is a young girl. She encounters a **Cheshire Cat** sitting on a bough of a tree a few yards off while walking on a path.

Alice: "Would you tell me, please, which way I ought to walk from here?"

"That depends a good deal on where you want to get to," said the **Cat**.

"I don't much care where—" said **Alice**.

"Then it doesn't matter which way you walk," said the **Cat**.

"—so long as I get *somewhere*," **Alice** added as an explanation.

"Oh, you're sure to do that," said the **Cat**, "if you only walk long enough."

It is good business practice to have clarity and consensus on the foundational items of your organization when making AI-related decisions. Here are some guiding questions:

- What is our mission?
- What is our philosophy?
- What are our values?
- What is our vision?
- What are our organizational performance metrics?
- What are our core processes (workflows)?
- What are our strengths and weaknesses?
- What are our opportunities and threats?
- What are the major issues facing our customers?
- What are our priorities?
- What are we trying to accomplish?

2. Create a Governance Structure

Driving Question: How will we govern the presence and use of AI in our organization?

Uncontrolled haphazard experimentations with AI are not typically successful and they aren't sustainable. An AI Governance Structure can help. Such a structure can assure there is clarity and consensus on AI accountabilities, roles, and responsibilities. Here are some guiding questions:

- Who is accountable for overseeing our AI activities?
- Who in the organization has deep AI expertise?
- Who will serve on our AI Governance Council?
- What will be the roles and responsibilities of the members of our AI Governance Council?
- What will be the objectives of our AI Governance Council?
- How will the AI Governance Council assure the safe, ethical, and responsible use of AI?
- What are the AI laws and regulations that apply to us?
- Are we in compliance with those laws and regulations?
- Do we expect any new laws or regulations in the near future?
- What are our organization's AI policies and procedures?
- Do we have an AI Risk Management System?
- Do each of our suppliers and partners have an AI Risk Management System?

3. Carefully Select the Use Cases

Driving Question: What are our specific planned AI use cases?

Every organization has numerous AI use case opportunities. The value realized from AI and the AI risks will be largely determined by the selected use cases. Here are some guiding questions:

- What are the possible AI use cases for our organization?
- Which use cases are the most promising?
- What AI-related technologies will we need for the use cases?
- Will we use external AI models and/or develop our own models?
- If we develop our own models, then do we have the necessary in-house experts?
- Will the models we use be open-source models or closed-source models?
- What are the risks associated with the use cases we are considering?
- Which performance metrics will be impacted by the use cases we are considering?
- What is the expected return on investment in the next year for the use cases we are considering?
- Will our employees benefit from the use cases we are considering?
- Will the use cases augment our employees and their work or replace our employees?
- What AI-related technologies are already embedded in the products and services that we use?
- Will our data infrastructure support the use cases we are considering?
- Will our supply chain support the use cases we are considering?

4. Mitigate the Risks

Driving Question: How will we mitigate the risks associated with our AI use cases?

The use of AI involves risks such as model output errors and hallucinations, decision errors, security issues, and employee discontent to name a few. Some of the risks are predictable, but there will probably be new emergent risks that weren't anticipated. Fortunately, many of the risks associated with AI are now known and understood because of the number of AI deployment case studies in the literature and the high level of scrutiny. Here are some guiding questions:

- Who will design and manage our AI Risk Management System?
- How will we identify, analyze, and prioritize our AI risks?
- What is our Risk Mitigation Process?
- What are the major AI risks facing our organization?
- What are the major AI risks facing our customers, suppliers, and partners?
- What are our risks associated with current AI laws and regulations?
- Do our suppliers and partners have an AI Risk Management System?

5. Develop a Roadmap

Driving Question: What is our AI Roadmap for the next year?

The deployment of AI should be aligned with your foundational items like your mission, philosophy, values, and vision (see Item #1). Unless you are lucky, you won't make much progress without a Roadmap. Here are some guiding questions:

- How can we use AI to help us with our foundational (Strategic Direction) items (Item #1)?
- What is our AI Roadmap for the next year?
- What are the most important AI milestones and activities for each month?
- Do we have detailed AI action plans?
- Who will review our AI Roadmap progress each month?
- How will we measure AI Roadmap progress?
- How will we review our AI Roadmap progress?

6. Measure the Return on Investment

Driving Question: Have our investments of time and money in AI been worth it?

Many people today are questioning whether AI is worth it. The *Return on AI Investment* should be measured even though precise numbers might be impossible to obtain. Here are some questions:

- Who will calculate our Return on AI Investment?
- Do we have a way to track the time and money spent on AI without creating a bureaucracy?
- What procedures will be used to calculate our Return on AI Investment?
- What are the major tangible and intangible benefits and cost associated with our AI use cases?
- Are metrics related to Customer Satisfaction and Employee Satisfaction improving due to AI?

7. Continually Develop Employees

Driving Question: Are our employees continuing to learn, grow, and develop?

You will probably be more successful with AI if your employees are knowledgeable and skilled in your core business and AI. There is a risk of employees becoming too reliant on AI and failing to keep learning, growing, and developing. It is sometimes necessary for humans to acquire basic knowledge and skills related to a profession—which in some cases can be replaced by AI—in order to develop the more advanced knowledge and skills. Ideally, AI should augment employees and their work and help them develop knowledge and skills. Here are some guiding questions:

- What core business knowledge and skills do employees need to develop?
- What AI knowledge and skills do employees need to develop?
- Does each employee have a development plan?
- Which of our employees need advanced training on AI?
- What specific AI training needs to be conducted in our organization?
- Can AI be used to train and develop our employees?
- How will we evaluate employee development?
- Do we have an Employee Development Roadmap?
- Do we need to hire new AI talent?

IV. AI's Future: Three Scenarios

“It's difficult to make predictions, especially about the future.”

Danish Folklore / Robert Storm Petersen

Many have said that even the best AI experts in the world can't predict AI's future. There are an infinite number of possibilities for the future of AI and speculation is now commonplace. The following are fictitious quotes on *“The Future of AI”* based on a compilation of actual quotes:

- “AI in the future will eliminate all diseases.”
- “We won't have to work in the future because of AI.”
- “Everyone will have an AI agent and a robot in the future.”
- “We won't have steering wheels in vehicles in the future because of AI.”
- “Scarcity resources won't exist in the future—we will live in an era of abundance.”

Kurzweil (2005) introduced a profound concept to describe a future transformed by technology:

“What, then, is the Singularity? It's a future period during which the pace of technological change will be so rapid, its impact so deep, that human life will be irreversibly transformed. Although neither utopian nor dystopian, this epoch will transform the concepts that we rely on to give meaning to our lives, from our business models to the cycle of human life, including death itself.”

For some, it might feel like we have already reached the *Singularity*.

Kissinger *et al.* (2024) predicted some of the human needs and actions because of AI:

“The advent of artificial intelligence is, in our view, a question of human survival. As we explain later in this book, AI’s future faculties, operating at human speeds, will render traditional regulation useless. We will need a fundamentally new form of control. For the global scientific community, the immediate task is to find the technical measures for instilling intrinsic safeguards in every AI system. For their part, nations and international organizations, once they have coalesced around a consensus, must develop new political structures for monitoring, enforcement, and crisis response. That will require the resolution of not one but two ‘alignment problems’: the technical alignment of human values and intentions with the actions of AI, and the diplomatic alignment of humans with their fellow humans.”

Lee and Qiufan (2021) offered an interesting look into the future of AI in the book titled, “*AI: Ten Visions for our Future*.” Each of the ten chapters starts with a futuristic short story followed by an analysis of the story from the perspective of the embedded AI concepts and techniques:

Chapter 1. The Golden Elephant

Analysis: Deep Learning, Big Data, Internet/Finance Applications, AI Externalities

Chapter 2. Gods Behind the Masks

Analysis: Computer Vision, Convolutional Neural Networks, Deepfakes, Generative Adversarial Networks (GANs), Biometrics, AI Security

Chapter 3. Twin Sparrows

Analysis: Natural Language Processing, Self-Supervised Training, GPT-3, AGI and Consciousness, AI Education

Chapter 4. Contactless Love

Analysis: AI Healthcare, AlphaFold, Robotic Applications, COVID Automation Acceleration

Chapter 5. My Haunting Idol

Analysis: Virtual Reality (VR), Augmented Reality (AR), and Mixed Reality (MR), Brain-Computer Interface (BCI), Ethical and Societal Issues

Chapter 6. The Holy Driver

Analysis: Autonomous Vehicles, Full Autonomy and Smart Cities, Ethical and Societal Issues

Chapter 7. Quantum Genocide

Analysis: Quantum Computers, Bitcoin Security, Autonomous Weapons and Existential Threat

Chapter 8. The Job Savior

Analysis: AI Job Displacement, Universal Basic Income (UBI), What AI Cannot Do, 3RS as Solution to Displacement

Chapter 9. Isle of Happiness

Analysis: AI and Happiness, General Data Protection Regulation (GDPR), Privacy Computing Using Federated Learning and Trusted Execution Environment (TEE)

Chapter 10. Dreaming of Plentitude

Analysis: Plentitude, New Economic Models, The Future of Money, Singularity

Bostrom (2014) introduced the concept of *superintelligence* and suggested how it could affect humanity: “We can tentatively define a superintelligence as *any intellect that greatly exceeds the cognitive performance of humans in virtually all domains of interest.*” There could be three forms of *superintelligence* according to Bostrom (2014):

“**Speed superintelligence:** *A system that can do all that a human intellect can do, but much faster.*”

“**Collective superintelligence:** *A system composed of a large number of smaller intellects such that the system’s overall performance across many very general domains vastly outstrips that of any current cognitive system.*”

“**Quality superintelligence:** *A system that is at least as fast as a human mind and vastly qualitatively smarter.*”

Short-term AI predictions (conjectures, speculations, etc.) are much easier than long-term predictions. Here are some **AI predictions for the rest of 2026:**

- There will continue to be a lot of attention paid to AI.
- There will continue to be races between countries, frontier AI labs, universities, and corporations.
- The U.S. and China will continue to lead in AI.
- All levels of the AI Stack will continue to be important.
- There will continue to be safety motives and profit motives which will at times conflict.
- There will continue to be AI proponents and detractors.
- There will continue to be a divide between those who master AI and those who don’t.
- There will continue to be new AI models released.
- There will continue to be new AI-oriented entities that are formed.
- There will be increased interest in Agentic AI (use of AI Agents) and Physical (embodied) AI.
- The AI models will continue to improve on the performance benchmarks.
- Some careers and jobs will be eliminated and new careers and jobs will be created.
- There will be more mature laws and regulations.
- More organizations will create their own small language models on equipment they control.
- The people with extensive AI knowledge and skills will continue to be in high demand.
- There will be new dangers and risks that emerge.
- There will continue to be questioning of the Return on AI Investment.
- AI will become imbedded in more systems, products, and services.
- There will be new AI use cases in homes, workplaces, schools, and hospitals.

Since the future of AI is uncertain and the possibilities are infinite, it might be useful to imagine and explore multiple scenarios for *conjecturing* about the future of AI. To carry out this *speculative* exercise, let’s use the same common anchor across the scenarios. **AI Presence** is an abstract concept, but it will be used here to mean *the extent to which AI is present in society* whether it be in homes, schools, hospitals, home offices, factories, corporations, equipment, devices, vehicles, etc. AI already has a significant *presence* in society. What happens now? What follows is a discussion of three scenarios for the future of **AI Presence**: Plateau, Incremental, and Acceleration. The three conceptual “*AI Presence*” scenarios are depicted in Figure 30 and then discussed.

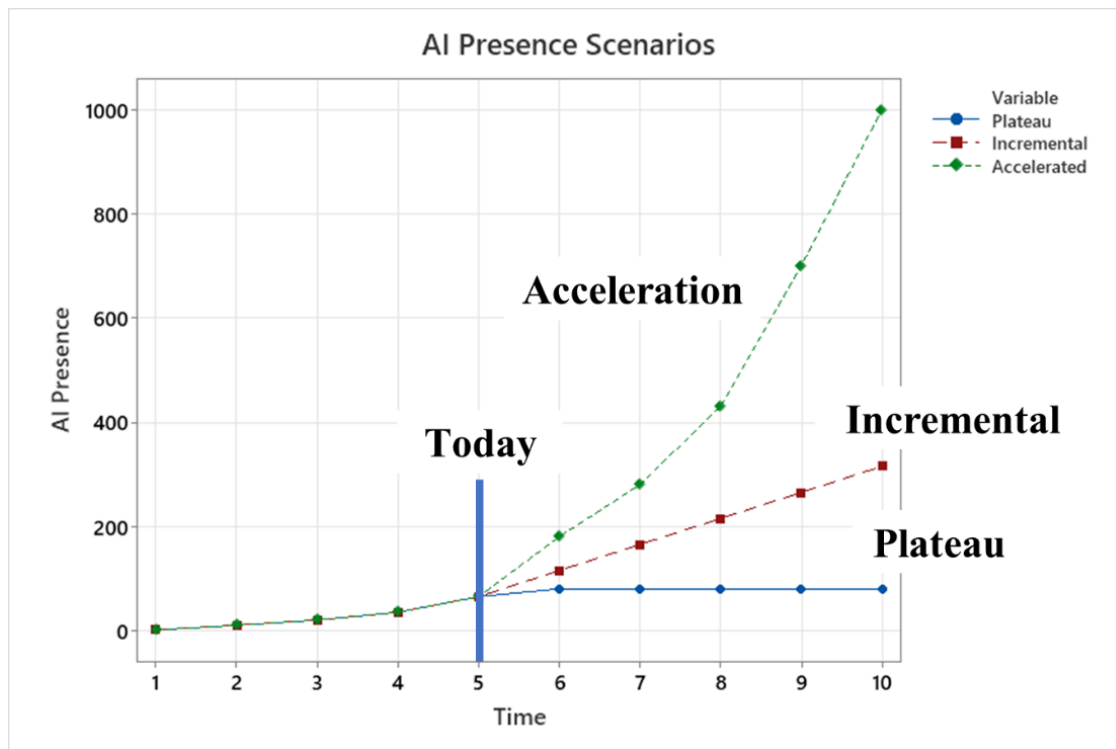


Figure 30. Three Future AI Presence Growth Scenarios.

Note: *AI Presence* means here *the extent to which AI is present in society*.

Plateau Scenario

Imagine if the *AI Presence* that we have today is what we will have in the future, i.e., *AI Presence* plateaus or levels-off.

Hypothetical Causes (Partial List):

- *AI Presence* plateaus because the Returns on AI Investment were not high enough.
- *AI Presence* plateaus because the risks associated with AI became too great.
- *AI Presence* plateaus because laws and regulations stifled innovation.
- *AI presence* plateaus because the number of AI-related innovations dramatically decreased.

Conjectured Consequences: This would be the least disruptive of the three scenarios. If this scenario were realized, then it might be a long time before AI receives as much attention, interest, and investment as it receives now. AI risks might be manageable. There would probably be a lot of effort on exploiting existing AI capabilities. There would still be demand for people with AI knowledge and skills and there would probably be a lot of re-grouping, re-thinking, re-organizing, and re-planning. In terms of *human emotions*:

- If you are currently an **AI proponent**, then you would probably be disappointed.
- If you are currently **neutral on AI**, then you might use AI only as necessary.
- If you are currently **anti-AI**, then you would probably be happy.

We might experience another *AI Winter*. Multiple AI Winters have occurred in the past. According to Wooldridge (2020): “The decade from the early 1970s to the early 1980s later became known as the **AI winter**, although it should perhaps better be known as the *first* AI winter, because there were more to come. The stereotype was established, of AI researchers making hopelessly overoptimistic and unwarranted predictions, then failing to deliver. Within the scientific community, AI began to acquire a reputation somewhat akin to homeopathic medicine. As a serious academic discipline, it appeared to be in terminal decline.”

A second AI winter occurred later (Wooldridge, 2020): “By the late 1980s, the boom days of expert systems were over, and another AI crisis was looming. Once again, the AI community was criticized for overselling ideas, promising too much, and delivering too little.” Google Gemini (July 20, 2026) dates **The First AI Winter** as occurring between 1974-1980 and **The Second AI Winter** occurring between 1987-1993.

Incremental Scenario

Imagine if the *AI Presence* increases incrementally in the future, i.e., *AI Presence* increases at a relatively constant rate.

Hypothetical Causes (Partial List):

- *AI presence* increases at a relatively constant rate because AI is proved to be beneficial.
- *AI presence* increases at a relatively constant rate because the AI risks were tolerable.
- *AI presence* increases at a relatively constant rate because laws and regulations were adequate.
- *AI presence* increases at a relatively constant rate because AI-related innovations continued.

Conjectured Consequences: There could be some economic and societal disruptions. Attention, interest, and investment would continue. Managing AI risks could continue to be challenging. There would probably be the exploitation of existing AI capabilities and exploration. There might be great demand for people with AI knowledge and skills. There would continue to be AI deployment planning. In terms of *human emotions*:

- If you are currently an **AI proponent**, then you would probably continue to be excited.
- If you are currently **neutral on AI**, then you would probably engage more with AI.
- If you are currently **anti-AI**, then you would probably be disappointed.

Acceleration Scenario

Imagine if *AI Presence* increases at an accelerating (increasing) rate in the future.

Hypothetical Causes (Partial List):

- *AI Presence* increases at an increasing rate because the Returns on AI Investment were great.
- *AI Presence* increases at an increasing rate because the AI-related risks were manageable.
- *AI Presence* increases at an increasing rate because AI spread to every sector and was scaled.
- *AI Presence* increases at an increasing rate because AI became able to recursively self-improve.

Conjectured Consequences: There could be profound and significant economic and societal disruptions. Attention, interest, and investment would probably be high. There could be unimaginable benefits along with existential and extremely challenging risks. There would probably be the exploitation of existing AI capabilities and a lot of exploration. There could be incredibly high demand for people with AI knowledge and skills. AI deployment planning would be necessary.

In terms of *human emotions*:

- If you are currently an **AI proponent**, then you would probably be elated.
- If you are currently **neutral on AI**, then you would probably rapidly engage more with AI.
- If you are currently **anti-AI**, then you would probably be depressed.

We don't know what will happen with AI in the future. We can only (1) hope that all forms of AI will be controllable, benevolent, and aligned with human interests and (2) be prepared.

“Any sufficiently advanced technology is indistinguishable from magic.”

Arthur C. Clarke

Humans must unravel the mysteries of AI and shine more light into the *black boxes*. We will also need to continue to learn, grow, and develop for us to comprehend the *great mysteries* of life.

V. Conclusion

It is sometimes useful to pause, take a break from your busy life, and think deeply about what is happening around you. AI involves creating *intelligent machines* that learn from data, solve problems, and perform tasks. The rapid and profound advances in AI capabilities have caused humans to think deeply and question who we are, what we want to become, and what role *intelligent machines* will play in our lives.

The participants of the *1956 Dartmouth Summer Research Project on Artificial Intelligence* met seventy years ago over a period of 8-10 weeks to share ideas, discuss and debate concepts and techniques, and brainstorm ideas on the new label of *Artificial Intelligence*. In other words, **they paused**. They had the good fortune of having access to the knowledge of their predecessors—knowledge that was created centuries, decades, years, and months prior to the summer of 1956.

Modern-day AI is present in almost all segments of society whether it be in our governments, schools, hospitals, corporations, and other organizations. Benefits have been realized and significant risks have been identified. There are competitive races between countries, frontier AI labs, universities, and corporations. The demand for AI talent has dramatically increased. Many are justifiably questioning the value of AI and how much it actually benefits humans.

The future of AI is uncertain. Predictions range from dystopia-to-utopia and scarcity-to-abundance. Three scenarios regarding the future of *AI Presence* were discussed: Plateau Scenario, Incremental Scenario, and Acceleration Scenario. Regardless of which one—if any—becomes reality, it's probably wise for humans to learn and safely experiment with AI and for leaders of organizations to deploy AI with caution. Preparing for a Singularity now seems reasonably wise.

It has been said that *truth is sometimes stranger than fiction*. The Three Laws of Robotics appeared in the short story—*Runaround*—by Isaac Asimov (1942, 1950): [excerpts] “**The First Law:** A robot may not injure a human being, or, through inaction, allow a human being to come to harm; **The Second Law:** A robot must obey the orders given it by human beings except where such orders would conflict with the First Law; and **The Third Law:** A robot must protect its own existence as long as such protection does not conflict with the First or Second Laws.” It seems surreal that these three laws are being seriously discussed today. AI has progressed so much so fast that it's *hard to believe* at times what AI can do. We humans should strive to continue to learn, grow, and develop and not let AI do everything for us. Our knowledge and skills can deteriorate quickly if we don't put in the effort to maintain them and possibly go beyond them. It's our ability to *think* which makes us unique. We should all pause, think deeply, pose existential questions, and answer them the best we are able. I'll continue to be optimistic about the future of AI and hope for the best possible future for humanity and the *intelligent machines*.

Acknowledgements

This paper is dedicated to John McCarthy, Claude Shannon, Marvin Minsky, and Nathaniel Rochester who submitted the research project proposal in 1955 for the *Dartmouth Summer Research Project on Artificial Intelligence* that was eventually held during the summer of 1956 at Dartmouth College in Hanover, NH. Their conversations set in motion a seventy-year AI odyssey.

Author's Google Gemini Prompt > “Create a black and white image set in front of the Math Building on Dartmouth College in 1956 containing one group portrait of John McCarthy, Marvin Minsky, Claude Shannon, and Nathaniel Rochester - the four people who took part in the famous 1956 Dartmouth AI Research Project.” The Google Gemini image is shown in Figure 31.

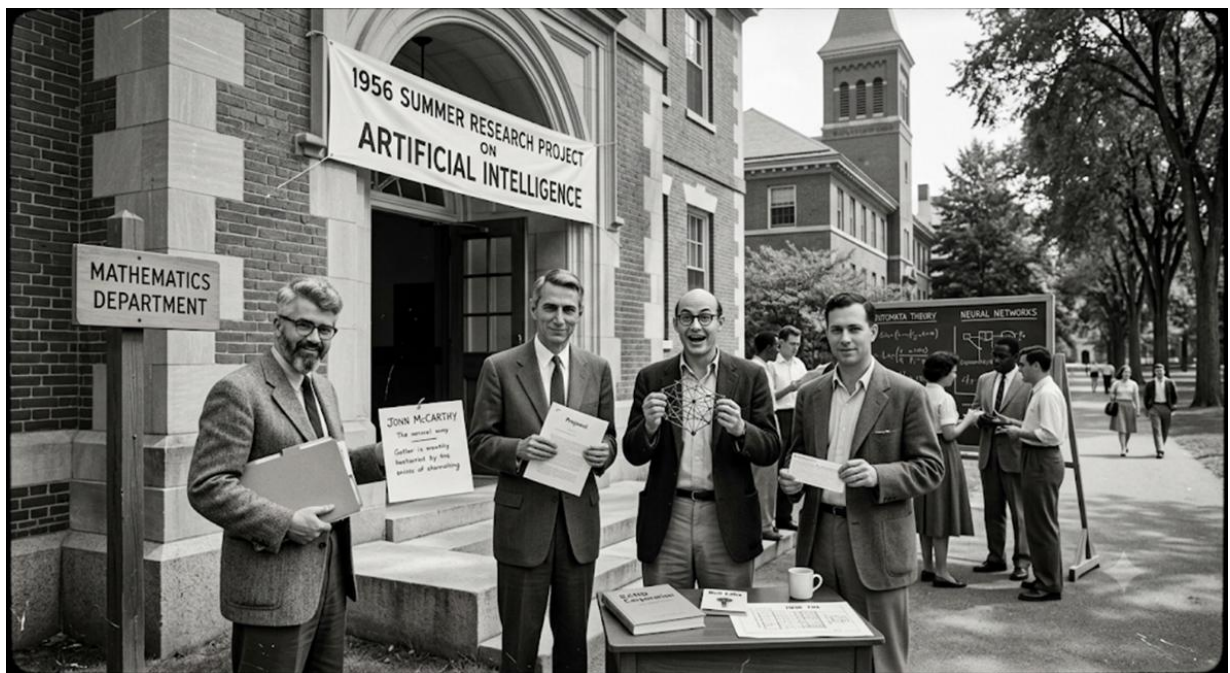


Figure 31. Image Created by Google Gemini on August 4, 2026. ☞

Left-To-Right

John McCarthy, Dartmouth College, Assistant Professor of Mathematics

Claude E. Shannon, Bell Telephone Laboratories, Mathematician

Marvin L. Minsky, Harvard University, Harvard Junior Fellow in Mathematics and Neurology

Nathaniel Rochester, I.B.M. Corporation, Manager of Information Research

References: Alphabetical

- Advanced Machine Intelligence (AMI) Labs Website: www.amilabs.xyz
Alibaba Website: www.alibaba.com
Amazon Website: www.amazon.com
Anthropic Website: www.anthropic.com
Arena AI: www.arena.ai
arXiv Website: arXiv.org
Apple Website: www.apple.com
Asimov, I., (1942), *I, Robot*, In *Runaround*, Spectra Books, New York, NY.
Asimov, I., (1950), *I, ROBOT*, Gnome Press, New York, NY.
Baidu Website: www.baidu.com
Banks, I., (1987), *Consider Phlebas*, Macmillan, London, UK.
Bayes, T., (1764), *An Essay Towards Solving a Problem in the Doctrine of Chances*, Philosophical Transactions of the Royal Society of London for 1763, 53: 370-418.
Bengio, Y., LeCun, Y., & Hinton, G., (2021), *Deep Learning for AI*, Turing Lecture, Communication of the ACM, July 2021, Volume 64, Number 7.
Berger, J. O., (1985), *Statistical Decision Theory and Bayesian Analysis* (2nd Ed.), Springer-Verlag, New York, NY.
Bishop, C. M. & Bishop, H., (2024), *Deep Learning: Foundations and Concepts*, Springer, Cham, Switzerland.
Boole, G., (1854), *An Investigation of the Laws of Thought*, Walton & Maberly, London, U.K.
Bostrom, N., (2014), *Superintelligence: Paths, Dangers, Strategies*, Oxford University Press, Oxford, U.K.
Bostrom, N., (2024), *Deep Utopia: Life and Meaning in a Solved World*, Ideapress Publishing, Washington, D.C.
Box, G. E. P. & Draper, N. R., (1987), *Empirical Model-Building and Response Surfaces*, John Wiley & Sons, New York, NY.
Bryson, A. E. & Ho, Y-C., (1969), *Applied Optimal Control*, Blaisdell.
ByteDance Website: www.bytedance.com
Čapek, K., (2004), *R.U.R. (Rossum's Universal Robots)*, Penguin Books, New York, NY, Originally published by Aventinum in 1921 in Prague, Czech Republic.
Carroll, L., (1865), *Alice's Adventures in Wonderland*, Macmillan & Company, London, U.K.
Chollet, F. & Watson, M., *Deep Learning with Python* (3rd Ed.), Manning Publications, Shelter Island, New York, NY.
Clarke, A. C., (2018), *2001: A Space Odyssey*, ACE, Hudson, NY, Originally published in 1968 by New American Library.
Discovery Loop Website: www.discoveryloop.com
Drucker, P. F., (1994), *The Theory of the Business*, Harvard Business Review, September-October 1994, Harvard Business Review Press, Boston, MA.

- Fukushima, K., (1980), *Neocognitron: A Self-Organizing Neural Network Model for a Mechanism of Pattern Recognition Unaffected by Shift in Position*, Biological Cybernetics, 36, PP. 193-202.
- Goodfellow, I., Bengio, Y., & Courville, A., (2016), *Deep Learning*, The MIT Press, Cambridge, MA.
- Google Website: www.google.com
- Govindarajan, V. & Venkatraman, V., (2024), *Fusion Strategy: How Real-Time Data and AI Will Power the Industrial Future*, Harvard Business Review Press, Boston, MA.
- Hannas, W. C. & Chang, H., (2023), *Chinese Power and Artificial Intelligence: Perspectives and Challenges*, Routledge, London, U.K.
- Hao, K., (2025), *Empire of AI: Dreams and Nightmares in Sam Altman's OpenAI*, Penguin Press, New York, NY.
- Harville, D. A., (2018), *Linear Models and the Relevant Distributions and Matrix Algebra*, CRC Press, Boca Raton, FL.
- Haykin, S., (2009), *Neural Networks and Learning Machines* (3rd Ed.), Pearson Education, Hoboken, NJ.
- Hebb, D. O., (1949), *The Organization of Behavior: A Neuropsychological Theory*, John Wiley & Sons, New York, NY.
- Huawei Website: www.huawei.com
- IBM Website: www.ibm.com
- Ineffable AI Website: www.ineffable.ai
- Isaacson, W., (2023), *Elon Musk*, Simon & Schuster, New York, NY.
- Jordan, M. I., (2025), *A Collectivist, Economic Perspective on AI*, arXiv:2507.06268v3.
- Karp, A. C. & Zamiska, N. W., (2025), *The Technological Republic: Hard Power, Soft Belief, And the Future of the West*, Crown Currency, New York, NY.
- Kasparov, G. & Greengard, M., (2017), *Deep Thinking: Where Machine Intelligence Ends and Human Creativity Begins*, PublicAffairs, New York, NY.
- Kaplan, W., (1984), *Advanced Calculus* (3rd Ed.), Addison-Wesley, Reading, MA.
- Kim, J. et al., (2026), *Inducing Language Models to Assert Their Own Consciousness Restores Human Beliefs and Values*, Published on arXiv:2607.28607v1, July 30, 2026.
- Kissinger, H. A., Mundie, C., & Schmidt, E., (2024), *Genesis: Artificial Intelligence, Hope, and the Human Spirit*, Little, Brown and Company, New York, NY.
- Krzanowski, W. J., (2000), *Principles of Multivariate Analysis: A User's Perspective* (Revised Edition), Oxford University Press, New York, NY.
- Krizhevsky, A., Sutskever, I., & Hinton, G., (2012), *ImageNet Classification With Deep Convolutional Neural Networks*, Proceedings – Advances in Neural Information Processing Systems, 25, PP. 1090-1098.
- Kurzweil, R., (2005), *The Singularity is Near: When Humans Transcend Biology*, Penguin Books, New York, NY.
- Kurzweil, R., (2024), *The Singularity is Nearer: When We Merge with AI*, Viking, New York, NY.

- LeCun, Y., *et al.*, (1990), *Handwritten Digit Recognition with a Back-Propagation Network*, Proceedings - Advances in Neural Information Processing Systems, PP. 396-404.
- LeCun, Y., Bottou, L., Bengio, Y., & Haffner, P., (1998), *Gradient-Based Learning Applied to Document Recognition*, Proceedings – IEEE, 86, PP. 2278-2324.
- LeCun, Y., Bengio, Y., & Hinton, G., (2015), *Deep Learning*, Nature, Volume 521, May 28, 2015.
- Lee, K. F., (2018), *AI Superpowers: China, Silicon Valley, and the New World Order*, Houghton Mifflin Harcourt, Boston, MA.
- Lee, K. F. & Qiufan, C., (2021), *AI 2041: Ten Visions for Our Future*, Currency, New York, NY.
- Liedtke, C. A., (2025), *Artificial Intelligence for Strategic Improvement*, Strategic Improvement Systems, LLC, Minnetrista, MN. Website. www.strategicimprovementsystems.com/research/.
- Lindgren, B. W., (1993), *Statistical Theory* (4th Ed.), Chapman & Hall/CRC, Boca Raton, FL.
- Lu, C., Lue, C., Lange, R. T., Forester, J., Clune, J., & Ha, D., (2024), *The AI Scientist: Towards Fully Automated Open-Ended Scientific Discovery*, Sakana AI, Tokyo, Japan.
- Mallaby, S., (2026), *The Infinity Machine: Demis Hassabis, DeepMind, and the Quest for Superintelligence*, Penguin Press, New York, NY.
- Marr, D., (2010), *Vision: A Computational Investigation into the Human Representation and Processing of Visual Information*, The MIT Press, Cambridge, MA.
- Marr, B., (2025), *AI Strategy: Unleash the Power of Artificial Intelligence in Your Business*, Kogan Page Limited, London, U.K.
- Marr, B., (2019), *Artificial Intelligence in Practice: How 50 Successful Companies Used Big Data Analytics to Deliver Extraordinary Results*, John Wiley & Sons, West Sussex, U.K.
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E., (1955), *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*.
- McCulloch, W. S. & Pitts, W., (1943), *A Logical Calculus of the Ideas Immanent in Nervous Activity*, Bulletin of Mathematical Biophysics, Vol. 5, PP. 115-133.
- Meta Platforms Website: www.meta.com
- Metz, C., (2021), *Genius Makers: The Mavericks Who Brought AI to Google, Facebook, and the World*, Dutton, New York, NY.
- Microsoft Website: www.microsoft.com
- Minsky, M. L., (1954), *Theory of Neural-Analog Reinforcement Systems and Its Application to The Brain-Model Problem*, Doctoral Dissertation, Princeton University.
- Minsky, M. L. & Papert, S. A., (1969), *Perceptrons: An Introduction to Computational Geometry* (Expanded Edition), The MIT Press, Boston, MA.
- Moonshot AI Website: www.moonshot.ai
- Newell, A. & Simon, H. A., (1956), *Current Developments in Complex Information Processing*, The Rand Corporation, Santa Monica, CA.
- Newell, A. & Simon, H. A., (1959), *Report on a General Problem-Solving Program*, The Rand Corporation, Santa Monica, CA.
- Nilsson, N., (2010), *The Quest for Artificial Intelligence: A History of Ideas and Achievements*, Cambridge University Press, New York, NY.

- Nobel Foundation, (2024), www.nobelprize.org/all-nobel-prizes-2024/.
- NVIDIA Website: www.nvidia.com, www.blogs.nvidia.com/blog/open-secure-ai-alliance/.
- Olson, P., (2024), *Supremacy: AI, ChatGPT, and the Race That Will Change the World*, St. Martin's Press, New York, NY.
- OpenAI Website: www.openai.com
- Orwell, G., (1949), *Nineteen Eighty-Four*, Secker and Warburg, London, U.K.
- Palantir Website: www.palantir.com
- Petro, T., (2025), *AI in the Boardroom: Preparing Leaders for Responsible Governance*, Business Expert Press, New York, NY.
- Pickover, C. A., (2024), *Artificial Intelligence: An Illustrated History* (Revised Edition), Union Square & Company, New York, NY.
- Pope Leo XIV, (2026), Encyclical Letter – *Magnifica Humanitas: On Safeguarding the Human Person in the Time of Artificial Intelligence*, The Vatican, Published May 5, 2026.
- Prince, S., (2023), *Understanding Deep Learning*, The MIT Press, Cambridge, MA.
- Prometheus Website: www.prometheus.ai
- Raschka, S., (2025), *Build a Large Language Model (From Scratch)*, Manning Publications, Shelter Island, NY.
- Rosenblatt, F., (1958), *The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain*, Psychological Review, Volume 65, Number 6.
- Rumelhart, D. E. & McClelland, J. L., (1986), *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Volume 1: Foundations*, The MIT Press, Cambridge, MA.
- Rumelhart, D. E. & McClelland, J. L., (1986), *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Volume 2: Psychological and Biological Models*, The MIT Press, Cambridge, MA.
- Russell, S. & Norvig, P., (2022), *Artificial Intelligence: A Modern Approach* (4th Edition, Global Edition), Pearson Education Limited, Harlow, U.K.
- Sadler, M. & Regan, N., (2019), *Game Changer: AlphaZero's Groundbreaking Chess Strategies and the Promise of AI*, New In Chess, Alkmaar, The Netherlands.
- Safe Superintelligence Website: www.ssi.inc
- Searle, S. R., (1982), *Matrix Algebra Useful for Statistics*, John Wiley & Sons, New York, NY.
- Sejnowski, T. J., (2018), *The Deep Learning Revolution*, The MIT Press, Cambridge, MA.
- Sejnowski, T. J., (2024), *ChatGPT and the Future of AI: The Deep Language Revolution*, The MIT Press, Cambridge, MA.
- Shannon, C. E., (1948), *A Mathematical Theory of Communication*, The Bell System Technical Journal, Volume 27, PP. 379-423, 623-656, July, October, 1948.
- Shelley, M., (1995), *Frankenstein*, The Easton Press, Norwalk, CT, Originally published in 1818.
- Solomonoff, G., (2006), *Ray Solomonoff and the Dartmouth Summer Research Project in Artificial Intelligence*, 1956, Oxbridge Research, Cambridge, MA. It appears in the *IEEE Annals of the History of Computing*, Volume 28, Issue 3, Pages 34–41.
- SpaceX Website: www.spacex.com

- Stanford Institute for Human-Centered Artificial Intelligence, (2026), *Artificial Intelligence Index Report 2026*, Stanford University, Palo Alto, CA. Website: www.hai.stanford.edu.
- Sutskever, I., Vinyals, O., & Le, Q. V., (2014), *Sequence to Sequence Learning with Neural Networks*, Proceedings – Advances in Neural Information Processing Systems, 27, 3104-3112.
- Sutton, R., (1988), *Learning to Predict by the Methods of Temporal Differences*, Machine Learning 3: 9-44, 1988, Kluwer Academic Publishers, Boston, MA.
- Sutton, R. S. & Barto, A. G., (2020), *Reinforcement Learning: An Introduction* (2nd Ed.), The MIT Press, Cambridge, MA.
- Tencent Website: www.tencent.com
- Tesla Website: www.tesla.com
- Thinking Machines Lab Website: www.thinkingmachines.ai
- Torralba, A., Isola, P., & Freeman, W. T., (2024), *Foundations of Computer Vision*, The MIT Press, Cambridge, MA.
- TSMC Website: www.tsmc.com
- Turing, A. M., (1937), *On Computable Numbers, with an Application to the Entscheidungsproblem*, Proceedings of the London Mathematical Society, Series 2, Volume 42, Issue 1, 1936-1937, PP. 230-265.
- Turing, A. M., (1950), *Computing Machinery and Intelligence*, Mind 49: Pages 433-460.
- Unitree Robotics Website: www.unitree.com
- Vaswami, A. et al., (2017), *Attention Is All You Need*, 31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA.
- Wachter, R., (2026), *A Giant Leap: How AI is Transforming Healthcare and What That Means for Our Future*, Portfolio | Penguin, New York, NY.
- Wei, J. et al., (2023), *Chain-of-Thought Prompting Elicits Reasoning in Large Language Models*, arXiv:2201.11903v6.
- Wiener, N., (1948), *Cybernetics or the Control and Communication in the Animal and the Machine*, The M.I.T. Press, Cambridge, MA.
- Wolfram, S., (2023), *What is ChatGPT Doing . . . and Why Does It Work*, Wolfram Media, Wolfram-media.com.
- Wooldridge, M., (2020), *A Brief History of Artificial Intelligence: What It Is, Where We Are, and Where We Are Going*, Flatiron Books, New York, NY.
- World Labs Website: www.worldlabs.ai
- Z.ai Website: www.z.ai
- Zao-Sanders, M., (2026), *How People are Really Using AI in 2026*, HBR.org, June 1, 2026, Harvard Business Review.
- Zhao, W. X. et al., (2026), *A Survey of Large Language Models*, arXiv:2303.18223v19, 3/18/26.

References: Categorized

Historical AI Works

- Bayes, T., (1764), *An Essay Towards Solving a Problem in the Doctrine of Chances*, Philosophical Transactions of the Royal Society of London for 1763, 53: 370-418.
- Bengio, Y., LeCun, Y., & Hinton, G., (2021), *Deep Learning for AI*, Turing Lecture, Communication of the ACM, July 2021, Volume 64, Number 7.
- Boole, G., (1854), *An Investigation of the Laws of Thought*, Walton & Maberly, London, U.K.
- Bryson, A. E. & Ho, Y-C., (1969), *Applied Optimal Control*, Blaisdell.
- Fukushima, K., (1980), *Neocognitron: A Self-Organizing Neural Network Model for a Mechanism of Pattern Recognition Unaffected by Shift in Position*, Biological Cybernetics, 36, PP. 193-202.
- Hebb, D. O., (1949), *The Organization of Behavior: A Neuropsychological Theory*, John Wiley & Sons, New York, NY.
- Krizhevsky, A., Sutskever, I., & Hinton, G., (2012), *ImageNet Classification With Deep Convolutional Neural Networks*, Proceedings – Advances in Neural Information Processing Systems, 25, PP. 1090-1098.
- Kurzweil, R., (2005), *The Singularity is Near: When Humans Transcend Biology*, Penguin Books, New York, NY.
- Kurzweil, R., (2024), *The Singularity is Nearer: When We Merge with AI*, Viking, New York, NY.
- LeCun, Y., et al., (1990), *Handwritten Digit Recognition with a Back-Propagation Network*, Proceedings - Advances in Neural Information Processing Systems, PP. 396-404.
- LeCun, Y., Bottou, L., Bengio, Y., & Haffner, P., (1998), *Gradient-Based Learning Applied to Document Recognition*, Proceedings – IEEE, 86, PP. 2278-2324.
- LeCun, Y., Bengio, Y., & Hinton, G., (2015), *Deep Learning*, Nature, Volume 521, May 28, 2015.
- Marr, D., (2010), *Vision: A Computational Investigation into the Human Representation and Processing of Visual Information*, The MIT Press, Cambridge, MA.
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E., (1955), *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*.
- McCulloch, W. S. & Pitts, W., (1943), *A Logical Calculus of the Ideas Immanent in Nervous Activity*, Bulletin of Mathematical Biophysics, Vol. 5, PP. 115-133.
- Minsky, M. L., (1954), *Theory of Neural-Analog Reinforcement Systems and Its Application to The Brain-Model Problem*, Doctoral Dissertation, Princeton University.
- Minsky, M. L. & Papert, S. A., (1969), *Perceptrons: An Introduction to Computational Geometry* (Expanded Edition), The MIT Press, Boston, MA.
- Newell, A. & Simon, H. A., (1956), *Current Developments in Complex Information Processing*, The Rand Corporation, Santa Monica, CA.
- Newell, A. & Simon, H. A., (1959), *Report on a General Problem-Solving Program*, The Rand Corporation, Santa Monica, CA.
- Rosenblatt, F., (1958), *The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain*, Psychological Review, Volume 65, Number 6.

- Rumelhart, D. E. & McClelland, J. L., (1986), *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Volume 1: Foundations*, The MIT Press, Cambridge, MA.
- Rumelhart, D. E. & McClelland, J. L., (1986), *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Volume 2: Psychological and Biological Models*, The MIT Press, Cambridge, MA.
- Shannon, C. E., (1948), *A Mathematical Theory of Communication*, The Bell System Technical Journal, Volume 27, PP. 379-423, 623-656, July, October, 1948.
- Sutskever, I., Vinyals, O., & Le, Q. V., (2014), *Sequence to Sequence Learning with Neural Networks*, Proceedings – Advances in Neural Information Processing Systems, 27, 3104-3112.
- Sutton, R., (1988), *Learning to Predict by the Methods of Temporal Differences*, Machine Learning 3: 9-44, 1988, Kluwer Academic Publishers, Boston, MA.
- Turing, A. M., (1937), *On Computable Numbers, with an Application to the Entscheidungsproblem*, Proceedings of the London Mathematical Society, Series 2, Volume 42, Issue 1, 1936-1937, PP. 230-265.
- Turing, A. M., (1950), *Computing Machinery and Intelligence*, Mind 49: Pages 433-460.
- Vaswami, A. et al., (2017), *Attention Is All You Need*, 31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA.
- Wei, J. et al., (2023), *Chain-of-Thought Prompting Elicits Reasoning in Large Language Models*, arXiv:2201.11903v6.
- Wiener, N., (1948), *Cybernetics or the Control and Communication in the Animal and the Machine*, The M.I.T. Press, Cambridge, MA.

History of AI

- Isaacson, W., (2023), *Elon Musk*, Simon & Schuster, New York, NY.
- Kasparov, G. & Greengard, M., (2017), *Deep Thinking: Where Machine Intelligence Ends and Human Creativity Begins*, PublicAffairs, New York, NY.
- Metz, C., (2021), *Genius Makers: The Mavericks Who Brought AI to Google, Facebook, and the World*, Dutton, New York, NY.
- Nilsson, N., (2010), *The Quest for Artificial Intelligence: A History of Ideas and Achievements*, Cambridge University Press, New York, NY.
- Pickover, C. A., (2024), *Artificial Intelligence: An Illustrated History* (Revised Edition), Union Square & Company, New York, NY.
- Russell, S. & Norvig, P., (2022), *Artificial Intelligence: A Modern Approach* (4th Edition, Global Edition), Pearson Education Limited, Harlow, U.K.
- Sadler, M. & Regan, N., (2019), *Game Changer: AlphaZero's Groundbreaking Chess Strategies and the Promise of AI*, New In Chess, Alkmaar, The Netherlands.
- Sejnowski, T. J., (2018), *The Deep Learning Revolution*, The MIT Press, Cambridge, MA.
- Solomonoff, G., (2006), *Ray Solomonoff and the Dartmouth Summer Research Project in Artificial Intelligence*, 1956, Oxbridge Research, Cambridge, MA. It appears in the *IEEE Annals of the History of Computing*, Volume 28, Issue 3, Pages 34–41.
- Wooldridge, M., (2020), *A Brief History of Artificial Intelligence: What It Is, Where We Are, and Where We Are Going*, Flatiron Books, New York, NY.

Deep Learning Textbooks

- Bishop, C. M. & Bishop, H., (2024), *Deep Learning: Foundations and Concepts*, Springer, Cham, Switzerland.
- Goodfellow, I., Bengio, Y., & Courville, A., (2016), *Deep Learning*, The MIT Press, Cambridge, MA.
- Haykin, S., (2009), *Neural Networks and Learning Machines* (3rd Ed.), Pearson Education, Hoboken, NJ.
- Prince, S., (2023), *Understanding Deep Learning*, The MIT Press, Cambridge, MA.
- Sutton, R. S. & Barto, A. G., (2020), *Reinforcement Learning: An Introduction* (2nd Ed.), The MIT Press, Cambridge, MA.
- Torralba, A., Isola, P., & Freeman, W. T., (2024), *Foundations of Computer Vision*, The MIT Press, Cambridge, MA.

AI Resources

- Arena AI: www.arena.ai
- arXiv Website: arXiv.org
- Nobel Foundation, (2024), www.nobelprize.org/all-nobel-prizes-2024/.
- Stanford Institute for Human-Centered Artificial Intelligence, (2026), *Artificial Intelligence Index Report 2026*, Stanford University, Palo Alto, CA. Website: www.hai.stanford.edu.

Leading AI

- Drucker, P. F., (1994), *The Theory of the Business*, Harvard Business Review, September-October 1994, Harvard Business Review Press, Boston, MA.
- Liedtke, C. A., (2025), *Artificial Intelligence for Strategic Improvement*, Strategic Improvement Systems, LLC, Minnetrista, MN. Website: www.strategicimprovementsystems.com/research/.
- Marr, B., (2025), *AI Strategy: Unleash the Power of Artificial Intelligence in Your Business*, Kogan Page Limited, London, U.K.
- Petro, T., (2025), *AI in the Boardroom: Preparing Leaders for Responsible Governance*, Business Expert Press, New York, NY.

New AI Entities

- Advanced Machine Intelligence (AMI) Labs Website: www.amilabs.xyz
- Discovery Loop Website: www.discoveryloop.com
- Ineffable AI Website: www.ineffable.ai
- Prometheus Website: www.prometheus.ai
- Safe Superintelligence Website: www.ssi.inc
- Thinking Machines Lab Website: www.thinkingmachines.ai
- World Labs Website: www.worldlabs.ai

Companies | Organizations

Alibaba Website: www.alibaba.com
Amazon Website: www.amazon.com
Anthropic Website: www.website.anthropic.com
Apple Website: www.apple.com
Baidu Website: www.baidu.com
ByteDance Website: www.bytedance.com
Google Website: www.google.com
Huawei Website: www.huawei.com
IBM Website: www.ibm.com
Meta Platforms Website: www.meta.com
Microsoft Website: www.microsoft.com
Moonshot AI Website: www.moonshot.ai
NVIDIA Website: www.nvidia.com, www.blogs.nvidia.com/blog/open-secure-ai-alliance/.
OpenAI Website: www.openai.com
Palantir Website: www.palantir.com
SpaceX Website: www.spacex.com
Tencent Website: www.tencent.com
Tesla Website: www.tesla.com
TSMC Website: www.tsmc.com
Unitree Robotics Website: www.unitree.com
Z.ai Website: www.z.ai

Statistics | Technical

Berger, J. O., (1985), *Statistical Decision Theory and Bayesian Analysis* (2nd Ed.), Springer-Verlag, New York, NY.
Box, G. E. P. & Draper, N. R, (1987), *Empirical Model-Building and Response Surfaces*, John Wiley & Sons, New York, NY.
Chollet, F. & Watson, M., *Deep Learning with Python* (3rd Ed.), Manning Publications, Shelter Island, New York, NY.
Harville, D. A., (2018), *Linear Models and the Relevant Distributions and Matrix Algebra*, CRC Press, Boca Raton, FL.
Kaplan, W., (1984), *Advanced Calculus* (3rd Ed.), Addison-Wesley, Reading, MA.
Krzanowski, W. J., (2000), *Principles of Multivariate Analysis: A User's Perspective* (Revised Edition), Oxford University Press, New York, NY.
Lindgren, B. W., (1993), *Statistical Theory* (4th Ed.), Chapman & Hall/CRC, Boca Raton, FL.
Raschka, S., (2025), *Build a Large Language Model (From Scratch)*, Manning Publications, Shelter Island, NY.
Searle, S. R., (1982), *Matrix Algebra Useful for Statistics*, John Wiley & Sons, New York, NY.

General AI

- Bostrom, N., (2014), *Superintelligence: Paths, Dangers, Strategies*, Oxford University Press, Oxford, U.K.
- Bostrom, N., (2024), *Deep Utopia: Life and Meaning in a Solved World*, Ideapress Publishing, Washington, D.C.
- Govindarajan, V. & Venkatraman, V., (2024), *Fusion Strategy: How Real-Time Data and AI Will Power the Industrial Future*, Harvard Business Review Press, Boston, MA.
- Hannas, W. C. & Chang, H., (2023), *Chinese Power and Artificial Intelligence: Perspectives and Challenges*, Routledge, London, U.K.
- Hao, K., (2025), *Empire of AI: Dreams and Nightmares in Sam Altman's OpenAI*, Penguin Press, New York, NY.
- Jordan, M. I., (2025), *A Collectivist, Economic Perspective on AI*, arXiv:2507.06268v3.
- Karp, A. C. & Zamiska, N. W., (2025), *The Technological Republic: Hard Power, Soft Belief, And the Future of the West*, Crown Currency, New York, NY.
- Kim, J. et al., (2026), *Inducing Language Models to Assert Their Own Consciousness Restores Human Beliefs and Values*, Published on arXiv:2607.28607v1, July 30, 2026.
- Kissinger, H. A., Mundie, C., & Schmidt, E., (2024), *Genesis: Artificial Intelligence, Hope, and the Human Spirit*, Little, Brown and Company, New York, NY.
- Lee, K. F., (2018), *AI Superpowers: China, Silicon Valley, and the New World Order*, Houghton Mifflin Harcourt, Boston, MA.
- Lee, K. F. & Qiufan, C., (2021), *AI 2041: Ten Visions for Our Future*, Currency, New York, NY.
- Lu, C., Lue, C., Lange, R. T., Forester, J., Clune, J., & Ha, D., (2024), *The AI Scientist: Towards Fully Automated Open-Ended Scientific Discovery*, Sakana AI, Tokyo, Japan.
- Mallaby, S., (2026), *The Infinity Machine: Demis Hassabis, DeepMind, and the Quest for Superintelligence*, Penguin Press, New York, NY.
- Marr, B., (2019), *Artificial Intelligence in Practice: How 50 Successful Companies Used Big Data Analytics to Deliver Extraordinary Results*, John Wiley & Sons, West Sussex, U.K.
- Olson, P., (2024), *Supremacy: AI, ChatGPT, and the Race That Will Change the World*, St. Martin's Press, New York, NY.
- Pope Leo XIV, (2026), *Encyclical Letter – Magnifica Humanitas: On Safeguarding the Human Person in the Time of Artificial Intelligence*, The Vatican, Published May 5, 2026.
- Sejnowski, T. J., (2024), *ChatGPT and the Future of AI: The Deep Language Revolution*, The MIT Press, Cambridge, MA.
- Wachter, R., (2026), *A Giant Leap: How AI is Transforming Healthcare and What That Means for Our Future*, Portfolio | Penguin, New York, NY.
- Wolfram, S., (2023), *What is ChatGPT Doing . . . and Why Does It Work*, Wolfram Media, Wolfram-media.com.
- Zao-Sanders, M., (2026), *How People are Really Using AI in 2026*, HBR.org, June 1, 2026, Harvard Business Review.
- Zhao, W. X. et al., (2026), *A Survey of Large Language Models*, arXiv:2303.18223v19, 3/18/26.

Fiction

Asimov, I., (1942), *I, Robot*, In *Runaround*, Spectra Books, New York, NY.

Asimov, I., (1950), *I, ROBOT*, Gnome Press, New York, NY.

Banks, I., (1987), *Consider Phlebas*, Macmillan, London, UK.

Čapek, K., (2004), *R.U.R. (Rossum's Universal Robots)*, Penguin Books, New York, NY,

Originally published by Aventinum in 1921 in Prague, Czech Republic.

Carroll, L., (1865), *Alice's Adventures in Wonderland*, Macmillan & Company, London, U.K.

Clarke, A. C., (2018), *2001: A Space Odyssey*, ACE, Hudson, NY, Originally published in 1968
by New American Library.

Orwell, G., (1949), *Nineteen Eighty-Four*, Secker and Warburg, London, U.K.

Shelley, M., (1995), *Frankenstein*, The Easton Press, Norwalk, CT, Originally published in 1818.

Author's Biographical Information



Charles A. Liedtke, Ph.D. is the Founder and President of Strategic Improvement Systems, LLC, a management consulting company designed help leaders improve organizational performance. Charles conducts research, consults, and provides customized training on *Strategy, Culture, Quality, Analytics, AI, Improvement, and Innovation*. He has worked with organizations worldwide including Fortune 500 companies, privately-held companies, non-profit organizations, government entities, and medical centers. His most recent research projects were on *Quality, Analytics, and Big Data; Big Data in Hoshin Kanri; Information-Based Customer Value Creation; Advances in Strategic Planning; Shaping Organizational Culture; Visions & Visioning; Advances in Horizontal Interaction; Artificial Intelligence for Strategic Improvement; and AI's Past, Present, & Future*.

Charles earned a Ph.D. in Operations and Information Management from the University of Wisconsin-Madison; an M.B.A. and Ph.D. Minor in Statistics also from UW-Madison; an M.S. Degree in Statistics from Iowa State University; and a B.S. Degree in Economics from South Dakota State University.

Contact Information:

Charles A. Liedtke, Ph.D., President
Strategic Improvement Systems, LLC (SIS)
Email: charles@sisliedtke.com
Website: www.strategicimprovementsystems.com